

AD-A119 250

SOUTHEASTERN CENTER FOR ELECTRICAL ENGINEERING EDUCAT--ETC F/6 12/1
A COMPARISON OF THE PENCIL-OF-FUNCTION METHOD WITH PRONY'S METH--ETC(U)
DEC 77 T K SARKAR, V K JAIN, J NEBAT F33615-77-C-2059

UNCLASSIFIED

NL

1 of 1
AD A
119250

END
DATE
FILMED

10-82
DTIC

2

AD A119240

A COMPARISON OF THE PENCIL-OF-FUNCTION METHOD WITH PRONY'S
METHOD, WIENER FILTERS AND OTHER IDENTIFICATION TECHNIQUES

By

Tapan K. Sarkar¹

Vijay K. Jain²

Joshua Nebat³

Donald D. Weiner³

Submitted to
Air Force Weapons Laboratory
Contract : F 33615-77-C-2059
Task : 12

December 1977

SEP 14 1982

E

DTIC FILE COPY

¹Dept. of Electrical Eng.
Rochester Institute of Tech.
Rochester, N.Y. 14623

²Dept. of Electrical Eng.
University of South Florida
Tampa, Florida 33620

³Dept. of Electrical
& Computer Eng.
Syracuse University
Syracuse, NY 13210

¹Presently on leave at
Gordon McKay Laboratory
Harvard University
Cambridge, Mass. 02138

82 09 13 105

Information
for public
distribution is indicated

PREFACE

The authors wish to thank Dr. C.E. Baum, Dr. J.P. Castillo, Capt. M. Harrison, Capt. H.G. Hudson, Capt. E. Fuller and Mr. William Prather of the Air Force Weapons Laboratory for a keen interest in all aspects of this work.

The contributors to this report are Dr. T.K. Sarkar, Dr. D.D. Weiner, Mr. J. Nebat and Dr. V.K. Jain. Dr. Sarkar is presently with Harvard University, Dr. Weiner and Mr. Nebat are with Syracuse University, and Dr. Jain is with the University of South Florida.

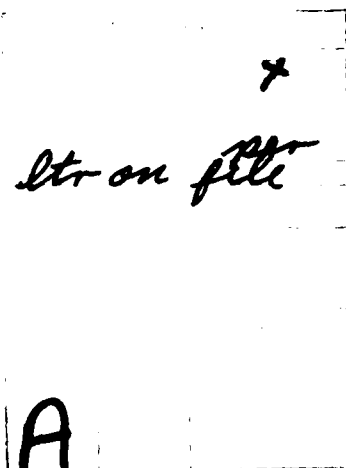


TABLE OF CONTENTS

<u>Section</u>	<u>Page</u>
I. INTRODUCTION	5
II. WIENER FILTER THEORY	8
2.1 Inverse Filter Formulation	14
2.2 Linear Prediction	15
2.3 Predictive Deconvolution or Spiking Filter Design	17
2.4 Recursive Filter Design	19
2.5 Kalman Filter Theory	23
III. STOCHASTIC METHODS APPLIED TO SYSTEM IDENTIFICATION.	26
3.1 Maximum Likelihood Estimation Theory	26
3.2 Minimum Predictor Error Variance	29
3.3 Maximum Entropy Spectral Analysis	30
IV. PRONY'S ALGORITHM AND ITS EXTENSIONS	36
4.1 Various Extensions to Prony's Method	37
V. ILL-POSED AND WELL-POSED PROBLEMS IN SYSTEM IDENTIFICATION	41
VI. PENCIL-OF-FUNCTIONS METHOD	51
VII. DISCUSSION	57
REFERENCES	58

LIST OF FIGURES

<u>Figure</u>		<u>Page</u>
1	Principles of Wiener Filtering	9
2	Transient Response of a Conducting Pipe Measured at the ATHAMAS - I EMP Simulator	54
3	True Response Vs. Reconstructed Response of a Sixth Order System for a 10 m long 1 m Diameter Conducting Pipe	56

I. INTRODUCTION

This paper deals with the identification/approximation of a linear system by poles and residues from a measured finite length input-output record of the system. The objective of this paper is threefold:

- 1) to illustrate that several different formulations for characterizing the impulse response of a system yield the same set of poles;
- 2) to show how different formulations regularize the ill-posed system identification problem, and
- 3) to demonstrate that relatively stable, consistent and reliable results in the identification of a system by poles and residues from a finite length input-output record can be achieved by the pencil-of-functions method.

Recognition that different formulations yield the same poles is not widely appreciated. In this paper it is shown that formulations based upon different assumptions result in identical sets of analysis equations. For example, it is not at all obvious that Prony's method (as derived by most numerical analysis formulations) may be interpreted in terms of predicting each value of data and

therefore, is a form of digital Wiener filter. Markel [1] and Markel and Gray [2] did recognize the fact that identical analysis equations can be obtained by several different techniques.

In our discussion of the various approaches, only references which are directly relevant are noted. No attempt has been made to cite the earliest sources. In many cases, additional references may be found in papers mentioned.

The problem of interest is to identify/approximate the transfer function of a system by its poles and residues when the noise contaminated input and output are specified. The signal and noise are considered to be stationary processes. When time limited signals are involved, the problem is converted to an equivalent stationary problem by convolving the time limited signal with white noise of unit power [3]. In the second section the classical method for extracting the signal from noise is discussed. This is the Wiener-Kolmogoroff theory [4]. The digital form of the Wiener-Hopf equation is derived. Topics associated with Wiener filters are also presented. They include inverse filter design, linear prediction, predictive deconvolution (or spiking) filter design, recursive filter design and Kalman filtering.

- 1 J.D. Markel, "Formant Trajectory Estimation from a Linear Least-Square Inverse Filter Formulation," SCRL-Monograph 7, Speech Communications Research Laboratory, Inc., Santa Barbara, October 1971.
- 2 J.D. Markel and A.H. Gray, "Linear Prediction of Speech," Springer-Verlag: Berlin.
- 3 E.A. Robinson and S. Traital, "Principles of Digital Wiener Filtering," Geophysical Prospecting, September 1967, pp. 311-333.
- 4 N. Levinson, "The Wiener RMS Error Criterion in Filter Design and Prediction," Journal of Mathematics and Physics, 1947 V. 25, pp. 261-278.

Both the popularly known covariance and autocorrelation methods are derived from the Wiener filter theory.

In the third section the various well-posed stochastic extensions of an ill-posed system identification/approximation problem are described. They include the maximum likelihood estimation theory, the minimum predictor error variance and the maximum entropy spectral analysis. It is demonstrated that identical analysis equations for parametric modeling of the system can be obtained.

The fourth section provides Prony's method in various forms. In particular, when a semi-least squares approach is applied to Prony's method, both the autocorrelation and the covariance method appear as special cases. Thus it is also a form of a digital Wiener filter.

The second objective of this paper is discussed in the fifth section. This section discusses the various concepts of ill-posed and well-posed problems in system identification. It is shown how the different techniques regularize the ill-posed system identification problem by introducing further limitations on the solution. Finally, it is shown how the pencil-of-functions method radically differs from the other formulations.

Finally, the third objective is demonstrated in section VI where results are presented to demonstrate the claim that relatively stable, reliable and consistent results are obtained for the location of the poles by the pencil-of-function method.

II. WIENER FILTER THEORY

Kolmogoroff (1942) and Wiener (1943) were the first to present a unified theory on extrapolation, interpolation and smoothing of stationary time series. The linear filter which performs the desired task is obtained by the solution of an integral equation known as the Wiener-Hopf equation [4]. For sampled data systems, the integral form of the Wiener-Hopf equation reduces to a finite sum. The present treatment describes how Wiener's concepts can be applied to the identification/approximation of linear systems. The basic model for this process consists of an input signal, a desired output signal and an actual output signal. If one minimizes the mean-squared error between the desired output signal and the actual output signal, it becomes possible to solve for the optimum system commonly known as the "Wiener" filter. The fundamental assumption underlying the procedure is that all processes are stationary.

A stationary time series is one whose statistical properties are time invariant. In particular, the statistics of the time series are independent of time. By definition, a stationary time series must be of infinite duration. However, in an actual experiment, we observe a times series over a finite interval. In order to apply the concepts of Wiener filtering, the finite length time series is convolved with a white noise series of unit power, to yield a stationary time series [3]. Moreover, in actual measurements only one waveform is often available

for computing various statistics. Thus ensemble averages are frequently evaluated by means of appropriate time averages. This is valid only when the random processes are ergodic.

The concept of Wiener filtering is well known but is included here for completeness. The fundamental elements of digital Wiener filtering are summarized in Figure 1 for the sampled data problem.

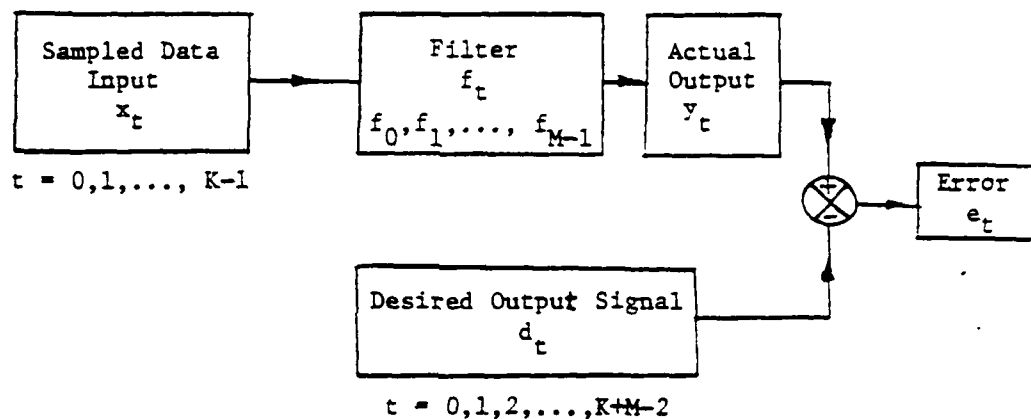


Figure 1. Principles of Wiener Filtering.

Given the input sequence x_t and a desired output sequence d_t , the problem is to find the linear filter coefficients f_t , whose output sequence $x_t * f_t$ [* denotes convolution] yields a minimum mean-squared error estimate of d_t . If $E\{\cdot\}$ denotes the expected value, then the error

$$I = E\{(d_t - y_t)^2\} \quad (2.1)$$

is to be minimized.

In this case an M -length filter $f_t = \{f_0, f_1, \dots, f_{M-1}\}$ converts, in a least error energy, sense a K -length input $x_t = \{x_0, x_1, \dots, x_{K-1}\}$

into a $K+M-1$ length sequence y_t which approximates the desired sequence $d_t = \{d_0, d_1, \dots, d_{K+M-2}\}$. The actual output is obtained as

$$\begin{aligned} y_t &= \{y_0, y_1, \dots, y_{K+M-2}\} \\ &= x_t * f_t = \sum_{\tau=0}^{M-1} f_{\tau} x_{t-\tau}. \end{aligned} \quad (2.2)$$

The problem of making the actual output y_t as close as possible to the desired output d_t can be interpreted in terms of minimizing the error energy

$$I = E\{(d_t - y_t)^2\} = E\left\{(d_t - \sum_{\tau=0}^{M-1} f_{\tau} x_{t-\tau})^2\right\}. \quad (2.3)$$

The error is minimized by evaluating the partial derivatives of I with respect to f_{τ} and equating them to zero. This results in a set of equations

$$\begin{aligned} \frac{dI}{df_j} &= E\left\{2(d_t - \sum_{\tau=0}^{M-1} f_{\tau} x_{t-\tau})(-x_{t-j})\right\} \\ &= -2E\{d_t x_{t-j}\} + 2 \sum_{\tau=0}^{M-1} f_{\tau} E\{x_{t-\tau} x_{t-j}\} \\ &= 0, \quad \text{for } j = 0, 1, 2, \dots, M-1 \end{aligned} \quad (2.4)$$

The unknown filter coefficients are obtained by solving the following set of simultaneous equations.

$$\begin{aligned} \sum_{\tau=0}^{M-1} f_{\tau} E\{x_{t-\tau} x_{t-j}\} &= E\{d_t x_{t-j}\} \\ \text{for } j &= 0, 1, \dots, M-1 \end{aligned} \quad (2.5)$$

In order to solve the above equations, it is necessary to compute the expected values in Equation (2.5). By assuming that the ensemble

averages can be evaluated in terms of time averages, one obtains [3]

$$E\{x_{t-\tau}x_{t-j}\} \stackrel{\Delta}{=} \frac{1}{K-M} \sum_{t=M}^{K-1} x_{t-\tau}x_{t-j} \stackrel{\Delta}{=} C_{\tau,j} \quad (2.6)$$

(in the covariance method)

$$\begin{aligned} &\stackrel{\Delta}{=} \frac{1}{K} \sum_{t=0}^{K-1+M} x_{t-\tau}x_{t-j} \\ &= \frac{1}{K} \sum_{t=0}^{K-|\tau-j|-1} x_t x_{t+|\tau-j|} \stackrel{\Delta}{=} r_{(\tau-j)} \end{aligned} \quad (2.7)$$

(in the autocorrelation method)

Similarly

$$E\{d_t x_{t-j}\} \stackrel{\Delta}{=} \frac{1}{K-M} \sum_{t=M}^{K-1} d_t x_{t-j}$$

(in the covariance method)

$$\stackrel{\Delta}{=} \frac{1}{K} \sum_{t=0}^{K-1+M} d_t x_{t-j} \quad (2.9)$$

(in the autocorrelation method)

It is important to stress that the terms "covariance" and "autocorrelation" are not based upon the standard usage of the terms as occur in the theory of stochastic processes. Rather, we follow the usage which is quite prevalent in the speech processing literature [5].

The following discussion is intended to clarify their interpretation.

It is clear that the covariance definition given by (2.6) yield an unbiased estimate since

5 J. Makhoul, "Linear Prediction: A Tutorial Review," Proc. IEEE, Vol. 63, No. 4, April 1975, pp. 561-580.

$$\begin{aligned}
E\{E(x_{t-\tau}x_{t-j})\} &\stackrel{\Delta}{=} E\left\{\frac{1}{K-M} \sum_{t=M}^{K-1} x_{t-\tau}x_{t-j}\right\} \\
&= \frac{1}{K-M} \sum_{t=M}^{K-1} E\{x_{t-\tau}x_{t-j}\} = E\{x_{t-\tau}x_{t-j}\}.
\end{aligned}$$

On the other hand, the autocorrelation definition given by (2.7) results in a biased estimate since

$$\begin{aligned}
E\{E(x_{t-\tau}x_{t-j})\} &\stackrel{\Delta}{=} E\left\{\frac{1}{K} \sum_{t=0}^{K-1+M} x_{t-\tau}x_{t-j}\right\} \\
&= \frac{1}{K} \sum_{t=0}^{K-|\tau-j|-1} E\{x_t x_{t+|\tau-j|}\} = \left[1 - \frac{|\tau-j|}{K}\right] E\{x_{t-\tau}x_{t-j}\}.
\end{aligned}$$

Since $(\tau-j) \leq M$, it is interesting to note that the bias is negligible when M (the order of the filter) is much less than K . The bias is significant, however, when K is only slightly larger than M . This explains why a large number of data points ($K \gg M$) is necessary for unbiased spectral estimation when using the autocorrelation function to obtain the power spectral density.

Because the covariance method gives an unbiased estimate, it might be assumed that it is a better estimate. However, the biased estimate provided by the autocorrelation method is often preferable. As an example consider the zero mean four data point sequence (4, -2, -1, -1) and $M = 2$. Using the unbiased estimator, the expected values are $C'_{00} = 1$, $C'_{10} = 1.5$, $C'_{20} = -1$, $C'_{30} = -2$. In contrast, the biased estimator results in $r_{(0)} = 5.5$, $r_{(1)} = -1.25$, $r_{(2)} = -0.5$, $r_{(3)} = -1.0$. Note that C'_{00} is less than C'_{10} whereas $r_{(0)}$ is guaranteed to be greater

than $r_{(\tau)}$ for $\tau > 0$.

Continuing with our development, the discrete Wiener-Hopf equation presented in (2.5) can be written in the following matrix form for the covariance method

$$[C_{ij}]_{M \times M} [F_i]_{M \times 1} = [D_j]_{M \times 1} \quad (2.10)$$

where $[C_{ij}]$ is a square matrix whose elements are given as

$$C_{ij} = \sum_{t=M}^{K-1} x_{t-i} x_{t-j}, \quad (2.11)$$

$[F_i]$ is a column matrix consisting of the M unknown filter coefficients $f_0, f_1, \dots, f_{M-1}$, and $[D_j]$ is a column matrix whose elements are given by

$$D_j = \sum_{t=M}^{K-1} d_t x_{t-j} \quad (2.12)$$

For the autocorrelation method, the unknown filter coefficients are obtained from the solution of the matrix equation

$$[R_{|i-j|}]_{M \times M} [F_i]_{M \times 1} = [D_j]_{M \times 1} \quad (2.13)$$

where $[R_{|i-j|}]$ is a square matrix whose elements are given as

$$R_{|i-j|} = \sum_{t=0}^{K-1+M} x_{t-i} x_{t-j} = \sum_{t=0}^{K-1-|i-j|} x_t x_{t+|i-j|} \quad (2.14)$$

and $[D_j]$ is a column matrix defined by (12).

Interestingly, most of the formulations for the solution of an unknown linear filter lead to analysis equations which can be formulated either in terms of the autocorrelation matrix equations (2.13) or in terms of the covariance matrix equations (2.10).

It is important to note that the Wiener filter is not always realizable as a causal rational function (in terms of poles and zeros). As a result, the Wiener filter is, in general, an infinite order filter. When the order of the filter is specified a priori, the resulting filter may no longer be optimum.

Various forms of the Wiener filter have appeared under different names and have been used in various geophysical, speech processing, and digital signal processing applications. Next, various modifications of the Wiener filter are presented.

2.1. Inverse Filter Formulation

The inverse filter attempts to transform the input signal into an impulse [6]. Assume that the input sequence x_t is transformed to an impulse of area σ by an all-zero filter of the form

$$F(z) = \sum_{i=0}^{M-1} f_i z^{-i}, \quad \text{with } f_0 = 1 \quad (2.15)$$

In terms of Figure 1, the desired waveform d_t is an impulse of area σ . It follows that the coefficients of the filter should be chosen such that

$$\sum_{\tau=0}^{M-1} f_{\tau} x_{t-\tau} = \sigma \delta_t \quad (2.16)$$

where $\delta_t = 1$ for $t = 0$ and zero otherwise. Multiplication of both

6 J.D. Markel, "Digital Inverse Filtering - A New Tool for Formant Trajectory Estimation," IEEE Trans. on Audio and Electroacoustics, Vol. AU-18, No. 2, June 1970, pp. 137-141.

sides by x_{t-j} and summing t from k to $K-1+l$, one obtains

$$\sum_{\tau=0}^{M-1} f_{\tau} \left[\sum_{t=k}^{K-1+l} x_{t-\tau} x_{t-j} \right] = \begin{cases} 0, & k > 0 \\ \sigma x_{-j}, & k = 0 \end{cases} \quad (2.17)$$

Since x_{-j} is zero for $j > 0$, the unknown coefficients f_{τ} for $\tau = 1, 2, \dots, M-1$ are obtained from the solution of the following equations

$$\sum_{\tau=0}^{M-1} f_{\tau} \left[\sum_{t=k}^{K-1+l} x_{t-\tau} x_{t-j} \right] = 0 \quad (2.18)$$

for $j = 1, 2, \dots, M-1$.

For $k = M$ and $l = 0$, the above equations reduce to that of the covariance equations as in (2.10). This is because

$$D_j = \sum_{t=M}^{K-1} \delta_t x_{t-j} = 0. \quad (2.19)$$

For $k = 0$ and $l = M$, equation (2.18) reduces to that of the autocorrelation equations. If the input signal is approximated by an all-pole model the poles of the input signal are obtained from the zeros of $F(z)$ i.e. from the solution of the polynomial equations

$$\sum_{\tau=0}^{M-1} f_{\tau} z^{-\tau} = 0. \quad (2.20)$$

In particular, if $(z^{-1})_i$ is the i th root of the above equation (2.20), then the i th pole is equal to $\ln[(z^{-1})_i]$.

2.2 Linear Prediction

The term "linear Prediction" was first used by Wiener in his classic work on prediction of stationary time series. Since its publication, it has found wide application in the determination of all-pole

models for the processing of speech signals [5].

The basic philosophy here is to take a part of the sampled waveform (say the first $M-1$ points from K data points) and predict the next data point on the waveform by proper choice of the predictor coefficients a_i . The linear predictor of step size one predicts the M th data point of the waveform when a $(M-1)$ order predictor filter is chosen. In the time domain, the predicted sample $\hat{x}_M$ is given by

$$\hat{x}_M = - \sum_{i=1}^{M-1} a_i x_{M-i}$$

where $(-a_1, -a_2, -a_3, \dots, -a_{M-1})$ are the predictor coefficients. Assuming the signal spectrum is to be modelled by an all-pole model, the coefficients a_i are the negative of the values of f_i presented in (2.15). A $(M-1)$ order linear predictor thus requires a linear combination of the previous $(M-1)$ samples. The error is then given by

$$\begin{aligned} e_M &= \hat{x}_M + x_M = - \sum_{i=1}^{M-1} a_i x_{M-i} + x_M \\ &= - \sum_{i=0}^{M-1} a_i x_{M-i} \quad \text{with } a_0 = -1 \end{aligned} \quad (2.21)$$

When $(K-l+l)$ different data points are predicted, the total squared error is defined by

$$\begin{aligned} I &= \sum_{t=k}^{K-l+l} [e_t]^2 = \sum_{t=k}^{K-l+l} \left[\sum_{i=0}^{M-1} a_i x_{t-i} \right]^2 \\ &= \sum_{i=0}^{M-1} \sum_{j=0}^{M-1} a_i a_j \left\{ \sum_{t=k}^{K-l+l} x_{t-i} x_{t-j} \right\} \\ &= \sum_{i=0}^{M-1} \sum_{j=0}^{M-1} f_i f_j \left\{ \sum_{t=k}^{K-l+l} x_{t-i} x_{t-j} \right\}. \end{aligned} \quad (2.22)$$

Minimization of I , with respect to the set of the filter coefficients leads to a set of simultaneous equations given by

$$\sum_{i=0}^{M-1} f_i \left(\sum_{t=k}^{K-1+l} x_{t-i} x_{t-j} \right) = 0 \quad (2.23)$$

for $j = 1, 2, \dots, M - 1$, from which the unknown filter coefficients f_i are obtained. When $k = M$ and $l = 0$, this amounts to minimization of the error only over data points of the waveform from M to $K - 1$. When $k = 0$ and $l = M$, this implies minimization over the entire waveform. Equations (2.23) are identical to the set of equations (2.19) obtained in the case of the inverse filter formulation in the previous section.

Linear prediction is thus equivalent to an all-pole model for the spectrum of the input x_t . The poles for this model are again obtained from the solution of the polynomial equation

$$\sum_{i=0}^{M-1} f_i z^{-i} = 0.$$

Thus if x_t represents the measured impulse response of the system, the set of poles obtained by linear prediction can be interpreted so as to parameterize the system in terms of an all-pole model.

2.3 Predictive Deconvolution or Spiking Filter Design

The general linear filtering problem involves the input x_t , the impulse response h_t and the output y_t . When it is desirable to evaluate h_t given x_t and y_t , the problem is referred to as deconvolution. In this sense the inverse filter problem discussed in section 2.1 is a deconvolution problem. Predictive deconvolution refers to the case in which the output y_t is assumed to be a delayed impulse

such that

$$\sum_{\tau=0}^{M-1} f_{\tau} x_{t-\tau} = \sigma \delta_{t-\alpha+1} \quad (2.24)$$

The unknown filter coefficients f_{τ} can be pursued in a manner analogous to the inverse filter approach by assuming the input to be represented as an all-pole model. However, we prefer to show that the filter coefficients can also be determined by interpreting the problem as a prediction problem. It is in this sense that the term predictive deconvolution is used [7].

Introduce the change of variables $t - \alpha + 1 = s$ in equation (2.24). The resulting equation is

$$\sum_{\tau=0}^{M-1} f_{\tau} x_{s+\alpha-\tau-1} = \sigma \delta_s \quad (2.25)$$

Next, assume the filter coefficients to be given by

$$\overbrace{(-1, 0, 0, \dots, 0)}^{\alpha-1 \text{ zeros}}, -a_1, -a_2, \dots, -a_{M-1} \quad (2.26)$$

The upper limit on the summation is now given by $M + \alpha - 2$. Equation (2.25) can now be written as

$$x_{t+\alpha-1} + \sum_{\tau=1}^{M+\alpha-2} f_{\tau} x_{t+\alpha-\tau-1} = x_{t+\alpha-1} - \sum_{\tau=1}^{M-1} a_{\tau} x_{t-\tau} = \sigma \delta_t \quad (2.27)$$

If the estimate of $x_{t+\alpha-1}$ is assumed to be given by $\hat{x}_{t+\alpha-1}$, then

$$\hat{x}_{t+\alpha-1} = - \sum_{\tau=1}^{M-1} a_{\tau} x_{t-\tau} \quad (2.28)$$

and $\sigma \delta_t$ can be interpreted as the error of the estimate.

7 K.L. Peacock and S. Treital, "Predictive Deconvolution: Theory and Practice," Geophysics, Vol. 34, No. 2, April 1969, pp. 135-169.

If the total squared error is minimized as was done for the linear prediction approach in section 2.2, a set of equations is obtained for the filter coefficients. When $\alpha = 1$, these equations reduce to the same equations as were obtained in (2.23)

2.4 Recursive Filter Design

Recursive filters are often used in digital filter design [8-12]. With respect to Fig. 1, the spectral estimation problem is now posed in the following manner. Assume the filter input is the impulse δ_t . Let the filter impulse response be determined so as to approximate the data sequence in a minimum mean squared error sense. If x_t is the input to the filter and y_t is the output, then they are related by

$$\sum_{k=0}^{M-1} a_k y_{t-k} = \sum_{r=0}^{L-1} b_r x_{t-r}$$

Application of the z-transform to both sides yields

-
- 8 A. Oppenheim and R. Shaefer, Digital Signal Processing. Englewood Cliffs, NJ: Prentice Hall, 1975.
 - 9 J.L. Shanks, "Recursion Filter for Digital Processing," Geophysics, Vol. XXXII, No. 1, February 1967, pp. 33-51.
 - 10 D.L. Fletcher and C.N. Weygandt, "A Digital Method of Transfer Function Calculation," IEEE Trans. on Circuit Theory, Jan. 1971, pp. 185-187.
 - 11 C.S. Burrus and T.W. Parks, "Time Domain Design of Recursive Digital Filters," IEEE Trans. on Audio and Electroacoustics, Vol. AU-18, No. 2, June 1970, pp. 137-141.
 - 12 S. Treital and E.A. Robinson, "The Design of High Resolution Digital Filters," IEEE Trans. on Geoscience Electronics, Vol. GE-4, No. 1, June 1966, pp. 25-38.

$$\sum_{k=0}^{M-1} a_k z^{-k} Y(z) = \sum_{r=0}^{L-1} b_r z^{-r} X(z)$$

or

$$F(z) = \frac{Y(z)}{X(z)} = \frac{\sum_{r=0}^{L-1} b_r z^{-r}}{\sum_{k=0}^{M-1} a_k z^{-k}} \equiv \frac{B(z)}{A(z)}. \quad (2.29)$$

Therefore, a recursive filter has a transfer function which is expressed as a ratio of two polynomials of the z-transform variable. The objective here is to synthesize the filter from the given impulse response of the system. It is important to note that, unlike the previous three types of filters, this is not an all-pole filter but a pole-zero model. In other words, given some desired filter operator $D(z)$ (which is the transfer function of a desired system having the impulse response d_t), we require the polynomials $A(z)$ and $B(z)$ of the filter

$$F(z) = \frac{B(z)}{A(z)} \quad (2.30)$$

such that

$$F(z) \approx D(z) = d_0 + d_1 z^{-1} + d_2 z^{-2} + \dots + d_{K-1} z^{-K+1}. \quad (2.31)$$

A technique for determining $A(z)$ and $B(z)$ is outlined next.

$$A(z) = 1 + a_1 z^{-1} + a_2 z^{-2} + \dots + a_{M-1} z^{-M+1} \quad (2.32)$$

and

$$B(z) = b_0 + b_1 z^{-1} + b_2 z^{-2} + \dots + b_{L-1} z^{-L+1} \quad (2.33)$$

where M and L are arbitrary numbers which fix the number of poles and zeros respectively for the filter. Also divide $B(z)$ by $A(z)$ so as to obtain the infinite series

$$F(z) = f_0 + f_1 z^{-1} + f_2 z^{-2} + \dots$$

Since $F(z) A(z) = B(z)$, it follows that

$$F(z) \{ 1 + a_1 z^{-1} + a_2 z^{-2} + \dots + a_{M-1} z^{-M+1} \}$$

$$= b_0 + b_1 z^{-1} + b_2 z^{-2} + \dots + b_{L-1} z^{-L+1}.$$

Since multiplication of z -transforms is equivalent to convolution of the discrete time series, the series of b_t coefficients is equal to the convolution of the f_t coefficients with the a_t coefficients. Or equivalently,

$$b_t = \sum_{j=0}^{M-1} a_j f_{t-j}.$$

By assumption, $a_0 = 1$. Therefore,

$$f_t = b_t - \sum_{j=1}^{M-1} a_j f_{t-j}.$$

As $b_t = 0$ for $t > L$ (from equation (2.33)), one may write

$$f_t = - \sum_{j=1}^{M-1} a_j f_{t-j}, \quad \text{for } t > L. \quad (2.34)$$

By assuming the approximation in (2.31) is valid, the coefficients f_t approximate the coefficients d_t for $t > L$. Equation (2.34) may then be approximated by

$$d_t \approx - \sum_{j=1}^{M-1} a_j d_{t-j}$$

for $t = L + 1, L + 2, \dots, K - 1$ (2.35)

The error in (2.35) is given by

$$e_t = d_t + \sum_{j=1}^{M-1} a_j d_{t-j}$$

$$= \sum_{j=0}^{M-1} a_j d_{t-j}; \text{ since } a_0 = 1$$

for $t = L + 1, L + 2, \dots, K - 1$ (2.36)

The a_j coefficients are chosen in such a way that the mean-squared error is minimized. In particular,

$$I = \sum_{t=L+1}^{K-1} e_t^2 = \sum_{t=L+1}^{K-1} \left[\sum_{j=0}^{M-1} a_j d_{t-j} \right]^2 \quad (2.37)$$

is minimized. By differentiating I of equation (2.37) with respect to a_j and equating the derivatives to zero, a set of simultaneous equations is obtained. They are given by

$$\sum_{j=0}^{M-1} a_j \left[\sum_{t=L+1}^{K-1} d_{t-j} d_{t-k} \right] = 0 \quad (2.38)$$

for $k = 1, 2, \dots, M-1$.

For a realizable filter the numerator polynomial is generally one degree lower than the denominator polynomial (if the poles are simple). Thus

$$L + 1 = M.$$

Hence, (2.38) reduces to

$$\sum_{j=0}^{M-1} a_j \left[\sum_{t=M}^{K-1} d_{t-j} d_{t-k} \right] = 0 \quad (2.39)$$

which is equivalent to the covariance equations as given by (2.10). The right side of (2.10) is zero because

$$D_j = \sum_{t=M}^{K-1} \delta_{t-j} = 0 \quad \text{for } j = 1, 2, \dots, M-1.$$

The poles for the filter are then obtained by the solution of the polynomial equation

$$\sum_{i=0}^{M-1} a_i z^{-i} = 0.$$

It was shown previously that the inverse filter, assuming an all-pole model for the signal spectrum, is also equivalent to (2.10). We conclude therefore that the poles for an all-pole model correspond to the identical set of poles for a pole zero model for the

same order filter.

Next the residues at the poles (or equivalently, the numerator polynomial) can be obtained by minimizing the mean-squared error given by $\sum_t \{f_t - d_t\}^2$.

It is interesting to note that when the numerator polynomial is realized directly in the form presented in (2.33), the problem reduces to the case of the Padé approximation [13]. Mathematically, Padé approximation results in an approximation of $D(z)$ by $F(z)$ such that the seminorm

$$\begin{aligned} \|D(z) - F(z)\| &= |D(1) - F(1)| + |D^1(1) - F^1(1)| \\ &+ \dots + |D^{L+M-1}(1) - F^{L+M-1}(1)| \quad (2.40) \end{aligned}$$

is made zero. Here $D^k(1)$ represent the k th derivative of $D(z)$ evaluated on the unit circle.

2.5 Kalman Filter Theory

Underlying Wiener filter design is the so-called Wiener-Hopf integral equation, its solution through spectral factorization, and the practical problem of synthesizing the theoretically optimal filter from its impulse response. The normal Wiener filter is derived from the Wiener-Hopf equation and in general this equation can be solved only in the steady state, i.e. when the observation interval is semi-infinite. The contribution of Kalman was recognition of the fact that the integral

13 R.N. McDonough, "Representation and Analysis of Signals, Part XV - Matched Exponents for the Representation of Signals," Johns Hopkins University, April 1963.

equation could be converted into a nonlinear differential equation whose solution contains all the necessary information for design of the optimal filter. The problem of spectral factorization in the Wiener filter is analogous to the requirement for solving $M(M+1)/2$ coupled nonlinear algebraic equations in the M -order Kalman filter. These equations can be solved numerically for transient type problems, where data is available only for a finite interval. This, in general, results in the Kalman filters being time-variant. However, in the steady-state the Kalman filter reduces to the time invariant Wiener filter [14]. The presentation by Sorensen [15] expresses the results of Kalman filter theory in a way that makes this comparison easier.

The problem involves estimating a signal $\{s_n\}$, from measured data $\{d_0, d_1, \dots, d_{K-1}\}$. If the estimate is computed as a linear combination of the d_n , then

$$\hat{s}_n = \sum_{i=0}^{M-1} A_i d_i. \quad (2.41)$$

The M coefficients A_i are chosen in such a way that the mean-squared error,

$$I = E[(s_n - \hat{s}_n)^T (s_n - \hat{s}_n)], \quad (2.42)$$

is minimized. Here T denotes the transpose of the row vector $(s_n - \hat{s}_n)$. This criterion is satisfied when the error in the estimate s_n is orthogonal to the measured data, or

$$E[(s_n - \hat{s}_n) d_i^T] = 0, \quad \text{for } i = 0, 1, \dots, M-1 \quad (2.43)$$

14 A. Gelb, "Applied Optimal Estimation," MIT Press, 1974.

15 H.W. Sorensen, "Least-squares Estimation from Gauss to Kalman," IEEE Spectrum, July 1970, pp. 63-68.

Expressed in a different way

$$E[s_n d_i^T] = \sum_{i=0}^{M-1} A_i E[d_i d_i^T] \quad \text{for } i = 0, 1, \dots, M-1. (2.44)$$

However, the matrix inversion that is required becomes computationally impractical when M is large. Wiener and Kolmogoroff assumed an infinite amount of data (that is, the lower limit of the summation is $-\infty$ rather than zero). Eq (2.43) is then the discrete form of the Wiener-Hopf equation which was solved using spectral factorization.

The basic difference between Wiener-Kolmogoroff theory and the Kalman filter theory is how equation (2.44) is solved. In 1955 J.W.Follin suggested a recursive approach to solve (2.44). It is clear (see reference 15. p. 65) that Follin's work provided a direct stimulus for the work of Richard Bucy, which led to his subsequent collaboration with Kalman in the total development of the "state space" approach for obtaining the filter equations.

2.6 Summary

As outlined above all forms of the digital Wiener filter lead either to the covariance or the autocorrelation equations. It is also interesting that the same set of poles is obtained whether one models the signal as an all-pole model or as a pole-zero model.

III. STOCHASTIC METHODS APPLIED TO SYSTEM IDENTIFICATION

Three different stochastic methods used in spectral estimation are presented in this section. They are maximum likelihood estimation, minimum predictor error variance estimation, and maximum entropy spectral analysis. All three models are considered as all-pole models. These methods have no relationship with discrete Wiener filtering theory. The methods presented in this section start with completely different assumptions but finally yield the identical set of either covariance or autocorrelation equations which characterize the system to be identified from the measured impulse response.

In these methods it is assumed that the data samples are part of a random process. The problem is to choose the parameters of the system impulse response so as to make the probability of occurrence of the actual observation most likely. In other words, the system parameters are chosen in such a way that the probability density function defining the parameters is maximized.

3.1 Maximum Likelihood Estimation Theory

In this approach the measured impulse response of the system is considered as a segment of a random process. It is further assumed that the impulse response can be generated by passing an uncorrelated noise sequence $\{e_n\}$ through an all-pole model of the form

$$\frac{1}{F(z)} = \frac{1}{\sum_{i=0}^{M-1} f_i z^{-i}} \quad \text{with } f_0 = 1 \quad (3.1)$$

Next the random process $\{e_t\}$ is assumed to be Gaussian with zero mean and variance σ_e^2 . Thus,

$$E\{e_t\} = 0 \quad \text{and} \quad E\{e_i e_j\} = \delta_{ij} \sigma_e^2 \quad (3.2)$$

As the measured impulse response x_t , $t = 0, 1, \dots, K-1$ has been assumed to be generated by passing the noise sequence $\{e_t\}$ through the all-pole model, it follows that

$$\sum_{i=0}^{M-1} f_i x_{t-i} = e_t. \quad (3.3)$$

From (3.2) and (3.3) it is clear that the sequence x_t is Gaussian with zero mean and a cross-correlation defined by

$$E\{x_i x_j\} = g_{i-j}. \quad (3.4)$$

This correlation sequence g_{i-j} would then be a function of the system parameters f_i , $i = 0, 1, \dots, M-1$ and σ_e^2 . Since x_t is Gaussian, a Gaussian multivariate probability density function is defined for the sequence of random variables $x_0, x_1, \dots, x_{K-1}$. The maximum likelihood theory assumes that the parameter values which make the measured observation of the impulse response most likely are the same values which maximize the joint probability density function of x_i , $i = 0, 1, \dots, K-1$. This can be achieved by differentiating the density function with each of the unknown variables, $f_1, f_2, \dots, f_{M-1}$ and σ_e^2 and then setting the first partial derivative equal to zero. The

solutions of the set of equations then yield the values for the unknown parameters. Even though the procedure is conceptually simple, the set of equations becomes extremely nonlinear for M greater than 2 and no exact solution for this problem exists [2].

However, Itakura and Saito [16,17] solved the maximum likelihood problem by making some additional assumptions. First, the number of data points K is made much greater than M , the order of the filter (i.e., $K \gg M$). Second, the joint probability density function for the sequence $x_0, x_1, x_2, \dots, x_{K-1}$ is approximated by

$$p(x_0, x_1, \dots, x_{K-1}) = [2\pi\sigma_e^2]^{-K/2} \exp[-\alpha/2\sigma_e^2] \quad (3.5)$$

where

$$\alpha = \sum_{t=0}^{K-1+M} \left[\sum_{i=0}^{M-1} f_i x_{t-i} \right]^2 \quad (3.6)$$

It is interesting to note that α has the identical form of the error defined by the autocorrelation equations in linear prediction (see equation (2.22)).

It has been shown [16, 17] that the results obtained for the unknown filter coefficients f_i are identical to (2.13) which utilizes the autocorrelation equations.

The corresponding equations for the covariance method are obtained by defining a conditional density function for the probability density. This is achieved by treating the M data points $x_0, x_1, \dots, x_{M-1}$ as a set of deterministic initial conditions and the remaining

16 F. Itakura and S. Saito, "Analysis Synthesis Telephony based on Maximum Likelihood Method," 6th International Congress on Acoustics, Tokyo, Japan, Aug. 21-28, 1968, C-5-5, pp. C-17-20.

17 F. Itakura and S. Saito, "Extraction of Speech Parameters based upon the Statistical Method," Proc. Speech Info. Process, Tohoku University, Sendai, Japan, 5.1, 5.12 (1971) (in Japanese).

K-M data points as a set of random variables. Under the above assumptions, the conditional probability density function is approximated as [17, 2]

$$p_c\{(x_M, x_{M+1}, \dots, x_{K-1}) | (x_0, x_1, \dots, x_{M-1})\} = (2\pi\sigma_e^2)^{-0.5(K-M)} \exp[-\alpha/2\sigma_e^2] \quad (3.7)$$

where

$$\alpha = \sum_{t=M}^{K-1} \left[\sum_{i=0}^{M-1} f_i x_{t-i} \right]^2 \quad (3.8)$$

Again it is clear that α has the same form as the error energy defined for the covariance equations for the case of linear prediction (see (2.22)). Maximization of the conditional probability density function is then achieved by maximizing p_c with respect to the unknown filter coefficients f_i and σ_e^2 . The set of equations (see (2.10)) obtained in this case are identical to those of the covariance method.

3.2 Minimum Predictor Error Variance

In this method the data samples are not considered as a part of a Gaussian process. In other words, the method remains the same as before, i.e. the measured impulse response is generated by passing a noise sequence $\{e_t\}$ through an all-pole model [2]. It is assumed that the error sequence is of zero mean and of variance given by

$$E[e_t^2] = \sum_{i=0}^M \sum_{j=0}^M f_i f_j E[x_{t-i} x_{t-j}] \quad (3.9)$$

Next the process is assumed to be stationary so that the expectation in (3.9) can be expressed as

$$E[x_{t-i}x_{t-j}] = g_{i-j} \quad (3.10)$$

and the error variance as

$$E[e_t^2] = \sum_{i=0}^M \sum_{j=0}^M f_i f_j g_{i-j} \quad (3.11)$$

The problem is to determine the filter coefficients so as to minimize the error variance.

An additional assumption is now made. Specifically, it is assumed that the process is ergodic so that the ensemble average E may be converted to a time average. Hence, the approximation

$$\begin{aligned} g_{i-j} &\triangleq \frac{1}{K-M} \sum_{t=M}^{K-1} x_{t-i}x_{t-j} \\ &= c'_{ij} \end{aligned} \quad (3.12)$$

leads to the covariance equations (2.10) and the approximation

$$\begin{aligned} g_{i-j} &\triangleq \frac{1}{K} \sum_{t=0}^{K-|i-j|-1} x_t x_{t+|i-j|} \\ &= r_{i-j} \end{aligned} \quad (3.13)$$

leads to the autocorrelation equations (2.13). Hence, this method yields a set of analysis equations identical to those for the discrete Wiener filter.

3.3 Maximum Entropy Spectral Analysis

An important aspect of time series analysis is the computation of the power spectral density which is primarily determined by the second order statistics. In an actual experiment, the number of

data points is always finite. Hence, for the problem of interest the data length may not be sufficient to obtain a specified degree of frequency resolution. Also, given a finite number of K data points, we can obtain at most approximations of the K autocorrelation functions $r_0, r_1, \dots, r_{K-1}$. In the previous autocorrelation formulations the data has been assumed to be zero outside the known interval. In some instances, this may be an unreasonable assumption about the extension of the data beyond the known interval. The question then arises as to what assumptions should be made about the data outside the finite sample and what assumption should be made about their second order statistics (i.e. the autocorrelation), since they determine the power spectral density.

Burg proposed an information theory approach to the problem. He suggested [18] that the most reasonable choice of the unknown autocorrelations is the one which adds no information or adds most randomness or maximizes the entropy. He then proceeded to select the power spectral density having the maximum entropy of all possible spectra that agrees with the known values of the autocorrelation function r_1 .

The information content of a random process is defined in terms of a quantity called entropy and is mathematically expressed as

$$H = - \sum_j P_j \ln P_j \quad (3.14)$$

where P_j is the probability of the j th event of a random process.

18 J.P. Burg, "Maximum Entropy Spectral Analysis," Ph.D. Thesis, Stanford University, Palo Alto, California 1975.

When the random variable takes on a continuum of values, the sum in the definition of the entropy is replaced by an integral. Since we are dealing with a time series $x_0, x_1, \dots, x_{K-1}$, the probability is replaced by the joint probability density function $p(x_0, x_1, \dots, x_{K-1})$. Thus

$$H = - \int p(x_0, x_1, \dots, x_{K-1}) \ln \{p(x_0, x_1, \dots, x_{K-1})\} dV \quad (3.15)$$

where dV is an element of volume in the space spanned by the random variables. Burg then proceeded to adjoin a hypothetical variable x_K to the available estimates of the autocorrelation function r_0, r_1, r_2 and so on. We may then consider the joint probability density available for the K data points and the adjoined x_K as

$$p(x_0, x_1, \dots, x_{K-1}, x_K) \quad (3.16)$$

This probability density function has an entropy

$$H = - \int p(x_0, x_1, \dots, x_{K-1}, x_K) \ln \{p(x_0, x_1, \dots, x_{K-1}, x_K)\} dV \quad (3.17)$$

Burg chose as (3.16) that probability density function which has its first K second order moments as $r_0, r_1, \dots, r_{K-1}$, and which under the given constraint maximized (3.17). The obvious choice for the probability density function in (3.16) is Gaussian since according to Shannon and Weaver [19, 20] the Gaussian distribution results in maximum entropy under a constant energy constraint. Thus

-
- 19 C.E. Shannon and W. Weaver, The Mathematical Theory of Communication. Urbana, Illinois: University of Illinois Press, 1962, pp. 56-57.
 - 20 R.N. McDonough, "Maximum-entropy spatial processing of array data," Geophysics, Vol. 39, No. 6, December 1974, pp. 843-851.

$$p(x_0, x_1, \dots, x_{K-1}, x_K) = \frac{\exp\{-\frac{1}{2} X' [R_K]^{-1} X\}}{\{(2\pi)^{\frac{K+1}{2}} \det[R_K]\}^{1/2}} \quad (3.18)$$

where X is the column vector of the x_i , the prime indicates the transpose, and the matrix $[R_K]$ is given by

$$[R_K] = \begin{bmatrix} r_0 & r_1 & \dots & r_{K-1} & r_K \\ r_1 & r_0 & & r_{K-2} & r_{K-1} \\ \vdots & \vdots & & \vdots & \vdots \\ r_{K-1} & & & & \\ r_K & \dots & \dots & \dots & r_0 \end{bmatrix} \quad (3.19)$$

The entropy can then be expressed as [19, 20]

$$H = \frac{1}{2} \ln\{(2\pi e)^{K+1} \det[R_K]\} \quad (3.20)$$

Now r_K is to be chosen in such a way that H in (3.20) is maximized.

Hence the value of r_K is the one which maximizes $\det [R_K]$.

In order for r_i to constitute a proper set of autocorrelation values, the matrix $[R_K]$ must be positive semi-definite [21]. Moreover, $\det [R_K]$ is a quadratic function in r_K . It follows that maximizing $\det [R_K]$ with respect to r_K yields the value of r_K obtained from the solution of the following equation

$$\det \begin{vmatrix} r_1 & r_0 & \dots & r_{K-2} \\ r_2 & r_1 & & r_{K-3} \\ \vdots & & & \vdots \\ r_K & r_{K-1} & \dots & r_1 \end{vmatrix} = 0 \quad (3.21)$$

21 A. Van den Bos, "Alternative Interpretation of Maximum Entropy Spectral Analysis," IEEE Trans. on Information Theory, Vol. IT-17, No. 4, 1971, pp. 493-494.

A Discussion of Various Approaches to the Identification/Approximation Problem

TAPAN K. SARKAR, SENIOR MEMBER, IEEE, DONALD D. WEINER, MEMBER, IEEE,
JOSHUA NEBAT, AND VIJAY K. JAIN, SENIOR MEMBER, IEEE

Abstract—The identification/approximation of a linear system by poles and residues from a measured finite length input-output record of the system is discussed. The objective of this paper is to illustrate that several different formulations for characterizing the impulse response of a system yield the same set of poles.

I. INTRODUCTION

IN RECENT years the singularity expansion method (SEM) has been used very successfully to study transient phenomena in electromagnetic radiation and scattering problems. With this approach, information is obtained about the electromagnetic structures from measured transient responses to known inputs. The information leads to a characterization of the impulse response of the electromagnetic system by a sum of damped exponentials. It is desirable to know the complex natural frequencies or poles with a high degree of accuracy. The problem of extraction of the poles from the measured transient response data is reduced to a system approximation/identification problem.

This paper deals with the identification/approximation of a linear system by poles and residues from a measured finite length input-output record of the system. The objective of this paper is to illustrate that several different formulations for characterizing the impulse response of a system yield the same set of poles.

Recognition that different formulations yield the same poles is not widely appreciated. It is shown here that formulations based upon different assumptions result in identical sets of analysis equations. For example, it is not at all obvious that Prony's method (as derived by most numerical analysis formulations) may be interpreted in terms of predicting each value of data and, therefore, is a form of digital Wiener filter. Markel [2] did recognize the fact that identical analysis equations can be obtained by several different techniques.

In our discussion of the various approaches, only references which are directly relevant are noted. No attempt has been made to cite the earliest sources. In many cases, additional references may be found in papers mentioned.

The problem of interest is to identify/approximate the transfer function of a system by its poles and residues when the noise-contaminated input and output are specified. The signal and noise are considered to be stationary processes. When time-limited signals are involved, the problem is con-

verted to an equivalent stationary problem by convolving the time limited signal with white noise of unit power [3]. In the second section the classical method for extracting the signal from noise is discussed. This is the Wiener-Kolmogoroff theory [4]. The digital form of the Wiener-Hopf equation is derived. Topics associated with Wiener filters are also presented. They include inverse filter design, linear prediction, predictive deconvolution (or spiking) filter design, recursive filter design, and Kalman filtering. Both the popularly known covariance and autocorrelation methods are derived from the Wiener filter theory.

In the third section the various stochastic extensions are described. They include the maximum likelihood estimation theory, the minimum predictor error variance and the maximum entropy spectral analysis. It is demonstrated that identical analysis equations for parametric modeling of the system can be obtained.

The fourth section provides Prony's method in various forms. In particular, when a semi-least-squares approach is applied to Prony's method, both the autocorrelation and the covariance method appear as special cases. Thus, it is also a form of a digital Wiener filter.

II. WIENER FILTER THEORY

Kolmogoroff in 1942 and Wiener in 1943 were the first to present a unified theory on extrapolation, interpolation, and smoothing of stationary time series. The linear filter which performs the desired task is obtained by the solution of an integral equation known as the Wiener-Hopf equation [4]. For sampled data systems, the integral form of the Wiener-Hopf equations reduces to a finite sum. The present treatment describes how Wiener's concepts can be applied to the identification/approximation of linear systems. The basic model for this process consists of an input signal, a desired output signal, and an actual output signal. If one minimizes the mean-squared error between the desired output signal and the actual output signal, it becomes possible to solve for the optimum system commonly known as the "Wiener" filter. The fundamental assumption underlying the procedure is that all processes are stationary.

A stationary time series is one whose statistical properties are time invariant. In particular, the statistics of the time series are independent of time. By definition, a stationary time series must be of infinite duration. However, in an actual experiment, we observe a time series over a finite interval. In order to apply the concepts of Wiener filtering, the finite length time series is convolved with a white noise series of unit power, to yield a stationary time series [3]. Moreover, in actual measurements only one waveform is often available for computing various statistics. Thus ensemble averages are frequently evaluated by means of appropriate time averages. This is valid only when the random processes are ergodic.

Manuscript received May 14, 1980; revised February 3, 1981. This work was supported in part by AFWL under Contract F33615-77-C-2059 and by the Office of Naval Research under Contract N00014-79-C-0598.

T. Sarkar is with the Department of Electrical Engineering, Rochester Institute of Technology, Rochester, NY 14623.

D. D. Weiner and J. Nebat are with the Department of Electrical Engineering, Syracuse University, Syracuse, NY 13210.

V. K. Jain is with the Department of Electrical Engineering, University of South Florida, Tampa, FL 33620.

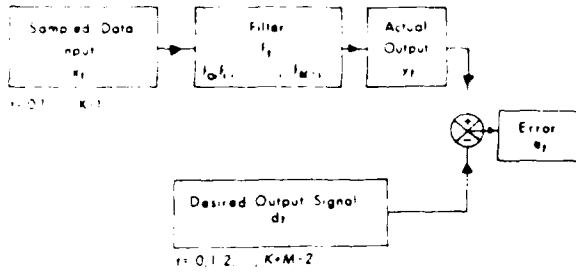


Fig. 1. Principles of Wiener filtering.

Although the concept of Wiener filtering is well-known, it is included here for completeness. The fundamental elements of digital Wiener filtering are summarized in Fig. 1 for the sampled data problem. Given the input sequence x_t and a desired output sequence d_t , the problem is to find the linear filter coefficients f_t , whose output sequence $x_t * f_t$ (the asterisk denotes convolution) yields a minimum mean-squared error estimate of d_t . If $E\{\cdot\}$ denotes the expected value, then the error

$$I = E\{(d_t - y_t)^2\} \quad (1)$$

is to be minimized.

In this case an M -length filter $f_t = \{f_0, f_1, \dots, f_{M-1}\}$ converts, in a least-error energy sense, a K -length input $x_t = \{x_0, x_1, \dots, x_{K-1}\}$ into a $K+M-1$ length sequence y_t which approximates the desired sequence $d_t = \{d_0, d_1, \dots, d_{K+M-2}\}$. The actual output is obtained as

$$y_t = \{y_0, y_1, \dots, y_{K+M-2}\} \\ = x_t * f_t = \sum_{\tau=0}^{M-1} f_\tau x_{t-\tau} \quad (2)$$

The problem of making the actual output y_t as close as possible to the desired output d_t can be interpreted in terms of minimizing the error "energy"

$$I = E\{(d_t - y_t)^2\} = E\left\{\left(d_t - \sum_{\tau=0}^{M-1} f_\tau x_{t-\tau}\right)^2\right\} \quad (3)$$

The error is minimized by evaluating the partial derivatives of I with respect to f_j and equating them to zero. This results in a set of equations

$$\frac{dI}{df_j} = E\left\{2\left(d_t - \sum_{\tau=0}^{M-1} f_\tau x_{t-\tau}\right)(-x_{t-j})\right\} \\ = -2E\{d_t x_{t-j}\} + 2 \sum_{\tau=0}^{M-1} f_\tau E\{x_{t-\tau} x_{t-j}\} \\ = 0, \quad \text{for } j = 0, 1, 2, \dots, M-1. \quad (4)$$

The unknown filter coefficients are obtained by solving the following set of simultaneous equations. This is the discrete Wiener-Hopf equation.

$$\sum_{\tau=0}^{M-1} f_\tau E\{x_{t-\tau} x_{t-j}\} = E\{d_t x_{t-j}\}$$

for $j = 0, 1, \dots, M-1$.

In order to solve the above equations, it is necessary to compute the expected values in (5). By assuming that the ensemble averages can be evaluated in terms of time averages, one obtains [3]

$$E\{x_{t-\tau} x_{t-j}\} \triangleq \frac{1}{K-M} \sum_{t=M}^{K-1} x_{t-\tau} x_{t-j} \triangleq C_{\tau j} \quad (6)$$

(in the covariance method)

$$\triangleq \frac{1}{K} \sum_{t=0}^{K-1+M} x_{t-\tau} x_{t-j} \\ = \frac{1}{K} \sum_{t=0}^{K-|\tau-j|-1} x_t x_{t+|\tau-j|} \triangleq r_{(\tau-j)} \quad (7)$$

(in the autocorrelation method).

Similarly

$$E\{d_t x_{t-j}\} \triangleq \frac{1}{K-M} \sum_{t=M}^{K-1} d_t x_{t-j} \quad (8)$$

(in the covariance method)

$$\triangleq \frac{1}{K} \sum_{t=0}^{K-1+M} d_t x_{t-j} \quad (9)$$

(in the autocorrelation method).

It is important to stress that the terms "covariance" and "autocorrelation" are not based upon the standard usage of the terms as occur in the theory of stochastic processes. Rather, we follow the usage which is quite prevalent in the speech processing literature [5]. The following discussion is intended to clarify their interpretation.

It is clear that the covariance definition given by (6) yields an unbiased estimate since

$$E[E\{x_{t-\tau} x_{t-j}\}] \triangleq E\left[\frac{1}{K-M} \sum_{t=M}^{K-1} x_{t-\tau} x_{t-j}\right] \\ = \frac{1}{K-M} \sum_{t=M}^{K-1} E\{x_{t-\tau} x_{t-j}\} \\ = E\{x_{t-\tau} x_{t-j}\}.$$

On the other hand, the autocorrelation definition given by (7) results in a biased estimate since

$$E[E\{x_{t-\tau} x_{t-j}\}] \triangleq E\left[\frac{1}{K} \sum_{t=0}^{K-1+M} x_{t-\tau} x_{t-j}\right] \\ = \frac{1}{K} \sum_{t=0}^{K-|\tau-j|-1} E\{x_t x_{t+|\tau-j|}\} \\ = \left[1 - \frac{|\tau-j|}{K}\right] E\{x_{t-\tau} x_{t-j}\}.$$

(5) Since $(\tau - j) \leq M$, it is interesting to note that the bias is

negligible when M (the order of the filter) is much less than K . The bias is significant, however, when K is only slightly larger than M . This explains why a large number of data points ($K \gg M$) is necessary for unbiased spectral estimation when using the autocorrelation function to obtain the power spectral density.

Because the covariance method gives an unbiased estimate, it might be assumed that it is a better estimate. However, the biased estimate provided by the autocorrelation method is often preferable. As an example consider the zero mean four data point sequence (4, -2, -1, -1) and $M = 2.0$. Using the unbiased estimator, the expected values are $C_{00}' = 1.0$, $C_{10}' = 1.5$, $C_{20}' = -1.0$, $C_{30}' = -2.0$. In contrast, the biased estimator results in $r_{(0)} = 5.5$, $r_{(1)} = -1.25$, $r_{(2)} = -0.5$, $r_{(3)} = -1.0$. Note that C_{00}' is less than C_{10}' whereas $r_{(0)}$ is guaranteed to be greater than $r_{(\tau)}$ for $\tau > 0$. Hence, the autocorrelation method yields nonnegative power spectral densities whereas the covariance method may not.

Continuing with our development, the discrete Wiener-Hopf equation presented in (5) can be written in the following matrix form for the covariance method:

$$[C_{ij}]_{M \times M} [F_i]_{M \times 1} = [D_j]_{M \times 1} \quad (10)$$

where $[C_{ij}]$ is a square matrix whose elements are given as

$$C_{ij} = \sum_{\tau=M}^{K-1} x_{t-i} x_{t-j}, \quad (11)$$

$[F_i]$ is a column matrix consisting of the M unknown filter coefficients $f_0, f_1, \dots, f_{M-1}$, and $[D_j]$ is a column matrix whose elements are given by

$$D_j = \sum_{\tau=M}^{K-1} d_{\tau} x_{t-j}, \quad (12)$$

For the autocorrelation method, the unknown filter coefficients are obtained from the solution of the matrix equation

$$[R_{|i-j|}]_{M \times M} [F_i]_{M \times 1} = [D_j]_{M \times 1} \quad (13)$$

where $[R_{|i-j|}]$ is a square matrix whose elements are given as

$$R_{|i-j|} = \sum_{\tau=0}^{K-1+M} x_{t-i} x_{t-j} = \sum_{\tau=0}^{K-1-|i-j|} x_{\tau} x_{\tau+|i-j|} \quad (14)$$

and $[D_j]$ is a column matrix defined by [12].

Interestingly, most of the formulations for the solution of an unknown linear filter lead to analysis equations which can be formulated either in terms of the autocorrelation matrix equations (13) or in terms of the covariance matrix equations (10).

It is important to note that the Wiener filter is not always realizable as a causal rational function (in terms of poles and zeros). As a result, the Wiener filter is, in general, an infinite order filter. When the order of the filter is specified *a priori*, the resulting filter may no longer be optimum.

Various forms of the Wiener filter have appeared under different names and have been used in various geophysical, speech processing, and digital signal processing applications. Next, various modifications of the Wiener filter are presented.

A. Inverse Filter Formulation

The inverse filter attempts to transform the input signal into an impulse [6]. Assume that the input sequence x_t is

transformed to an impulse of area σ by an all-zero filter of the form

$$F(z) = \sum_{l=0}^{M-1} f_l z^{-l}, \quad \text{with } f_0 = 1. \quad (15)$$

In terms of Fig. 1, the desired waveform d_t is an impulse of area σ . It follows that the coefficients of the filter should be chosen such that

$$\sum_{\tau=0}^{M-1} f_{\tau} x_{t-\tau} = \sigma \delta_t \quad (16)$$

where $\delta_t = 1$ for $t = 0$ and zero otherwise. Multiplication of both sides by x_{t-j} and summing t from k to $K-1+l$, one obtains

$$\sum_{\tau=0}^{M-1} f_{\tau} \left[\sum_{t=k}^{K-1+l} x_{t-\tau} x_{t-j} \right] = \begin{cases} 0 & k > 0 \\ \sigma x_{-j}, & k = 0 \end{cases} \quad (17)$$

Since x_{-j} is zero for $j > 0$, the unknown coefficients f_{τ} for $\tau = 1, 2, \dots, M-1$ are obtained from the solution of the following equations:

$$\sum_{\tau=0}^{M-1} f_{\tau} \left[\sum_{t=k}^{K-1+l} x_{t-\tau} x_{t-j} \right] = 0 \quad (18)$$

for $j = 1, 2, \dots, M-1$. For $k = M$ and $l = 0$, the above equations reduce to that of the covariance equations as in (10). This is because

$$D_j = \sum_{\tau=M}^{K-1} \delta_{\tau} x_{t-j} = 0. \quad (19)$$

For $k = 0$ and $l = M$, (18) reduces to that of the autocorrelation equations. If the input signal is approximated by an all-pole model, the poles of the input signal are obtained from the zeros of $F(z)$, i.e., from the solution of the polynomial equations

$$\sum_{\tau=0}^{M-1} f_{\tau} z^{-\tau} = 0. \quad (20)$$

In particular, if $(z^{-1})_i$ is the i th root of the above equation (20), then the i th pole is equal to $\ln [(z^{-1})_i]$.

B. Linear Prediction

The term "linear prediction" was first used by Wiener in his classic work on prediction of stationary time series. Since its publication, it has found wide application in the determination of all-pole models for the processing of speech signals [5].

The basic philosophy here is to take a part of the sampled waveform (say the first $M-1$ points from K data points) and predict the next data point on the waveform by proper choice of the predictor coefficients a_i . The linear predictor of step size one predicts the M th data point of the waveform when a $(M-1)$ order predictor filter is chosen. In the time domain, the predicted sample $\hat{x}_M$ is given by

$$\hat{x}_M = - \sum_{i=1}^{M-1} a_i x_{M-i}$$

where $(-a_1, -a_2, -a_3, \dots, -a_{M-1})$ are the predictor coefficients. Assuming the signal spectrum is to be modeled by an all-pole model, the coefficients a_i are the negative of the values of f_i presented in (15). A $(M-1)$ order linear predictor thus requires a linear combination of the previous $(M-1)$ samples. The error is then given by

$$\begin{aligned} e_M &= \hat{x}_M + x_M \\ &= -\sum_{i=1}^{M-1} a_i x_{M-i} + x_M \\ &= -\sum_{i=0}^{M-1} a_i x_{M-i}, \quad \text{with } a_0 = -1. \end{aligned} \quad (21)$$

When $(K-1+l)$ different data points are predicted, the total squared error is defined by

$$\begin{aligned} I &= \sum_{t=k}^{K-1+l} [e_t]^2 = \sum_{t=k}^{K-1+l} \left[\sum_{i=0}^{M-1} a_i x_{t-i} \right]^2 \\ &= \sum_{i=0}^{M-1} \sum_{j=0}^{M-1} a_i a_j \left\{ \sum_{t=k}^{K-1+l} x_{t-i} x_{t-j} \right\} \\ &= \sum_{i=0}^{M-1} \sum_{j=0}^{M-1} f_i f_j \left\{ \sum_{t=k}^{K-1+l} x_{t-i} x_{t-j} \right\}. \end{aligned} \quad (22)$$

Minimization of I , with respect to the set of the filter coefficients leads to a set of simultaneous equations given by

$$\sum_{i=0}^{M-1} f_i \left\{ \sum_{t=k}^{K-1+l} x_{t-i} x_{t-j} \right\} = 0 \quad (23)$$

for $j = 1, 2, \dots, M-1$, from which the unknown filter coefficients f_i are obtained. When $k = M$ and $l = 0$, this amounts to minimization of the error only over data points of the waveform from M to $K-1$. When $k = 0$ and $l = M$, this implies minimization over the entire waveform. Equation (23) are identical to the set of equations (19) obtained in the case of the inverse filter formulation in the previous section.

Linear prediction is, therefore, equivalent to an all-pole model for the spectrum of the input x_t . The poles for this model are again obtained from the solution of the polynomial equation

$$\sum_{i=0}^{M-1} f_i z^{-i} = 0.$$

Thus, if x_t represents the measured impulse response of the system, the set of poles obtained by linear prediction can be interpreted so as to parameterize the system in terms of an all-pole model.

C. Predictive Deconvolution or Spiking Filter Design

The general linear filtering problem involves the input x_t , the impulse response h_t and the output y_t . When it is desirable to evaluate h_t given x_t and y_t , the problem is referred to as deconvolution. In this sense the inverse filter problem discussed in Section II-A is a deconvolution problem. Predictive deconvolution refers to the case in which the output

y_t is assumed to be a delayed impulse such that

$$\sum_{\tau=0}^{M-1} f_\tau x_{t-\tau} = \sigma \delta_{t-\alpha+1}. \quad (24)$$

The unknown filter coefficients f_τ can be pursued in a manner analogous to the inverse filter approach by assuming the input to be represented as an all-pole model. However, we prefer to show that the filter coefficients can also be determined by interpreting the problem as a prediction problem. It is in this sense that the term predictive deconvolution is used [7].

Introduce the change of variables $t - \alpha + 1 = \beta$ in (24). The resulting equation is

$$\sum_{\tau=0}^{M-1} f_\tau x_{\beta+\alpha-\tau-1} = \sigma \delta_\beta. \quad (25)$$

Next, assume the filter coefficients to be given by

$$\begin{aligned} &\alpha-1 \text{ zeros} \\ &(-1, 0, 0, \dots, 0, -a_1, -a_2, \dots, -a_{M-1}). \end{aligned} \quad (26)$$

The upper limit on the summation is now given by $M + \alpha - 2$. Equation (25) can now be written as

$$\begin{aligned} x_{t+\alpha-1} + \sum_{\tau=1}^{M+\alpha-2} f_\tau x_{t+\alpha-\tau-1} \\ = x_{t+\alpha-1} - \sum_{\tau=1}^{M-1} a_\tau x_{t-\tau} \\ = \sigma \delta_t. \end{aligned} \quad (27)$$

If the estimate of $x_{t+\alpha-1}$ is assumed to be given by $\hat{x}_{t+\alpha-1}$, then

$$\hat{x}_{t+\alpha-1} = -\sum_{\tau=1}^{M-1} a_\tau x_{t-\tau} \quad (28)$$

and $\sigma \delta_t$ can be interpreted as the error of the estimate.

If the total squared error is minimized, as was done for the linear prediction approach in Section II-B, a set of equations is obtained for the filter coefficients. When $\alpha = 1$, these equations reduce to the same equations as were obtained in (23).

D. Recursive Filter Design

Recursive filters are often used in digital filter design [8]-[12]. With respect to Fig. 1, the spectral estimation problem is now posed in the following manner. Assume the filter input is the impulse δ_t . Let the filter impulse response be determined so as to approximate the data sequence in a minimum mean squared error sense. If x_t is the input to the filter and y_t is the output, then they are related by

$$\sum_{k=0}^{M-1} a_k y_{t-k} = \sum_{r=0}^{L-1} b_r x_{t-r}.$$

Application of the z -transform to both sides yields

$$\sum_{k=0}^{M-1} a_k z^{-k} Y(z) = \sum_{r=0}^{L-1} b_r z^{-r} X(z)$$

or

$$F(z) = \frac{Y(z)}{X(z)} = \frac{\sum_{r=0}^{L-1} b_r z^{-r}}{\sum_{k=0}^{M-1} a_k z^{-k}} \equiv \frac{B(z)}{A(z)} \quad (29)$$

Therefore, a recursive filter has a transfer function which is expressed as a ratio of two polynomials of the z -transform variable. The objective here is to synthesize the filter from the given impulse response of the system. It is important to note that, unlike the previous three types of filters, this is not an all-pole filter but a pole-zero model. In other words, given some desired filter operator $D(z)$ (which is the transfer function of a desired system having the impulse response d_t), we require the polynomials $A(z)$ and $B(z)$ of the filter

$$F(z) = \frac{B(z)}{A(z)} \quad (30)$$

such that

$$F(z) \approx D(z) = d_0 + d_1 z^{-1} + d_2 z^{-2} + \dots + d_{K-1} z^{-K+1} \quad (31)$$

A technique for determining $A(z)$ and $B(z)$ is outlined next.

$$A(z) = 1 + a_1 z^{-1} + a_2 z^{-2} + \dots + a_{M-1} z^{-M+1} \quad (32)$$

and

$$B(z) = b_0 + b_1 z^{-1} + b_2 z^{-2} + \dots + b_{L-1} z^{-L+1} \quad (33)$$

where M and L are arbitrary numbers which fix the number of poles and zeros, respectively, for the filter. Also, divide $B(z)$ by $A(z)$ so as to obtain the infinite series

$$F(z) = f_0 + f_1 z^{-1} + f_2 z^{-2} + \dots$$

Since $F(z)A(z) = B(z)$, it follows that

$$F(z)(1 + a_1 z^{-1} + a_2 z^{-2} + \dots + a_{M-1} z^{-M+1}) = b_0 + b_1 z^{-1} + b_2 z^{-2} + \dots + b_{L-1} z^{-L+1}$$

Since multiplication of z -transforms is equivalent to convolution of the discrete time series, the series of b_t coefficients is equal to the convolution of the f_t coefficients with the a_t coefficients. Or, equivalently,

$$b_t = \sum_{j=0}^{M-1} a_j f_{t-j}$$

By assumption, $a_0 = 1$. Therefore,

$$f_t = b_t - \sum_{j=1}^{M-1} a_j f_{t-j}$$

As $b_t = 0$ for $t > L$ (from (33)), one may write

$$f_t = - \sum_{j=1}^{M-1} a_j f_{t-j}, \quad \text{for } t > L. \quad (34)$$

By assuming the approximation in (31) is valid, the coefficients f_t approximate the coefficients d_t for $t > L$. Equation (34) may then be approximated by

$$d_t \approx - \sum_{j=1}^{M-1} a_j d_{t-j}, \quad \text{for } t = L+1, L+2, \dots, K-1. \quad (35)$$

The error in (35) is given by

$$e_t = d_t + \sum_{j=1}^{M-1} a_j d_{t-j} = \sum_{j=0}^{M-1} a_j d_{t-j}; \text{ since } a_0 = 1, \\ \text{for } t = L+1, L+2, \dots, K-1. \quad (36)$$

The a_j coefficients are chosen in such a way that the mean-squared error is minimized. In particular,

$$J = \sum_{t=L+1}^{K-1} e_t^2 = \sum_{t=L+1}^{K-1} \left[\sum_{j=0}^{M-1} a_j d_{t-j} \right]^2 \quad (37)$$

is minimized. By differentiating J of (37) with respect to a_j and equating the derivatives to zero, a set of simultaneous equations is obtained. They are given by

$$\sum_{j=0}^{M-1} a_j \left[\sum_{t=L+1}^{K-1} d_{t-j} d_{t-k} \right] = 0, \\ \text{for } k = 1, 2, \dots, M-1. \quad (38)$$

For a realizable filter the numerator polynomial is generally one degree lower than the denominator polynomial (if the poles are simple). Thus

$$L+1 = M.$$

Hence, (38) reduces to

$$\sum_{j=0}^{M-1} a_j \left[\sum_{t=M}^{K-1} d_{t-j} d_{t-k} \right] = 0 \quad (39)$$

which is equivalent to the covariance equations as given by (10). The right side of (10) is zero because

$$D_j = \sum_{t=M}^{K-1} \delta_{t-j} = 0, \quad \text{for } j = 1, 2, \dots, M-1.$$

The poles for the filter are then obtained by the solution of

the polynomial equation

$$\sum_{i=0}^{M-1} a_i z^{-i} = 0.$$

It was shown previously that the inverse filter, assuming an all-pole model for the signal spectrum, is also equivalent to (10). We conclude, therefore, that the poles for an all-pole model correspond to the identical set of poles for a pole-zero model for the same order filter.

Next the residues at the poles (or equivalently, the numerator polynomial) can be obtained by minimizing the mean-squared error given by $\sum_i \{f_i - d_i\}^2$.

It is interesting to note that when the numerator polynomial is realized directly in the form presented in (33), the problem reduces to the case of the Padé approximation [13]. Mathematically, Padé approximation results in an approximation of $D(z)$ by $F(z)$ such that the seminorm

$$\|D(z) - F(z)\| = |D(1) - F(1)| + |D^1(1) - F^1(1)| + \dots + |D^{L+M-1}(1) - F^{L+M-1}(1)| \quad (40)$$

is made zero. Here $D^k(1)$ represent the k th derivative of $D(z)$ evaluated on the unit circle.

E. Kalman Filter Theory

Underlying Wiener filter design is the so-called Wiener-Hopf integral equation, its solution through spectral factorization, and the practical problem of synthesizing the theoretically optimal filter from its impulse response. The normal Wiener filter is derived from the Wiener-Hopf equation and, in general, this equation can be solved only in the steady state, i.e., when the observation interval is semi-infinite. The contribution of Kalman was recognition of the fact that the integral equation could be converted into a nonlinear differential equation whose solution contains all the necessary information for design of the optimal filter. The problem of spectral factorization in the Wiener filter is analogous to the requirement for solving $M(M+1)/2$ coupled nonlinear algebraic equations in the M -order Kalman filter. These equations can be solved numerically for transient type problems, where data is available only for a finite interval. This, in general, results in the Kalman filters being time variant. However, in the steady-state the Kalman filter reduces to the time invariant Wiener filter [14]. The presentation by Sorensen [15] expresses the results of Kalman filter theory in a way that makes this comparison easier.

The problem involves estimating a signal $\{s_n\}$, from measured data $\{d_0, d_1, \dots, d_{K-1}\}$. If the estimate is computed as a linear combination of the d_n , then

$$\hat{s}_n = \sum_{i=0}^{M-1} A_i d_i \quad (41)$$

The M coefficients A_i are chosen in such a way that the mean-squared error,

$$I = E[(s_n - \hat{s}_n)^T (s_n - \hat{s}_n)], \quad (42)$$

is minimized. Here T denotes the transpose of the row vector $(s_n - \hat{s}_n)$. This criterion is satisfied when the error in the

estimate $\hat{s}_n$ is orthogonal to the measured data, or

$$E[(s_n - \hat{s}_n) d_i^T] = 0, \quad \text{for } i = 0, 1, \dots, M-1. \quad (43)$$

Expressed in a different way

$$E[s_n d_i^T] = \sum_{i=0}^{M-1} A_i E[d_i d_i^T], \quad \text{for } i = 0, 1, \dots, M-1. \quad (44)$$

However, the matrix inversion that is required becomes computationally impractical when M is large. Wiener and Kolmogoroff assumed an infinite amount of data (that is, the lower limit of the summation is $-\infty$ rather than zero). Equation (43) is then the discrete form of the Wiener-Hopf equation which was solved using spectral factorization.

The basic difference between Wiener-Kolmogoroff theory and the Kalman filter theory is how equation (44) is solved. In 1955, J. W. Follin suggested a recursive approach to solve (44). It is clear [see 15, p. 65] that Follin's work provided a direct stimulus for the work of Richard Bucy, which led to his subsequent collaboration with Kalman in the total development of the "state space" approach for obtaining the filter equations.

F. Summary

As outlined above, all forms of the digital Wiener filter lead either to the covariance or the autocorrelation equations. It is also interesting that the same set of poles is obtained whether one models the signal as an all-pole model or as a pole-zero model.

III. STOCHASTIC METHODS APPLIED TO SYSTEM IDENTIFICATION

Three different stochastic methods used in spectral estimation are presented in this section. They are maximum likelihood estimation, minimum predictor error variance estimation, and maximum entropy spectral analysis. All three models are considered as all-pole models. These methods have no relationship with discrete Wiener filtering theory. The methods presented in this section start with completely different assumptions but finally yield the identical set of either covariance or autocorrelation equations which characterize the system to be identified from the measured impulse response.

In these methods it is assumed that the data samples are part of a random process. The problem is to choose the parameters of the system impulse response so as to make the probability of occurrence of the actual observation most likely. In other words, the system parameters are chosen in such a way that the probability density function defining the parameters is maximized.

A. Maximum Likelihood Estimation Theory

In this approach the measured impulse response of the system is considered as a segment of a random process. It is further assumed that the impulse response can be generated by passing an uncorrelated noise sequence $\{e_i\}$ through an all-pole model of the form

$$\frac{1}{F(z)} = \frac{1}{\sum_{i=0}^{M-1} f_i z^{-i}} \quad \text{with } f_0 = 1. \quad (45)$$

Next the random process $\{e_t\}$ is assumed to be Gaussian with zero mean and variance σ_e^2 . Thus,

$$E\{e_t\} = 0 \quad \text{and} \quad E\{e_i e_j\} = \delta_{ij} \sigma_e^2. \quad (46)$$

As the measured impulse response x_t , $t = 0, 1, \dots, K-1$ has been assumed to be generated by passing the noise sequence $\{e_t\}$ through the all-pole model, it follows that

$$\sum_{i=0}^{M-1} f_i x_{t-i} = e_t. \quad (47)$$

From (46) and (47) it is clear that the sequence x_t is Gaussian with zero mean and a cross correlation defined by

$$E[x_i x_j] = g_{i-j}. \quad (48)$$

This correlation sequence g_{i-j} would then be a function of the system parameters f_i , $i = 0, 1, \dots, M-1$ and σ_e^2 . Since x_t is Gaussian, a Gaussian multivariate probability density function is defined for the sequence of random variables $x_0, x_1, \dots, x_{K-1}$. The maximum likelihood theory assumes that the parameter values which make the measured observation of the impulse response most likely are the same values which maximize the joint probability density function of x_i , $i = 0, 1, \dots, K-1$. This can be achieved by differentiating the density function with each of the unknown variables, $f_1, f_2, \dots, f_{M-1}$ and σ_e^2 and then setting the first partial derivative equal to zero. The solutions of the set of equations then yield the values for the unknown parameters. Even though the procedure is conceptually simple, the set of equations becomes extremely nonlinear for M greater than 2 and no exact solution for this problem exists [2].

However, Itakura and Saito [16], [17] solved the maximum likelihood problem by making some additional assumptions. First, the number of data points K is made much greater than M , the order of the filter (i.e., $K \gg M$). Second, the joint probability density function for the sequence $x_0, x_1, x_2, \dots, x_{K-1}$ is approximated by

$$p\{x_0, x_1, \dots, x_{K-1}\} = [2\pi\sigma_e^2]^{-K/2} \exp[-\alpha/2\sigma_e^2] \quad (49)$$

where

$$\alpha = \sum_{t=0}^{K-1+M} \left[\sum_{i=0}^{M-1} f_i x_{t-i} \right]^2. \quad (50)$$

It is interesting to note that α has the identical form of the error defined by the autocorrelation equations in linear prediction (see (22)).

It has been shown [16], [17] that the results obtained for the unknown filter coefficients f_i are identical to (13) which utilizes the autocorrelation equations.

The corresponding equations for the covariance method are obtained by defining a conditional density function for the probability density. This is achieved by treating the M data points $x_0, x_1, \dots, x_{M-1}$ as a set of deterministic initial conditions and the remaining $K-M$ data points as a set of random variables. Under the above assumptions, the conditional probability density function is approximated as [17], [2]:

$$p_c\{(x_M, x_{M+1}, \dots, x_{K-1}) | (x_0, x_1, \dots, x_{M-1})\} = (2\pi\sigma_e^2)^{-0.5(K-M)} \exp[-\alpha/2\sigma_e^2] \quad (51)$$

where

$$\alpha = \sum_{t=M}^{K-1} \left[\sum_{i=0}^{M-1} f_i x_{t-i} \right]^2. \quad (52)$$

Again it is clear that α has the same form as the error energy defined for the covariance equations for the case of linear prediction (see (22)). Maximization of the conditional probability density function is then achieved by maximizing p_c with respect to the unknown filter coefficients f_i and σ_e^2 . The set of equations (see (10)) obtained in this case are identical to those of the covariance method.

B. Minimum Predictor Error Variance

In this method the data samples are not considered as a part of a Gaussian process. In other words, the method remains the same as before, i.e., the measured impulse response is generated by passing a noise sequence $\{e_t\}$ through an all-pole model [2]. It is assumed that the error sequence is of zero mean and of variance given by

$$E\{e_t^2\} = \sum_{i=0}^M \sum_{j=0}^M f_i f_j E[x_{t-i} x_{t-j}]. \quad (53)$$

Next the process is assumed to be stationary so that the expectation in (53) can be expressed as

$$E[x_{t-i} x_{t-j}] = g_{i-j} \quad (54)$$

and the error variance as

$$E\{e_t^2\} = \sum_{i=0}^M \sum_{j=0}^M f_i f_j g_{i-j}. \quad (55)$$

The problem is to determine the filter coefficients so as to minimize the error variance.

An additional assumption is now made. Specifically, it is assumed that the process is ergodic so that the ensemble average E may be converted to a time average. Hence, the approximation

$$g_{i-j} \triangleq \frac{1}{K-M} \sum_{t=M}^{K-1} x_{t-i} x_{t-j} = C_{ij} \quad (56)$$

leads to the covariance equation (10) and the approximation

$$g_{i-j} \triangleq \frac{1}{K} \sum_{t=0}^{K-|i-j|-1} x_t x_{t+|i-j|} = r_{i-j} \quad (57)$$

leads to the autocorrelation equations (13). Hence, this method yields a set of analysis equations identical to those for the discrete Wiener filter.

C. Maximum Entropy Spectral Analysis

An important aspect of time series analysis is the computation of the power spectral density which is primarily determined by the second-order statistics. In an actual experiment, the number of data points is always finite. Hence, for the problem of interest the data length may not be sufficient to obtain a specified degree of frequency resolution. Also, given a finite number of K data points, we can obtain at

most approximations of the K autocorrelation functions $r_0, r_1, \dots, r_{K-1}$. In the previous autocorrelation formulations the data has been assumed to be zero outside the known interval. In some instances, this may be an unreasonable assumption about the extension of the data beyond the known interval. The question then arises as to what assumptions should be made about the data outside the finite sample and what assumption should be made about their second order statistics (i.e., the autocorrelation), since they determine the power spectral density.

Burg proposed an information theory approach to the problem. He suggested [18] that the most reasonable choice of the unknown autocorrelations is the one which adds no information or adds most randomness or maximizes the entropy. He then proceeded to select the power spectral density having the maximum entropy of all possible spectra that agrees with the known values of the autocorrelation function r_i .

The information content of a random process is defined in terms of a quantity called entropy and is mathematically expressed as

$$H = - \sum_j P_j \ln P_j \quad (58)$$

where P_j is the probability of the j th event of a random process. When the random variable takes on a continuum of values, the sum in the definition of the entropy is replaced by an integral. Since we are dealing with a time series $x_0, x_1, \dots, x_{K-1}$, the probability is replaced by the joint probability density function $p(x_0, x_1, \dots, x_{K-1})$. Thus

$$H = - \int p(x_0, x_1, \dots, x_{K-1}) \cdot \ln \{p(x_0, x_1, \dots, x_{K-1})\} dV \quad (59)$$

where dV is an element of volume in the space spanned by the random variables. Burg then proceeded to adjoin a hypothetical variable x_K to the available estimates of the autocorrelation function r_0, r_1, r_2 , and so on. We may then consider the joint probability density available for the K data points and the adjoined x_K as

$$p(x_0, x_1, \dots, x_{K-1}, x_K) \quad (60)$$

This probability density function has an entropy

$$H = - \int p(x_0, x_1, \dots, x_{K-1}, x_K) \cdot \ln \{p(x_0, x_1, \dots, x_{K-1}, x_K)\} dV \quad (61)$$

Burg chose as (60) that probability density function which has its first K second-order moments as $r_0, r_1, \dots, r_{K-1}$, and which under the given constraint maximized (61). The obvious choice for the probability density function in (60) is Gaussian since according to Shannon and Weaver [19], [20] the Gaussian distribution results in maximum entropy under a constant energy constraint. Thus

$$p(x_0, x_1, \dots, x_{K-1}, x_K) = \frac{\exp \{-\frac{1}{2} X' [R_K]^{-1} X\}}{\left\{ (2\pi)^{\frac{K+1}{2}} \det [R_K] \right\}^{1/2}} \quad (62)$$

where X is the column vector of the x_i , the prime indicates the transpose, and the matrix $[R_K]$ is given by

$$[R_K] = \begin{bmatrix} r_0 & r_1 & \dots & r_{K-1} & r_K \\ r_1 & r_0 & & r_{K-2} & r_{K-1} \\ \vdots & \vdots & & \vdots & \vdots \\ r_{K-1} & & & & \\ r_K & & \dots & & r_0 \end{bmatrix} \quad (63)$$

The entropy can then be expressed as [19], [20]

$$H = \frac{1}{2} \ln \{ (2\pi e)^{K+1} \det [R_K] \} \quad (64)$$

Now r_K is to be chosen in such a way that H in (64) is maximized. Hence the value of r_K is the one which maximizes $\det [R_K]$.

In order for r_i to constitute a proper set of autocorrelation values, the matrix $[R_K]$ must be positive semi-definite [21]. Moreover, $\det [R_K]$ is a quadratic function in r_K . It follows that maximizing $\det [R_K]$ with respect to r_K yields the value of r_K obtained from the solution of the following equation:

$$\det \begin{bmatrix} r_1 & r_0 & \dots & r_{K-2} \\ r_2 & r_1 & & r_{K-3} \\ \vdots & \vdots & & \vdots \\ r_K & r_{K-1} & \dots & r_1 \end{bmatrix} = 0 \quad (65)$$

Alternatively, if an all-pole $K-1$ order model is chosen, then from the previous sections we know that the autocorrelation functions for this problem are related by the K unknowns f_i as $r_j = r_{-j}$ and

$$r_j + \sum_{k=1}^{K-1} f_k r_{j-k} = 0, \quad \text{for } j = 1, 2, \dots, K \quad (66)$$

This is identical to (23), for $f_0 = 1$. The set of K equations in $K-1$ unknowns in (66) indeed has a solution which is found by solving the first $K-1$ equations. The last equation can be seen to be a linear combination of the first $K-1$ equations. In fact, the determinant of the above set of equations in (66) is identical to that of (65). In this sense, (66) is consistent with the maximum entropy method.

Thus it is shown that the extrapolated autocorrelation functions coincide with those functions which would have been predicted by the model of equation (66). Hence this procedure is equivalent to the all-pole model described by the maximum likelihood estimation [20]–[23].

Since so far as the second-order statistics are concerned, the sampled data x_i may be modeled arbitrarily closely by an all-pole model of order K , we may view the above process as an autoregressive process with input (white noise) and output (x_i) where the filter is described by the transfer function [22], [23]:

$$H(z) = \frac{1}{1 + \sum_{i=1}^{K-1} f_i z^{-i}} \quad (67)$$

Also

$$S_0(f) = S_i(f) |H(j\omega)|^2 \quad (68)$$

where $S_0(f)$ and $S_i(f)$ are the output and input power spectral density, respectively. The power spectral density of the process x_t has been shown to be [20]–[23]

$$S_0(f) = \frac{S_i(f)}{\left| 1 + \sum_{i=1}^K f_i \exp[-j2\pi f_i \Delta t] \right|^2} \quad (69)$$

where $S_i(f)$ is the power spectrum of the white noise driving the filter.

Thus, for $M = K$, identical analysis equations are obtained by the maximum entropy spectral analysis and by the maximum likelihood estimation theory. It has also been shown elsewhere that this is indeed so [20]–[23].

IV. PRONY'S ALGORITHM AND ITS EXTENSIONS

In the previous two sections the system identification problem was treated as a stochastic problem. The measured impulse response was characterized by a random process. In this section a different approach is taken in that the noise contaminated impulse response is processed as though it were a deterministic process. The problem is to determine the poles and residues which characterize the measured impulse response.

Historically, Prony was the first to make an attempt at fitting experimental data with complex exponentials. In 1795 Prony postulated that the basic laws dealing with gas expansion can be expressed as a sum of exponentials. He demonstrated that, given $2M$ data points, it is possible to fit exactly M exponentials to the data at those points. Prony must have experienced great frustration when he applied his method due to the extreme sensitivity of the exponent to the accuracy of the measured data. Accuracy requirements have been investigated by Lanczos [24]. However, in some cases of EMP problems, the exponentials encountered are complex and they are approximately at harmonic frequencies. Moreover, the real parts of the complex exponentials are much smaller than the imaginary parts. Hence, the damped exponentials in the case of electromagnetic pulse problems are more closely orthogonal and, thereby, create fewer problems with regard to accuracy than is evidenced in the example by Lanczos. The contamination of data by noise creates grave problems for extracting M correct exponents from $2M$ data points. Hence, more data are necessary and a semi-least squares approach to Prony's method is taken. The details of Prony's method are well known [25]–[31] and are omitted in this paper. The final equations that ultimately result are the same as either the autocorrelation or covariance equations of linear prediction. McDonough [13] and Van Blaricum [25], in their Ph.D. dissertations, and Markel and Gray [2] used the covariance equations (10). Thus, the semi-least squares Prony's method is equivalent to a M -length Wiener prediction filter. The term semi-least squares has been applied because the true least squares problem would give rise to a set of coupled nonlinear equations [26]. The true least squares problem has been defined as in [28].

Prony's method has found wide applications in finding poles and residues of transient waveforms through the works of Miller, Brittingham, and Willows [27].

V. CONCLUSION

We have demonstrated that identical analysis equations can be obtained by several different techniques even though they start with formulations based on different assumptions.

However a major objection with all these techniques is that there seems to be no compelling reason for the choice of $a_0 = 1$ rather than $a_{M-1} = 1$, or for that matter the choice of any a_i to be unity. Yet each of these choices leads to a different set of exponents [13]. As an example consider the signal $\exp[-0.00035t] \cos(0.25t)$. For this signal we have $s_{1,2} = \exp[-0.00035 \pm j0.25] = 0.97 \pm j0.25$. Let the length of data points be $K = 5$ and the order of the filter $M = 2$. Then $\{x_t\}$ is the sequence $\{1.00, 0.97, 0.88, 0.73, 0.54\}$ for the sequence $\{t = 0, 1, 2, 3, 4\}$. To this we add 1 percent zero-mean additive noise to obtain the sequence $\{x_t^{\text{noise}}\} = \{1.01, 0.96, 0.89, 0.72, 0.54\}$. Using the covariance method, $s_{1,2}$ for this signal are obtained from the two roots of the polynomial equation $a_0 s^2 + a_1 s + a_2 = 0$. If one assumes $a_0 = 1$, then the computed $s_{1,2} = 0.87 \pm j.23$. If one assumes $a_1 = 1$, then the computed $s_{1,2} = 0.98 \pm j.22$. If one assumes $a_2 = 1$, then the computed $s_{1,2} = 1.1 \pm j.14$.

The above example thus illustrates the dependence of the exponents on the choice of a_i to be unity.

ACKNOWLEDGMENT

The authors wish to thank Drs. C. E. Baum, Dr. J. P. Castillo, Capt. M. Harrison, Capt. H. G. Hudson, Capt. E. Fuller and W. Prather of AFWL and Dr. H. Mullaney of ONR for a keen interest in all aspects of this work. The constructive criticisms offered by Dr. R. M. Bevensee in preparing this manuscript are deeply appreciated.

REFERENCES

- [1] J. D. Markel, "Formant trajectory estimation from a linear least-squares inverse filter formulation," SCRL-Monograph 7, Speech Commun. Res. Lab., Inc., Santa Barbara, CA, Oct. 1971.
- [2] J. D. Markel and A. H. Gray, *Linear Prediction of Speech*, Berlin: Springer-Verlag.
- [3] E. A. Robinson and S. Treitel, "Principles of digital Wiener filtering," *Geophys. Prospecting*, pp. 311–333, Sept. 1967.
- [4] N. Levinson, "The Wiener R.M.S. error criterion in filter design and prediction," *J. Math. Phys.*, vol. 26, pp. 261–278, 1947.
- [5] J. Makhoul, "Linear prediction: A tutorial review," *Proc. IEEE*, vol. 63, no. 4, pp. 561–580, Apr. 1975.
- [6] J. D. Markel, "Digital inverse filtering—A new tool for formant trajectory estimation," *IEEE Trans. Audio Electroacoust.*, vol. AU-18, no. 2, pp. 137–141, June 1970.
- [7] K. L. Peacock and S. Treitel, "Prediction deconvolution. Theory and practice," *Geophys.*, vol. 34, no. 2, pp. 155–169, Apr. 1969.
- [8] A. Oppenheim and R. Schaefer, *Digital Signal Processing*, Englewood Cliffs, NJ: Prentice Hall, 1975.
- [9] J. L. Shanks, "Recursion Filter for Digital Processing," *Geophys.*, vol. 32, no. 1, pp. 33–51, Feb. 1967.
- [10] D. L. Fletcher and C. N. Weygandt, "A digital method of transfer function calculation," *IEEE Trans. Circuit Theory*, pp. 185–187, Jan. 1971.
- [11] C. S. Burrus and T. W. Parks, "Time domain design of recursive digital filters," *IEEE Trans. Audio Electroacoust.*, vol. AU-18, No. 2, pp. 137–141, June 1970.
- [12] S. Treitel and E. A. Robinson, "The design of high resolution digital filters," *IEEE Trans. Geosci. Electron.*, vol. GE-4, no. 1, pp. 25–38, June 1966.
- [13] R. N. McDonough, "Representation and analysis of signals. Part XV—Matched exponents for the representation of signals," Johns Hopkins Univ., Apr. 1963.
- [14] A. Gelb, *Applied Optimal Estimation*, Cambridge, MA: MIT, 1974.
- [15] H. W. Sorenson, "Least-squares estimation. From Gauss to Kalman," *IEEE Spectrum*, pp. 63–68, July 1970.

- [16] F. Itakura and S. Saito, "Analysis synthesis telephony based on maximum likelihood method," presented at the 6th Int. Congress Acoustics, Tokyo, Japan, Aug. 21-28, 1968, C-5-5, pp. C-17-20.
 - [17] F. Itakura and S. Saito, "Extraction of speech parameters based upon the statistical method," *Proc. Speech Info. Process.*, Tohoku Univ., Sendai, Japan, 5.1, 5.12, 1971 (in Japanese).
 - [18] J. P. Burg, "Maximum entropy spectral analysis," Ph.D. dissertation Stanford Univ., Palo Alto, CA, 1975.
 - [19] C. E. Shannon and W. Weaver, *The Mathematical Theory of Communication*. Urbana, IL: Univ. Illinois Press, 1962, pp. 56-57.
 - [20] R. N. McDonough, "Maximum-entropy spatial processing of array data," *Geophys.*, vol. 39, no. 6, pp. 843-851, Dec. 1974.
 - [21] A. van den Bos, "Alternative interpretation of maximum entropy spectral analysis," *IEEE Trans. Inform. Theory*, vol. IT-17, no. 4, pp. 493-494, 1971.
 - [22] D. E. Smylie, G. K. C. Clarke, and T. J. Ulrych, "Analysis of irregularities in the earth's rotation," in *Methods in Computational Physics*, vol. 13. New York: Academic, 1973, pp. 391-430.
 - [23] S. L. Marple, "Conventional Fourier, autoregressive and spectral methods of spectral analysis," Ph.D. dissertation, Stanford Univ., Palo Alto, CA, 1976.
 - [24] C. Lanczos, *Applied Analysis*. Englewood Cliffs, NJ: Prentice-Hall, pp. 272-279, 1956.
 - [25] M. Van Blancum and R. Mitra, "Techniques for extracting the complex resonances of a system directly from its transient response," *Interaction Note* 301, Dec. 1975. (Also in *IEEE Trans. Antennas Propagat.*, vol. AP-23, no. 6, Nov. 1975.)
 - [26] R. N. McDonough and W. H. Huggins, "Best least-squares representation of signals by exponentials," *IEEE Trans.*, AC-13, no. 4, pp. 405-412, Aug. 1968.
 - [27] E. K. Miller, J. N. Brittingham, and J. L. Willows, "Part I: The Derivation of Simple Poles in a Transfer Function From Real Frequency Information," "Part II: Results from Real EM Data," "Part III: Object classification and Identification," Lawrence Livermore Laboratory, Livermore, CA, Tech. Reps. UCRL 52050, UCRL 52118, UCRL 52211.
 - [28] T. K. Sarkar, "The best least squares rational approximation of transfer functions by solution of a linear uncoupled eigenvalue problem," *IEEE Trans. Acoust., Speech, Signal Processing*, to be published.
- Tapan K. Sarkar** (S'69-M'76-SM'81), for a photograph and biography please see page 379 of the March 1981 issue of this TRANSACTIONS.
- Donald D. Weiner** (S'54-M'60), for a photograph and biography please see page 379 of the March 1981 issue of this TRANSACTIONS.
- Joshua Nebat**, photograph and biography not available at time of publication.
- Vijay K. Jain** (M'63-SM'75), for a photograph and biography please see page 379 of the March 1981 issue of this TRANSACTIONS.

Alternatively, if an all pole K-1 order model is chosen, then from the previous sections we know that the autocorrelation functions for this problem are related by the K unknowns f_i as $r_j = r_{-j}$ and

$$r_j + \sum_{k=1}^{K-1} f_k r_{j-k} = 0, \quad \text{for } j = 1, 2, \dots, K. \quad (3.22)$$

This is identical to (2.23), for $f_0 = 1$. The set of K equations in K-1 unknowns in (3.22) indeed has a solution which is found by solving the first K-1 equations. The last equation can be seen to be a linear combination of the first K-1 equations. In fact, the determinant of the above set of equations in (3.22) is identical to that of (3.21). In this sense, (3.22) is consistent with the maximum entropy method.

Thus it is shown that the extrapolated autocorrelation functions coincide with those functions which would have been predicted by the model of equation (3.22). Hence this procedure is equivalent to the all-pole model described by the maximum likelihood estimation [20-23].

Since so far as the second order statistics are concerned, the sampled data x_t may be modeled arbitrarily closely by an all-pole model of order K, we may view the above process as an autoregressive process with input (white noise) and output (x_t) where the filter is described by the transfer function [22, 23]

$$H(z) = \frac{1}{1 + \sum_{i=1}^{K-1} f_i z^{-i}}. \quad (3.23)$$

-
- 22 D.E. Smylie, G.K.C. Clarke and T.J. Ulrych, "Analysis of Irregularities in the Earth's Rotation," Methods in Computational Physics, Vol. 13, New York: Academic Press, 1973, pp. 391-430.
 - 23 S.L. Marple, "Conventional Fourier, Autoregressive and Spectral Methods of Spectral Analysis," Ph.D. Dissertation, Stanford University, Palo Alto, California, 1976.

Also

$$S_o(f) = S_i(f) |H(j\omega)|^2 \quad (3.24)$$

where $S_o(f)$ and $S_i(f)$ are the output and input power spectral density, respectively. The power spectral density of the process x_c has been shown to be [20-23]

$$S_o(f) = \frac{S_i(f)}{\left| 1 + \sum_{i=1}^{K-1} f_i \exp[-j2\pi f i \Delta t] \right|^2} \quad (3.25)$$

where $S_i(f)$ is the power spectrum of the white noise driving the filter.

Thus, for $M = K$, identical analysis equations are obtained by the maximum entropy spectral analysis and by the maximum likelihood estimation theory. It has also been shown elsewhere that this is indeed so [20-23].

IV. PRONY'S ALGORITHM AND ITS EXTENSIONS

In the previous two sections the system identification problem was treated as a stochastic problem. The measured impulse response was characterized by a random process. In this section a different approach is taken in that the noise contaminated impulse response is processed as though it was a deterministic process. The problem is to determine the poles and residues which characterize the measured impulse response.

Historically, Prony was the first to make an attempt at fitting experimental data with complex exponentials. In 1795 Prony postulated that the basic laws dealing with gas expansion can be expressed as a sum of exponentials. He demonstrated that, given $2M$ data points, it is possible to fit exactly M exponentials to the data at those points. Prony must have experienced great frustration when he applied his method due to the extreme sensitivity of the exponent to the accuracy of the measured data. Accuracy requirements have been investigated by Lanczos [24]. However, in some cases of EMP problems the exponentials encountered are complex and they are approximately at harmonic frequencies. Moreover, the real parts of the complex exponentials are much smaller than the imaginary parts. Hence, the damped exponentials in the case of EMP problems are more closely orthogonal and, thereby, create fewer problems with regard to accuracy

24 C. Lanczos, Functional Analysis. Englewood Cliffs, N.J.: Prentice Hall, 1956, pp. 272-279.

than is evidenced in the example by Lanczos. The contamination of data by noise creates grave problems for extracting M correct exponents from $2M$ data points. Hence, more data are necessary and a semi-least squares approach to Prony's method is taken. The details of Prony's method are well known and are omitted in this report. The final equations that ultimately result are the same as either the autocorrelation or covariance equations of linear prediction. McDonough [13] and Van Blaricum [25], in their Ph.D. dissertations, and Markel and Gray [2] used the covariance equations (2.10). Thus, the semi-least squares Prony's method is equivalent to a M -length Wiener prediction filter. The term semi-least squares has been applied because the true least squares problem would give rise to a set of coupled nonlinear equations [26]. The true least squares problem has been defined as in [26].

4.1 Various Extensions to Prony's Method

The reason for presenting this section is to show that for a particular extension to Prony's method a procedure similar to the pencil-of-functions method is posed in a Hilbert space. Yet each of them yields a different answer. Hence, the way in which a problem is developed is extremely important.

-
- 25 M. Van Blaricum and R. Mittra, "Techniques for Extracting the Complex Resonances of a System directly from its Transient Response; Interaction Note 301, December 1975. (Also in IEEE Trans. on Antennas and Propagation, Vol. AP-23, No. 6, Nov. 1975.)
 - 26 R.N. McDonough and W.H. Huggins, "Best Least-Squares Representation of Signals by Exponentials," IEEE Trans. AC-13, No. 4, pp. 405-412, August 1968.

Tuttle [27] extended the Mth order difference equation encountered in Prony's method to an Mth order differential equation but then confined attention to the single point $t = 0$. Kautz [28] then extended the technique to the semi-infinite interval $[0, \infty)$. If a continuous function x_t is a sum of complex exponentials, then it can be shown that the various derivatives of x_t satisfy a constant coefficient homogeneous differential equation

$$\sum_{i=0}^{M-1} a_i x_t^{(i)} = 0 \quad \text{with } a_0 = 1 \quad (4.1)$$

where $x_t^{(i)}$ is the i th derivative with respect to t of x_t . But if the data are noisy, then the right-hand side of the above equation is no longer zero. We write

$$\sum_{i=0}^{M-1} a_i x_t^{(i)} = e_t \quad (4.2)$$

where e_t is the error term. Kautz then proceeded to solve for $\{a_i\}$ by trying to reduce the error

$$I = \int_0^{\infty} [e_t]^2 dt \quad (4.3)$$

The exponents $\{s_i\}$ used for fitting x_t are then obtained from the zeros of the characteristic equation

$$\sum_{i=0}^{M-1} a_i s^{-i} = 0 \quad (4.4)$$

-
- 27 D.F. Tuttle, "Network Synthesis for Prescribed Transient Response," D.Sc. Dissertation, M.I.T., 1948.
 - 28 W.H. Kautz, "Approximation over a Semi-infinite Interval," M.S. Thesis, M.I.T., 1948.

The expansion of x_t as a linear combination of its derivatives is inappropriate if the data x_t are not everywhere smooth (for example when it is sampled). Carr [29] extended this technique by integration of the differential equation (4.1) k times. This leads to

$$\sum_{i=0}^{M-1} a_i x_t^{(i-k)} = 0 \quad \text{with } a_0 = 1 \quad (4.5)$$

where

$$x_t^{(i-k)} = \int_{-\infty}^t \int_{-\infty}^{t_{k-1}} \dots \int_{-\infty}^{t_1} x_{t_1}^{(i)} dt_1 dt_2 \dots dt_{k-1} dt \quad (4.6)$$

[Note: The lower limit of the integral is $-\infty$ and it is assumed that

$$x_t^{(i)}(-\infty) = 0 \quad \text{for } i = 0, 1, \dots, M-1.] \quad (4.7)$$

Again, if the data are noisy, the coefficients $\{a_i\}$ are obtained from the minimization of the function

$$\int_0^{\infty} \left[\sum_{i=0}^{M-1} a_i x_t^{(i-k)} \right]^2 dt \quad (4.8)$$

The exponents s_i are obtained as before from the solution of the polynomial equation (4.4).

It is interesting to observe that this approach is very similar to the pencil-of-function method as discussed in section VI. For $k = M-1$, it is obvious that the data x_t is an element of a Hilbert space spanned by the data and its successive integrals [30].

-
- 29 J.W. Carr, "An Analytic Investigation of Transient Synthesis by Exponentials," M.S. thesis, M.I.T., 1949.
 - 30 M.J. Narasimha et al, "A Hilbert Space Approach to Linear Predictive Analysis of Speech Signals, Tech. Report 3606-10, Radioscience Lab, Stanford Electronics Lab, Stanford University, California, 1974.

The solution to the system identification problem by any difference equation leads to a regularized ill-posed problem which has been regularized in terms of how close the actual and the measured responses are instead of the natural frequencies of the waveform. This is described in detail in the next section.

A major objection with all these techniques is that there seems to be no compelling reason for the choice of $a_0 = 1$ rather than $a_{M-1} = 1$, or for that matter the choice of any a_i to be unity. Yet each of these choices leads to a different set of exponents [13]. As an example consider the signal $\exp[-0.0035t] \cos(0.25t)$. For this signal we have $s_{1,2} = \exp[-0.0035 \pm j0.25] = 0.97 \pm j0.25$. Let the length of data points be $K = 5$ and the order of the filter $M = 2$. Then $\{x_t\}$ is the sequence $\{1.00, 0.97, 0.88, 0.73, 0.54\}$ for the sequence $\{t = 0, 1, 2, 3, 4\}$. To this we add 1% zero mean additive noise to obtain the sequence $\{x_t^{\text{Noise}}\} = \{1.01, 0.96, 0.89, 0.72, 0.54\}$. The $s_{1,2}$ for this signal are obtained from the two roots of the polynomial equation $a_0 s^2 + a_1 s + a_2 = 0$. If one assumes $a_0 = 1$, then the computed $s_{1,2} = 0.87 \pm j.23$. If one assumes $a_1 = 1$, then the computed $s_{1,2} = 0.98 \pm j.22$. If one assumes $a_2 = 1$, then the computed $s_{1,2} = 1.1 \pm j.14$.

The above example thus illustrates the dependence of the exponents on the choice of a_i to be unity.

V. ILL-POSED AND WELL-POSED PROBLEMS OF SYSTEM IDENTIFICATION

The system identification problem is almost always ill-posed. (This is reflected by the fact that two impulse responses with drastically different natural frequencies may yield almost identical outputs for the same input.) An ill-posed problem can be regularized however by imposing additional constraints on the system.

We begin our discussion by defining the concept of a well-posed problem along the lines of Tykhonov [31-32] and Lavrentiev [33]. In particular, consider the operator equation

$$Xh = y \tag{5.1}$$

where the operator X maps an element in the space H to an element in the space Y . The problem of solving (5.1) for h given X and y is said to be well-posed if the following conditions are satisfied:

- 1) The solution to (5.1) exists for each element in the space Y .
- 2) The solution to (5.1) is unique in H .
- 3) Small perturbations in y result in small perturbations in the solution to (5.1) without the need to impose additional constraints.

If any of these conditions is violated, the problem is said to be

-
- 31 A.N. Tykhonov, "On the Solution of Incorrectly Formulated Problems and the Regularization Methods," Soviet Mathematics, 4, 1963 pp.1035-103
 - 32 A.N. Tykhonov, "Regularization of Incorrectly Posed Problems," Soviet Mathematics, 4, 1963, pp. 1624-1627.
 - 33 M.M. Lavrentiev, "Some Improperly Posed Problems of Mathematical Physics," Springer-Verlag Tracts in Natural Philosophy, Vol. II, Springer-Verlag, Berlin, 1967.

ill-posed. It is important to realize that uncertainty in data due to measurement error may cause a problem to become ill-posed. Specifically, this results when a noisy measurement of y produces a waveform which does not belong to the space Y .

When (5.1) was introduced, it was assumed that the operator X was known exactly. When there is uncertainty in X , in addition to uncertainty in y , the problem is said to be well-posed in the wide-sense provided condition (3) is generalized to require that the solution h depends continuously on both X and y (i.e. small perturbations in both X and y should produce only small perturbations in h). For example, in linear least-squares problems where (5.1) is a matrix equation in a finite dimensional space, the solution is given by

$$h = (X^T X)^{-1} X^T y. \quad (5.2)$$

Since the generalized inverse of a matrix does not depend continuously on its matrix elements, the problem is ill-posed in the wide sense. Interestingly enough, this problem is well-posed in the narrow sense. If the determinant of the matrix is very small or the condition number $(\|X\| \|X^{-1}\|)$ is very large, then the problem is numerically ill-conditioned [34]. Another example of an ill-posed problem is the integral equation of the first kind.

Hadamard introduced the notion of a well-posed (correct, properly posed) problem at the beginning of this century when he studied the Cauchy problem in connection with the solution of Laplace's equation [32]. He observed that the solution did not depend continuously on the data. On the basis of this, Hadamard concluded that something was

34 M.Z. Nashed, "Some Aspects of Regularization and Approximation Solutions of Ill-Posed Operator Equations," Proceedings of the 1972 Army Numerical Analysis Conference, pp. 163-181.

wrong with the problem formulation because solutions exhibiting such type of discontinuous dependence do not correspond to physical systems, i.e. they do not arise in the study of natural phenomena. Other mathematicians of that time, such as Petrovsky, also reached the same conclusion.

Mathematicians, such as Hadamard and Petrovsky, reasoned that the mathematical models associated with the ill-posed problems must be incorrect. However, today it is recognized that their definition of a well-posed problem is lacking. In fact, using that definition, many "inverse" problems of mathematical physics are ill-posed. This includes most radiation and scattering problems in antenna theory.

In order to avoid difficulties associated with the original definition, Tykhonov suggested that the three conditions be restated differently. In addition to the metric spaces H and Y and the operator X , let there be given some closed set $H_c \subset H$. We call the problem for the solution of (5.1) properly posed according to Tykhonov if the following conditions are fulfilled:

- 1) It is required that the solution h exists for some class of data y and belongs to the given set H_c , $h \in H_c$.
- 2) The solution is unique for the class of solutions belonging to H_c .
- 3) Arbitrarily small changes of y which do not carry the solution outside the metric space H_c correspond to arbitrarily small changes in the solution h .

We denote by H_c^A the image of H_c after application to the space H of the operator X . Requirement 3) can now be restated in the following manner,

3) The solution of equation (5.1) depends continuously on the right-hand side y which is a member of the set H_C^A .

If H_C^A is a compact set, the following statement holds. If equation (5.1) satisfies the requirements 1), 2) of a well-posed problem due to Tykhonov, then there exists a function $\alpha(\tau)$ such that [32]

a) $\alpha(\tau)$ is a continuous nondecreasing function with $\alpha(0) = 0$.

b) For any $h_1, h_2 \in H_C$ satisfying the inequality $d(Xh_1, Xh_2) \leq \varepsilon$, then $d(h_1, h_2) \leq \alpha(\varepsilon)$

Thus the requirement of continuous dependence is satisfied if 1) and 2) are satisfied.

We note that, if a problem is properly posed according to Tykhonov and we replace the metric spaces H and Y by their subspaces H_C and H_C^A , then the problem becomes properly posed in the usual sense.

The necessity of examining spaces H , Y together with H_C , H_C^A is due to the fact that in real problems the errors committed in the determination of the right-hand side y , usually lead to y outside of H_C^A . The consideration of the problem according to Tykhonov's formulation gives the possibility of constructing an approximate solution with a certain guaranteed degree of accuracy in spite of the fact that an exact solution of (5.1) with approximate data either does not exist at all or may strongly deviate from the "true" solution.

The new set of three conditions may be summarized as follows. The first condition guarantees the existence of X^{-1} in the sense that a solution may proceed by choosing a complete basis from the compact set H_C in order to project y and Xh (for $h \in H_C$) into H_C^A . The uniqueness of the solution is guaranteed by condition two. Condition three requires the continuity of the solution in the space H_C .

Given an ill-posed problem, Tykhonov has regularized the problem by redefining what is meant by an acceptable solution. Basically, the idea is to make constructive use of the notions we have with regard to a physical problem by which we determine a certain class of acceptable answers having more-or-less acceptable magnitudes and degrees of smoothness. Regularization of an ill-posed problem need not be confined to the method of Tykhonov. Various schemes proposed in the literature have involved one or more of the following concepts: [34]

- a) a change in the definition of a solution
- b) a change in the space to which the solution belongs
- c) a change of the operator X
- d) the introduction of regularizing operators
- e) probabilistic methods or well-posed stochastic extensions of ill-posed problems.

Note that it may be possible to regularize an ill-posed problem with respect to one set of variables but not another. Thus, the choice of piecewise triangles or piecewise sine functions as a basis for expanding the current distribution on an antenna by the method of moments results in a regularized ill-posed problem with respect to the current distribution on the antenna structure. However, the problem is not regularized with respect to charge because the charge distribution obtained in this manner is discontinuous. As a point of interest, the method of moments regularizes an ill-posed scattering or a radiation problem by the introduction of concepts (b) and (c).

In a system identification problem, the objective is to find the impulse response $h(t)$ of a linear time-invariant system when the input

$x(t)$ and the output data $y(t)$ to a system are known. The input-output relationship for a causal system is described by

$$y(t) = \int_0^t x(t-\tau)h(\tau)d\tau = Xh \quad (5.3)$$

In problems arising in system identification we are usually certain of the existence of the function $h(t)$ that appears in the integrand in equation (5.3). Its uniqueness can also be guaranteed. However even if the solution exists and is unique, eq. (5.3) can have for a specific X a peculiarity which makes the problem an incorrectly posed one. This peculiarity arises from the "smoothing" action of the convolution operator X . This is illustrated with the following example.

Consider two continuous functions $h_1(t) = h(t)$ and $h_2(t) = h_1(t) + C \sin \omega t$. It is clear that even for a very large value of C , we can choose sufficiently a large value of ω such that the difference between $y_1 = Xh_1$ and $y_2 = Xh_2$ is less in absolute value than any previously given (arbitrarily small) number ϵ , i.e the operator X "smoothes" out a very intense, but adequately high-frequency component, to an extremely small level. The presence of disturbances accompanying the function $y(t)$ makes the problem ill-posed. For instance, assume that experimental conditions permit agreement of the measurement $d(t)$ with the measured function $y(t)$ only to within an error δ

$$\max_{0 \leq t < \infty} \left| d(t) - y(t) \right| \leq \delta. \quad (5.4)$$

It is easy to see that if the operator X of (5.3) has the smoothing action we have described, then we can always find two functions $h_1(t)$ and $h_2(t)$ whose transforms $y_1 = Xh_1$ and $y_2 = Xh_2$ both satisfy (5.4). Accordingly, there are at least two different functions that satisfy (5.3) with $d(t)$

measured by experiment to within an error δ . In fact, there exists an infinite set of such functions, whose members may differ from each other by as much as we please. It is in this situation that the "incorrectness" of the problem formulation 3.1 actually lies.

In antenna problems, the situation is not that severe due to the highly peaked shape of the kernel $\exp(-jkr)/r$.

Suppose now the measured function is $d(t) = y(t) + e(t)$ where $e(t)$ is the noise of the system. The noise is assumed to be a stationary random process with zero mean and correlation function $\phi(\tau)$.

The solution of the problem

$$y(t) = \int_{-\infty}^{\infty} x(t-\tau) h(\tau) d\tau \quad (5.5)$$

is formally obtained in terms of Fourier transforms of the output and input by means of the expression

$$h(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{j\omega t} \frac{\tilde{y}(\omega)}{\tilde{x}(\omega)} d\omega \quad (5.6)$$

where the symbol $\sim$ denotes the Fourier transform of the corresponding function. Of interest is the variance of the function $h(t)$ when instead of y we use $d = y + e$ in (5.5). The variance is derived in [35] as

$$\sigma^2(h) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \frac{\tilde{\phi}(\omega)}{|\tilde{x}(\omega)|^2} d\omega \quad (5.7)$$

where $\tilde{\phi}(\omega)$ is the power spectrum of the noise. Note for finite energy signals that

$$|\tilde{x}(\omega)| \rightarrow 0 \quad \text{for} \quad |\omega| \rightarrow \infty. \quad (5.8)$$

In order that the variance of the solution remains finite, the power spectrum $\tilde{\phi}(\omega)$ of the noise must fall off sufficiently rapidly as $|\omega| \rightarrow \infty$.

35 V.F. Turchin et al, "The Use of Mathematical Statistics Methods in the Solution of Incorrectly Posed Problems," Soviet Physics Uspekli 19, 1971, pp. 681-703.

This imposes severe restrictions on the class of processes $e(t)$ that are admissible as noise. Commonly, these conditions are not satisfied; since the noise is usually assumed to contain a background "white noise" component. Consequently as $|\omega| \rightarrow \infty$, the spectrum $\tilde{\phi}(\omega)$ approaches a nonzero constant limit. Then the variance in (5.7) is infinite. Hence, unsatisfactory solutions are obtained when the experimentally found function $d(t)$ is substituted for $y(t)$ in (5.5). The source of the difficulty is obvious: the high frequency components of $d(t)$, which arise from the presence of noise and which are not present in the true function $y(t)$, produce large oscillations in the solution.

It is useful to examine the situation needed for (5.5) to be a correctly posed problem in the presence of white noise. In particular, if $\tilde{x}(\omega)$ is a rational function, the numerator polynomial must be of higher degree than the denominator polynomial. This requires in $x(t)$ the presence of singularity functions such as doublets, triplets, etc., all of which have infinite energy.

The classical Wiener problem is also ill-posed. The solution is determined from the orthogonality principle, which states that the linear minimum mean square error estimator is chosen to make the error orthogonal to the data. However, the solution is not unique because, in general, other estimators which are also orthogonal to the data can be added to the solution without upsetting the orthogonality condition.

Maximum likelihood, minimum variance, and maximum entropy spectral estimation regularize what would otherwise be an ill-posed problem by using statistical techniques to estimate the solution as opposed to solving for

an exact solution. The main feature of a statistical regularization scheme is that an estimation "rule" is prescribed for the observed data. Given the noisy data, one simply applies the estimation "rule" in order to achieve the estimate. The quality of the estimate depends on the goodness of the estimation rule chosen as well as the accuracy of the a priori knowledge concerning the statistics of the underlying process. This leads to the replacement of the exact solution of the equation by an approximate "regularized" solution. Different strategies, both optimal and suboptimal, may be suitable for different problems. However, they all result in a statistical regularization of the problem and, in general, yield estimates of varying quality.

One disadvantage of the statistical regularization approach is that considerable a priori information is usually needed if a particular strategy is to be successfully applied. Nevertheless, the following advantages hold:

First, the probabilistic approach is the natural way to describe measurement noise which is often responsible for a problem becoming ill-posed.

Second, the probabilistic method allows more complete use of previous experience, by including it in the a priori distributions.

Third, when there is no such experience, the probabilistic method still allows one to proceed by making use of extremely weak assumptions about the unknown processes.

The various techniques discussed in section III describe various strategies to regularize in a statistical way the ill-posed system identification problem.

The pencil-of-functions regularizes the identification problem by generating the compact set H_c to which the solution belongs. This is achieved by introduction of the operator S which integrates a function from ∞ to t . For a discrete-time system the integral reduces to a sum. The pencil-of-function method makes use of the simple sequence

$$\eta = \{ A_i r_i^k \}$$

The sequence is indexed on k and $r_i = \exp(j2\pi f_i)$ and $\{A_i\}$ are the residues. Application of the operator S on η reduces to

$$S\eta = \{ A_i r_i^k / (1 - r_i) \}.$$

It follows that $[1 - (1 - r_i)S]\eta = \{0\}$. It can also be shown that the operator S maps the space onto itself while preserving the poles of the sequence. It is because of this factor that the poles and zeros obtained by this method are extremely stable, reliable, consistent and accurate.

The vectors which span the compact set H_c to which the pole r_i belongs are generated by successive applications of the operator $[1 - (1 - r_i)S]$ to η . Note that in this operator r_i is an unknown quantity. The r_i 's are obtained from the linear dependence of the spanning vectors of the compact set H_c . The detailed mathematical derivations may be obtained in [36].

VI. PENCIL-OF-FUNCTIONS METHOD

A useful mathematical entity arises by combining two given functions defined on a common interval together with a scalar parameter as

$$f(t, \lambda) = \lambda g(t) + h(t) \quad (6.1)$$

The entity f is called a pencil of functions where $g(t)$ and $h(t)$ are parameterized by λ . For example, if $h(t)$ is composed of

$$h(t) = A \exp(-st)$$

and

$$g(t) = \int_0^t h(t) dt = \frac{-A}{s} \exp(-st) \quad [\text{for } s > 0]$$

then the pencil $\lambda g(t) + h(t)$ can be formed. The pencil of functions is linearly dependent only when $\lambda = -s$. Therefore, the value of λ can be computed from $h(t)$ and its integral using their inner product. The main result thus concerning the linear dependence of the pencil sets is that the parameter λ satisfy a polynomial equation. The details of this method may be found in references [36-38]. An advantage of this technique is the generation of the successive

-
- 36 V.K. Jain, "On System Identification and Approximation," Florida State University, Tallahassee, Eng. Res. Rep., SS-I 1, 1970.
 - 37 V.K. Jain, "Filter Analysis by Use of Pencil of Functions: Part I & II," IEEE Trans. on Circuits and Systems, Vol. CAS-21, No. 6, September 1974.
 - 38 T.K. Sarkar et al, "Suboptimal System Approximation/Identification with known Error," Report No. AFWL-TR-77-200 on Contract F33615-77-C-2059, September 1977.

integrals of the function. So assuming the function itself is in L_2 space, then the set generated by the successive integrals forms a compact set in L_2 space. This is how the ill-posed system identification problem is regularized by converting the function to L_2 . Since we are interested first in finding the values of λ for the pencil, generation of successive integrals of the function forces the solution to belong to the compact set spanned by the integrals. That is why it has been possible to estimate an error bound on the location of the poles [36-38].

Another added advantage of formulating the problem this way is that the effect of conventional filtering can be greatly reduced. This can be achieved through successively smoothing the function by passing it through a band pass filter with transfer function, $[(as + b)/(cs + d)]$, as opposed to pure integration. The integrator is then a special case of a band pass filter for $a = d = 0$ and $b = c = 1$. This can increase the frequency resolution of the identification technique.

Moreover, as the poles are obtained from a polynomial equation whose coefficients form the minors of the Grammian of the pencil of functions, noise corrections can be done easily. Thus, in order to make the estimate of the poles unbiased, the entries in the gram-matrix can be altered in a systematic way to yield an unbiased estimate for the poles [25,26].

As an example consider the transient response of a conducting pipe tested at the ATHAMAS-I EMP simulator. The conducting pipe is 10 m long and 1 m in diameter. Hence, the true resonance of the pipe is expected to be in the neighborhood of 14 MHz. Also, the pipe has

been excited in such a way that it is reasonable to expect only odd harmonics at the scattered fields. The data which have been measured are the integral of the E-field and hence is available in terms of a voltage. Thus, in addition to the frequencies of the conducting pipe one should also observe a very dominant low frequency pole. The same transient data as depicted in Figure 2 {Figure 10 of reference [38]} is used for analysis. The results for a fifth and a seventh order system are as follows:

For $n = 5$, the poles in radians/sec are

$$\begin{array}{ll} (-0.0029 \mp j0.083) \times 10^9 & (=13.33 \text{ MHz}) \\ (-0.0428 \mp j0.217) \times 10^9 & (=35.20 \text{ MHz}) \\ (-0.0098 \quad \quad \quad) \times 10^9 & (= 1.56 \text{ MHz}) \end{array}$$

For $n = 7$, the poles in radians/sec are

$$\begin{array}{ll} (-0.0058 \mp j0.084) \times 10^9 & (=13.40 \text{ MHz}) \\ (-0.0270 \mp j0.219) \times 10^9 & (=35.10 \text{ MHz}) \\ (-0.0270 \mp j0.550) \times 10^9 & (=87.60 \text{ MHz}) \\ (-0.0012 \quad \quad \quad) \times 10^9 & (=0.19 \text{ MHz}) \end{array}$$

It is interesting to observe that the real pole due to the integrator has been obtained. This pole is a very dominant pole as the data was recorded after having passed through an integrator. The above results display a dynamic range of approximately 1000:1 for the values of poles of the conducting pipe.

Next the data were differentiated to get rid of the undesirable dominant pole of the integrator. The differentiation was done numerically. For a fourth and a sixth order system the above results have been

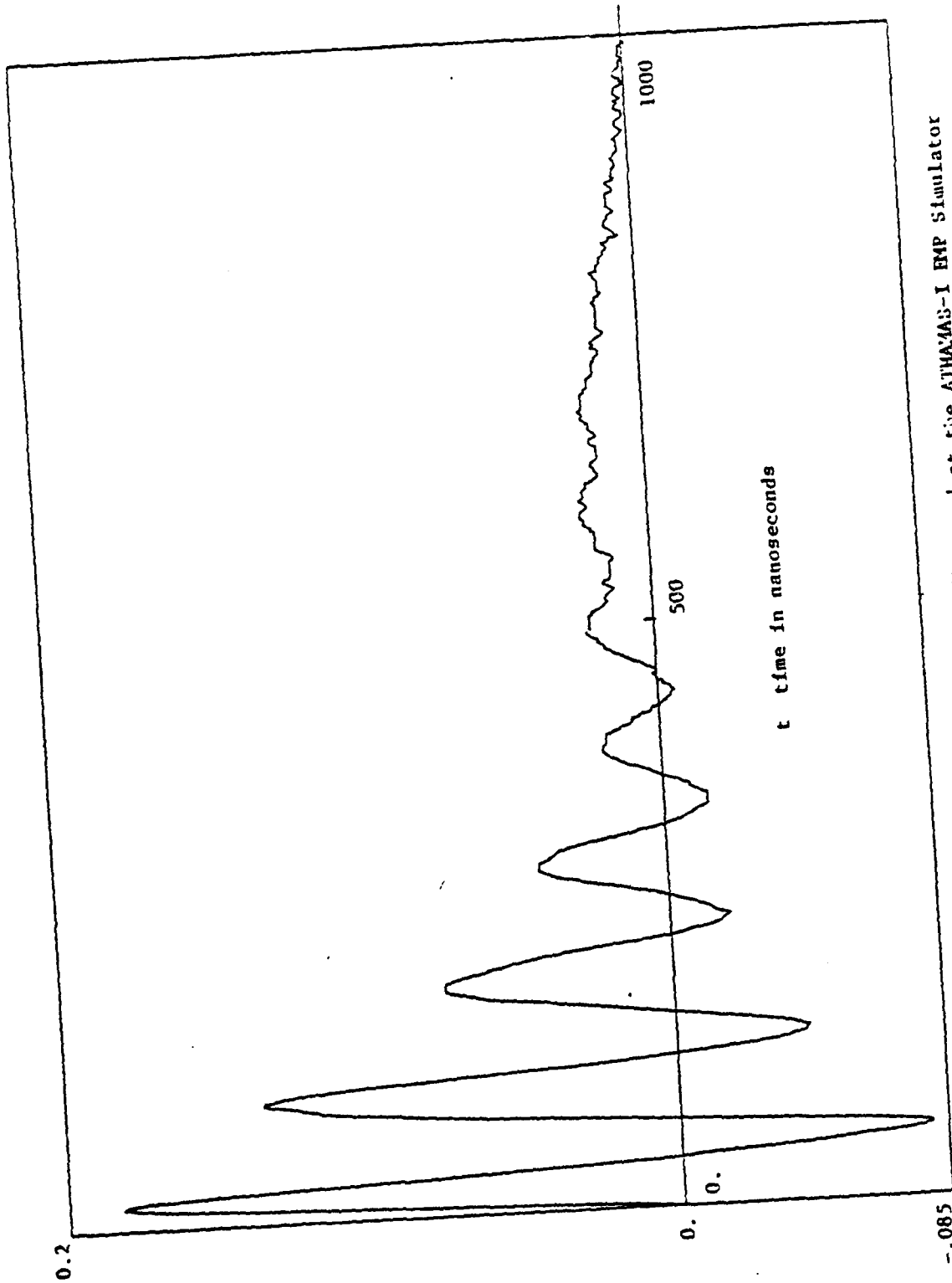


Fig. 2. Transient Response of a Conducting Pipe Measured at the ATHAS-I EMP Simulator

recalculated as follows:

For $n = 4$, the poles in radians/sec are

$$(-0.0026 \pm j0.086) \times 10^9 \quad (=13.70 \text{ MHz})$$

$$(-0.0480 \pm j0.235) \times 10^9 \quad (=37.47 \text{ MHz})$$

For $n = 6$, the poles in radians/sec are

$$(-0.005 \pm j0.083) \times 10^9 \quad (=13.23 \text{ MHz})$$

$$(-0.034 \pm j0.221) \times 10^9 \quad (=35.59 \text{ MHz})$$

$$(-0.071 \pm j0.406) \times 10^9 \quad (=65.9 \text{ MHz})$$

Here a good approximation to the poles has been obtained with only four poles. Also, there seems to be a good agreement in the pole locations obtained from the original integrated data and the numerically differentiated data. It is also interesting to observe that indeed the poles are occurring approximately at odd harmonics of the fundamental. Hence, the pencil-of-functions method does provide stable, reliable, consistent and accurate values of poles from noise contaminated measured responses of electromagnetic systems.

In Figure 3, the true numerically differentiated data is plotted against the reconstructed response of a sixth order system. The plot has been normalized to unity amplitude. It is interesting to note that there is a close agreement even in the very early times of the two waveforms.

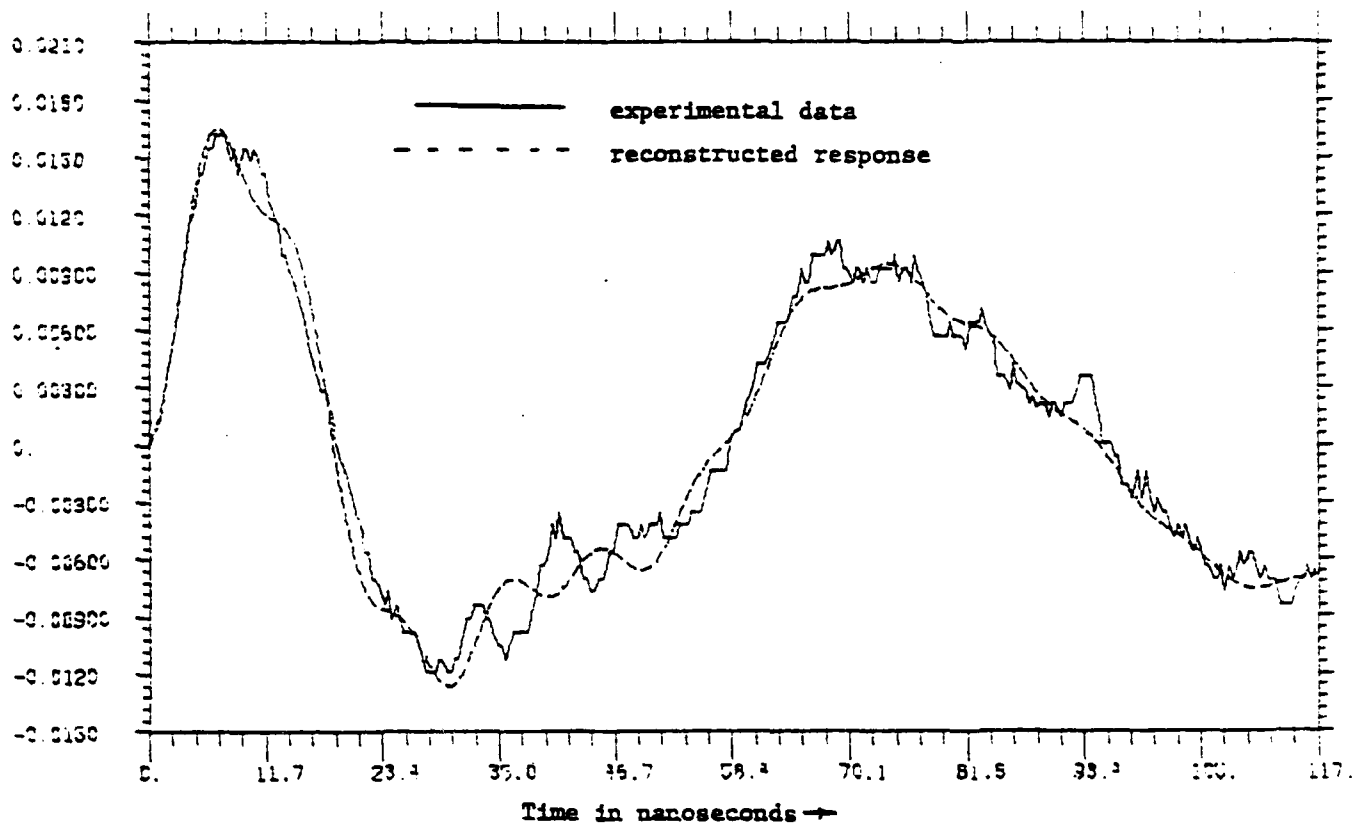


Fig. 3. True Response Vs. Reconstructed Response of a Sixth Order System for a 10 m Long 1 m Diameter Conducting Pipe.

VII. DISCUSSIONS

The previous sections demonstrate that the use of a proper mathematical model is extremely important in regularizing an ill-posed problem. It has also been shown that the pencil-of-functions differs radically from the other existing schemes of finding poles and residues of a finite length noise contaminated record. It is because of the use of a completely different regularizing scheme that analytical error bounds on the pole locations are possible. This is why reliable, consistent, stable and accurate results for the poles and residues have been obtained for this method. Thus, the pencil-of-functions method shows a great promise for the analysis of poles and residues from measured transient responses of a finite-size conducting body in free space.

REFERENCES

1. J.D. Markel, "Formant Trajectory Estimation from a Linear Least-Squares Inverse Filter Formulation," SCRL-Monograph 7, Speech Communications Research Laboratory, Inc., Santa Barbara, Oct. 1971.
2. J.D. Markel and A.H. Gray, Linear Prediction of Speech. Berlin: Springer-Verlag.
3. E.A. Robinson and S. Treital, "Principles of Digital Wiener Filtering," Geophysical Prospecting, September 1967, pp. 311-333.
4. N. Levinson, "The Wiener R.M.S. error criterion in Filter Design and Prediction," Journal of Mathematics and Physics, Vol. 26, 1947, pp. 261-278.
5. J. Makhoul, "Linear Prediction: A Tutorial Review," Proc. IEEE, Vol. 63, No. 4, April 1975, pp. 561-580.
6. J.D. Markel, "Digital Inverse Filtering - A New Tool for Formant Trajectory Estimation," IEEE Trans. on Audio and Electroacoustics, Vol. AU-18, No. 2, June 1970, pp. 137-141.
7. K.L. Peacock and S. Treital, "Predictive Deconvolution: Theory and Practice," Geophysics, Vol. 34, No. 2, April 1969, pp. 155-169.
8. A. Oppenheim and R. Schaefer, Digital Signal Processing. Englewood Cliffs, N.J.: Prentice Hall, Inc., 1975.
9. J.L. Shanks, "Recursion Filter for Digital Processing," Geophysics, Vol. XXXII, No. 1, February 1967, pp. 33-51.
10. D.L. Fletcher and C.N. Weygandt, "A Digital Method of Transfer Function Calculation," IEEE Trans. on Circuit Theory, January 1971, pp. 185-187.
11. C.S. Burrus and T.W. Parks, "Time Domain Design of Recursive Digital Filters," IEEE Trans. on Audio and Electroacoustics, Vol. AU-18, No. 2, June 1970, pp. 137-141.
12. S. Treital and E.A. Robinson, "The Design of High Resolution Digital Filters," IEEE Trans. on Geoscience Electronics, Vol. GE-4, No. 1, June 1966, pp. 25-38.
13. R.N. McDonough, "Representation and Analysis of Signals, Part XV - Matched Exponents for the Representation of Signals," Johns Hopkins University, April 1963.

14. A. Gelb, Applied Optimal Estimation. M.I.T. Press, 1974.
15. H.W. Sorensen, "Least-Squares Estimation: from Gauss to Kalman," IEEE Spectrum, July 1970, pp. 63-68.
16. F. Itakura and S. Saito, "Analysis Synthesis Telephony Based on Maximum Likelihood Method," 6th International Congress on Acoustics, Tokyo, Japan, August 21-26, 1968, C-5-5, pp. C-17-20.
17. F. Itakura and S. Saito, "Extraction of Speech Parameters Based upon the Statistical Method," Proc. Speech Info. Process, Tohoku University, Sendai, Japan, 5.1, 5.12, 1971 (in Japanese).
18. J.P. Burg, "Maximum Entropy Spectral Analysis," Ph.D. Dissertation, Stanford University, Palo Alto, California, 1975.
19. C.E. Shannon and W. Weaver, The Mathematical Theory of Communication. Urbana, Illinois: University of Illinois Press, 1962, pp. 56-57.
20. R.N. McDonough, "Maximum-entropy Spatial Processing of Array Data," Geophysics, Vol. 39, No. 6, December 1974, pp. 843-851.
21. A. Van den Bos, "Alternative Interpretation of Maximum Entropy Spectral Analysis," IEEE Trans. on Inf. Theory, Vol. IT-17, no. 4, 1971, pp. 493-494.
22. D.E. Smylie, G.K.C. Clarke and T.J. Ulrych, "Analysis of Irregularities in the Earth's Rotation," in Methods in Computational Physics, Vol. 13, Academic Press, New York, 1973, pp. 391-430.
23. S.L. Marple, "Conventional Fourier, Autoregressive and Spectral Methods of Spectral Analysis," Ph.D. Dissertation, Stanford University, Palo Alto, California, 1976.
24. C. Lanczos, Functional Analysis. Englewood Cliffs, N.J.: Prentice-Hall, Inc., pp. 272-279.
25. M. Van Blaricum and R. Mittra, "Techniques for Extracting the Complex Resonances of a System Directly from its Transient Response," Interaction Note 301, December 1975 (Also in IEEE Trans. on Ant. and Propagation, Vol. AP-23, No. 6, November 1975).
26. R.N. McDonough and W.H. Huggins, "Best Least-Squares Representation of Signals by Exponentials," IEEE Trans. AC-13, No. 4, pp. 405-412, August 1968.
27. D.F. Tuttle, "Network Synthesis for Prescribed Transient Response," D.Sc. Dissertation, M.I.T., 1948.
28. W.H. Kautz, Approximations over a Semi-infinite Interval," M.Sc. Thesis, M.I.T., 1948.

29. J.W. Carr, "An Analytic Investigation of Transient Synthesis by Exponentials," M.S. Thesis, M.I.T., 1949.
30. M.J. Narasimha et. al, "A Hilbert Space Approach to Linear Predictive Analysis of Speech Signals," Tech. Report 3606-10, Radioscience Lab., Stanford Electronics Lab., Stanford University, California, 1974.
31. A.N. Tykhonov, "On the Solution of Incorrectly Formulated Problems and the Regularization Methods," Soviet Mathematics, 4, 1963, pp. 1033-1038.
32. A.N. Tykhonov, "Regularization of Incorrectly Posed Problems," Soviet Mathematics, 4, 1963, pp. 1624-1627.
33. M.M. Lavrentiev, "Some Improperly Posed Problems of Mathematical Physics," Springer-Verlag Tracts in Natural Philosophy, Vol. II, Springer-Verlag, Berlin, 1967.
34. M.Z. Nashed, "Some Aspects of Regularization and Approximation Solutions of Ill-posed Operator Equations," Proceedings of the 1972 Army Numerical Analysis Conference, pp. 163-181.
35. V.F. Turchin et. al., "The Use of Mathematical Statistics Methods in the Solution of Incorrectly Posed Problems," Soviet Physics Uspekhir 13, 1971, pp. 681-703.
36. V.K. Jain, "On System Identification and Approximation," Florida State University, Tallahassee, Eng. Res. Rep., SS-11, 1970.
37. V.K. Jain, "Filter Analysis by Use of Pencil of Functions: Part I & II," IEEE Trans. on Circuits and Systems, Vol. CAS-21, No. 5, September 1974.
38. T.K. Sarkar et. al., "Suboptimal System Approximation/Identification with Known Error," Report No. AFWL-TR-77-200 on Contract F33615-77-C-2059, September 1977.

DATE
FILME