

This dissertation has been 65-12,995
microfilmed exactly as received

HEYMANN, Michael, 1936-
A TIME DOMAIN TECHNIQUE FOR
LINEAR SYSTEM IDENTIFICATION.

The University of Oklahoma, Ph.D., 1965
Engineering, chemical

University Microfilms, Inc., Ann Arbor, Michigan

THE UNIVERSITY OF OKLAHOMA
GRADUATE COLLEGE

A TIME DOMAIN TECHNIQUE FOR LINEAR
SYSTEM IDENTIFICATION

A DISSERTATION
SUBMITTED TO THE GRADUATE FACULTY
in partial fulfillment of the requirements for the
degree of
DOCTOR OF PHILOSOPHY

BY
MICHAEL HEYMANN
Norman, Oklahoma

1965

A TIME DOMAIN TECHNIQUE FOR LINEAR
SYSTEM IDENTIFICATION

APPROVED BY

C. M. Shepewich

George M. Ewing

J. D. Roberts

Hung-Ta Ho

M. & M. F. J. J.

DISSERTATION COMMITTEE

ABSTRACT

Several serious limitations have been encountered in the past with the common frequency domain techniques for chemical process identification methods. These include such methods as direct frequency response testing, random testing and pulse testing. Most of the difficulties are associated directly with noise problems and with the fact that frequency domain methods do not permit accurate determination of the system's parameters. Furthermore, none of the available identification methods determines the reliability of the obtained model.

The major goal of this work was to devise a modeling technique that simultaneously evaluates the system parameters along with their associated uncertainties. The proposed technique for linear system identification is based exclusively on time domain analysis. Fundamental mathematical relations of linear differential equations are utilized for the identification and an extensive error propagation analysis permits the determination of the uncertainty associated with each of the system parameters.

The procedure is based on analysis of transient output responses of the process to be identified. Arbitrary input pulses can be utilized to force the system from steady state.

For identification of an n^{th} order model, n linearly independent response functions are required. It is assumed that the order of the system to be identified is either known or strongly suspected beforehand. The correct order of the system is then determined through proper interpretation of the error propagation analysis.

The method was investigated by analog computer simulation studies of second and third order systems with and without superimposed noise. Also, actual chemical plant test data were utilized for identification.

The identification procedure was programmed and carried out on a digital computer. Medium size digital computers are sufficient for effective execution of the analysis. Computational time requirements are insignificant.

The time domain identification technique was found to be superior to conventional methods in several respects: (1) Accurate determination of natural frequencies was obtained for systems with natural frequencies ranging from very closely separated to ratios up to 1000! (2) Noise problems were not very severe and did not generally corrupt the results significantly. (3) The error propagation analysis successfully determined the uncertainty associated with the models.

ACKNOWLEDGMENT

The author wishes to take this opportunity to express his sincere gratitude and appreciation to his advisers, Dr. C. M. Sliepcevich and Dr. M. L. McGuire, for their advice and stimulation given throughout this work.

Likewise the author expresses his appreciation to his colleague, Mr. R. H. Luecke, who contributed to this research through many stimulating discussions. Dr. K. A. Bishop's work in the Process Dynamics Laboratory at the University of Oklahoma inspired the initiation of this study. The assistance of other personnel of the Process Dynamics Laboratory is also sincerely appreciated.

Thanks are due to Mr. Earl Gray, Dr. W. S. Stewart and their associates from the Phillips Petroleum Company for their cooperation in providing the experimental plant test data, and for the many valuable discussions.

For the successful completion of this investigation a great deal of high speed digital computation was required. The assistance of the entire staff of the Osage Computer Center is sincerely appreciated. In particular, the author extends his thanks to Mr. Paul Johnson, for his patient assistance and valuable advice and to Mr. B. K. Smith for the many hours of assistance, given so readily, in

operation of the computer.

The financial support of the National Science Foundation and of the Phillips Petroleum Company is gratefully acknowledged.

Finally, I wish to express my deep gratitude to my wife, Alima, for her patience and understanding which made this work possible.

Michael Heymann

TABLE OF CONTENTS

	Page
LIST OF TABLES	ix
LIST OF ILLUSTRATIONS	x
 Chapter	
I. INTRODUCTION	1
II. THEORY	19
General Considerations Derivation of the Identification Technique Perturbation Models and Steady State Operation Preliminary Evaluation of the Technique	
III. ERROR PROPAGATION ANALYSIS	58
Error Bounds of the Inverse of $\Phi_{n-1}(t_{n-1}-\delta)$ Error Bounds of $e^{\delta A}$ Error Bounds of the Eigenvalues of $e^{\delta A}$ Error Bounds of the System's Poles Discussion	
IV. ERROR ESTIMATION FOR INTEGRALS OF RESPONSE CURVES	87
Continuous Analysis Discrete Analysis	
V. FORMALIZATION OF THE TECHNIQUE	133
VI. ANALOG COMPUTER STUDIES	141
Analog Computer Simulations Procedure of Simulation Studies	
VII. DISCUSSION OF RESULTS	152
Identification of Second Order Systems with Real Poles Identification of Second Order Systems with Complex Poles	

Chapter	Page
Effect of Noise and Error in Response Data on Quality of Identification	
Effect of Error Estimation on Quality of Identification	
Methods for Securing Linear Independent Responses and Effects of Weak Linear Independence	
Effect of Utilization of Higher Order Φ Matrices	
Identification of Third Order Systems	
The Problem of Misidentification	
Identification of Actual Chemical Plant Test Data	
VIII. CONCLUSIONS AND RECOMMENDATIONS	198
Conclusions	
Range of Applicability	
Effect of Noise	
Determination of System's Order	
Comparison to Frequency Domain Identification Methods	
Limitations	
Recommendations for Future Work	
REFERENCES	202
APPENDICES	
A. NOMENCLATURE	207
B. GENERAL MATHEMATICAL RELATIONSHIPS FOR IDENTIFICATION OF LINEAR SYSTEMS WITH PERIODICALLY TIME VARYING PARAMETERS	216
C. DERIVATION OF EQUATION (4-42) FOR THE MEAN ZERO CROSSING INTERVAL OF BAND PASS NOISE	222
D. DERIVATION OF EQUATION (4-45) FOR THE RATIO BETWEEN THE STANDARD DEVIATION AND THE MEAN ZERO CROSSING INTERVAL OF BAND PASS NOISE	225
E. COMPUTER PROGRAM LISTINGS FOR IDENTIFICATION PROCESS	229

LIST OF TABLES

Table	Page
1. Results of Analog Computer Studies of Error Propagation into Integrals of Response Curves .	125
2. Example of Potentiometer and Gain Settings for Third Order System	149
3. Final Results of Identification of Second Order Systems with Real Poles	164
4. Identification of a Second Order System with Identical Poles ($\xi=1$): $\rho_1=\rho_2=0.1[1/\text{sec}]$. .	169
5. Final Results of Identification of Second Order Systems with Complex Poles	171
6. Identification of Chemical Plant Test Data (Based on Response Records of $y_1(t)$ in Figure 27)	194
7. Identification of Chemical Plant Test Data (Based on Response Records of $y_2(t)$ in Figure 27)	195

LIST OF ILLUSTRATIONS

Figure	Page
1. General Chemical Process	20
2. Scalar Linear Process	41
3. Scalar Linear Process with General Input Function	42
4. Error Propagation for Real Poles as a Function of $(-\delta\rho_1)$: Normalized Plot of Equation (3-54) .	83
5. Error Propagation for Real Poles as Function of δ	84
6. Typical Noisy Response Function	98
7. Typical Cosine-Exponential Autocorrelation Function	109
8. Typical Spectral Density Corresponding to a Cosine Exponential Correlation Function	109
9. Comparison Between Spectral Densities of Cosine Exponential and Band-Pass Approximation Noise Models	110
10. Typical Zero Crossing Distribution for Band-Pass Noise Model Computed by Means of $P(\tau) =$ $r''(\tau)/4\beta$ as Compared to $P^{(2)}(\tau)$ Proposed by Longuent-Higgins. Experimental Results of Raniel Plotted for Comparison	113
11. Numerical Solution of Equation (4-45)	115
12. Numerical Solution of Equation (4-45) Low Range .	116
13. Analog Computer Circuit Diagram for Studies of Error Propagation in Noise Integration . . .	119
14. Sample Record of Random Function and its Integral	120
15. Comparison Between Experimental Band-Pass Approximation and Theoretical Spectral Density Functions	123

Figure	Page
16. Discrete Sampling of Random Noise	127
17. Formalized Block Diagram of the Identification Technique	137
18. Analog Computer Circuit Diagram and Digital Output Assembly for Simulation Studies	144
19. Typical Set of Linearly Independent Response Records and Forcing Pulses for Identification of a Second Order System with Real Poles: $\rho_1 = -0.002$ [1/sec], $\rho_2 = -0.1$ [1/sec]. No Superimposed Noise	153
20. Identification of Second Order System with Real Poles: $\rho_1 = -0.002$ [1/sec], $\rho_2 = -0.1$ [1/sec] No Superimposed Noise	154
21. Effect of Variation of t_{n-1} on Quality of Identification	159
22. Identification of Second Order System with Real Poles: $\rho_1 = -0.0002$ [1/sec], $\rho_2 = -0.2$ [1/sec] No Superimposed Noise	160
23. Identification of Second Order System with Real Poles: $\rho_1 = -0.0667$ [1/sec], $\rho_2 = -0.1$ [1/sec] No Superimposed Noise	162
24. Typical Set of Linearly Independent Responses and Their Integrals for Identification of Second Order System with Complex Poles: $\rho_{1,2}$ $= -0.06 \mp j 0.19078$ [1/sec]. No Superimposed Noise	165
25. Identification of Second Order System with Complex Poles: $\rho_{1,2} = -0.06 \mp j 0.19078$ [1/sec]. No Superimposed Noise	166
26. Error Propagation Characteristics of Second Order System with Complex Poles: $\rho_{1,2}$ $= -0.14 \mp j 0.142828$ [1/sec]. No Superimposed Noise	168
27. a. Sample of High Speed Noise Recording, b.&c. Noisy Responses of Second Order System with Real Poles: $\rho_1 = -0.002$ [1/sec], $\rho_2 = -0.2$ [1/sec]. Estimated Noise Level 3.4 Volt (3.4 percent)	172

28.	Identification of Second Order System with Real Poles: $\rho_1 = -0.002$ [1/sec], $\rho_2 = -0.2$ [1/sec]. Estimated Noise Level: 3.4 percent in Responses, 2.7 percent in Integrals	173
29.	Identification of Second Order System with Complex Poles: $\rho_{1,2} = -0.1 \mp j 0.1732$ [1/sec]. Effect of Noise Magnitude on Quality of Identification	175
30.	Noisy Responses of Second Order System with Complex Poles: $\rho_{1,2} = -0.1 \mp j 0.1732$ [1/sec] Estimated Noise Level: 8 volt (8 percent) . . .	176
31.	Identification of Second Order System with Real Poles $\rho_1 = -0.002$ [1/sec], $\rho_2 = -0.2$ [1/sec]. Effect of Error Estimation on Quality of Identification. No Superimposed Noise	179
32.	Identification of Second Order System with Real Poles $\rho_1 = -0.02$ [1/sec], $\rho_2 = -0.2$ [1/sec]. Effect of Error Estimation on Quality of Identification. Estimated Noise Level: 3.5% in Responses, 2% in Integrals	181
33.	Identification of Third Order System with Real Poles: $\rho_1 = -0.0054027$ [1/sec], $\rho_2 = -0.036722$ [1/sec], $\rho_3 = -0.05356$ [1/sec]. No Superimposed Noise	183
34.	Identification of Third Order System with One Real Pole $\rho_1 = -0.10636$ [1/sec] and a Pair of Complex Poles $\rho_{2,3} = -0.012517 \mp j 0.01705$ [1/sec]. No Superimposed Noise	184
35.	Identification of Third Order System with One Real Pole $\rho_1 = -0.10607$ [1/sec] and a Pair of Complex Poles $\rho_{2,3} = -0.00490459 \mp j 0.01361$ [1/sec]. No Superimposed Noise	185
36.	Identification of a Third Order Model for a Second Order System	189
37.	Plant Pulse Test Data	191

A TIME DOMAIN TECHNIQUE FOR LINEAR SYSTEM IDENTIFICATION

CHAPTER I

INTRODUCTION

Considerable progress has been made in recent years in the area of process control and dynamics. The number of recent publications of research papers in this field is remarkable and of particular note is the relative number of theoretical papers as compared to reports of experimental research. Although the impact of various developments in systems engineering by other disciplines such as electrical and aerospace engineering is still apparent, chemical engineers seem to have finally realized that the problems in chemical process control are sufficiently different from problems in other disciplines to justify original and independent research. The availability of high capacity and powerful computing machinery (digital and analog) has removed many computational barriers and enabled researchers to attack problems which were not solvable previously because of computational limitations. Particularly notable is the work done in analysis of chemical reactor dynamics and stability [1,2,4,34] and developments of new and sophisticated control philosophies such

as optimal control techniques [7] and methods which utilize the invariance principle as a design guide [6,22,41]. Also, because of increased computing capability considerable work has been done in analyzing overall plant performance and stability. Results of this nature have been obtained for distillation columns, heat exchangers and the like.

In contrast to conventional "rule of thumb" and semi-empirical methods for control system design, these new techniques are characterized by the systematic consideration of all aspects of the system dynamics and in particular system constraints and limitations. By the establishment of mathematically defined, optimal control schemes evaluation of an overall control strategy can be carried out. Although criteria for the existence of an optimum are not unique in nature and the use of engineering judgement is required in their selection as discussed by Kalman [26], a large amount of the ambiguity has been removed by utilization of the new techniques.

One of the basic requirements of advanced control design (and analysis) techniques is the complete understanding of the plant dynamic behavior, i.e. a reliable mathematical model of the plant dynamic characteristics is of fundamental importance. There exists a general trend in the use of advanced design techniques to become more and more "sophisticated" and "scientific". With conventional feedback control schemes little has to be known about the dynamics of the controlled system in order to obtain stable satisfactory control

action. However, it cannot be expected that an optimum control design will be obtained without an adequate description of the process.

One can generally state that the more information that is given about a system, the better its control can be designed. A feedback controller monitors the error in the controlled variable and acts on a "manipulative" variable to counteract the cause of error, i.e. the disturbance. However, the amount of control effort is directly related to the existing error, thereby requiring the presence of an error before any control action can be taken. This type of control can be designed effectively with little knowledge about the system dynamics. However, overdesign of the system and the controls has often been the price for this uncertainty. In the sophisticated "optimum" control schemes an effort is made to use information about the system dynamics in order to anticipate the required control action, thus being able to apply corrective action prior to the propagation of the disturbance to the controlled variable. This type of control is realized in the form of "feedforward" control. There is little evidence of any application of the novel "overall" design techniques to actual industrial plants. In fact, little has even been reported of experimental laboratory work of this nature.

There appears to be a considerable lag between theory and application [8,45]. One can speculate as to the reasons for this lag. It has been argued that a lag between theory and application always exists and is a natural phenomenon,

often even an advantage! In this particular case, however, the lag is not really between theory and application; it is rather a lag between a suggestion for application and the application itself. An important reason for this delay appears to be the uncertainty associated with plant application of the mathematical methods. Information about process behavior is required which is not readily available. Although theoretical considerations may be useful in evaluating the general topology of plant equations and thereby providing formal models, it is frequently impossible to evaluate the actual parameters numerically with any degree of confidence. Another difficulty is that even if mathematical models can be evaluated from theoretical considerations, they are, as has been mentioned by Rosenbrock [45], frequently too complicated to be handled. However, the most serious limitation is that usually the overall plant performance is not adequately understood and even a reliable formal model cannot be obtained from theoretical considerations.

Since the theoretical approach may not yield satisfactory dynamic models of plant operation, one must frequently resort to experimental methods for model evaluation (often called "system identification"). By "experimental methods" is meant the evaluation of plant dynamic behavior from actual plant test data.

Experimental data are obtained from measurements of responses of plant output (state) variables, to programmed or unprogrammed variations in input (disturbance) variables.

In addition, information about steady state operating conditions is often useful. Present day methods for system identification are quite useful for evaluating simple mathematical models of chemical plants, but their value from the control standpoint is limited due to several serious drawbacks.

Several testing and analysis techniques for system identification have been developed in the past. All testing methods are based (as has been mentioned) on the analysis of input-output relationships and differ essentially in the types of test signals applied and the mathematical techniques used for data manipulation. The basic methods can be classified as: (1) sinusoidal (frequency response) testing, (2) statistical or random testing, and (3) pulse testing and Fourier transformation. All of these methods for identification can be considered as frequency domain techniques in the sense that at some stage of the analysis the data is transformed into the frequency domain. All of the above methods depend on frequency domain representations, usually magnitude and phase lag versus frequency diagrams (Bode plots), for the final plant evaluation.

Direct frequency response testing has been widely used in the past [11,18,35,38]. It was originally adopted from electrical engineering (as have been most fundamental concepts in control and systems engineering) for utilization in the chemical plant. Although frequency response testing is natural and convenient in electrical engineering applications, it did not generally prove satisfactory in the process

industries as was pointed out by Hougen [25]. The basic virtue of the method is the simplicity with which the data can be handled. Reasonable accuracy of the results can be obtained if noise effects are minimized [18,19,49,50]. General drawbacks of the method which have been widely recognized are: (1) the tests are long and time consuming and may cause the system to deviate drastically from its normal mode of operation, thereby resulting in production of off specification products or even damaging the equipment [9,25], (2) a large number of tests over a wide range of frequencies is required regardless of the complexity of the system, (3) sinusoidal function generators may be difficult and expensive to build and operate, (4) only a limited class of plant variables can be sinusoidally forced with any expected degree of success, and at a reasonable cost, (5) system noise strongly limits the general applicability of the method since the frequency range of testing is seriously reduced and the results corrupted, and (6) even mild nonlinearities will corrupt the expected sinusoidal response [50]. The need to evaluate the system model from the Bode plots or other frequency domain representations is a serious limitation (as will be discussed later). Although work has been done and is still being published on the utilization of direct frequency response testing [31,47,48,49], it has usually been used for model verification rather than identification. Comparison of theoretical models (of varying complexity) to experimental data has been the main reason for this type of testing. Most of the

publications of direct frequency response testing report of laboratory scale work, and only a few of plant scale work. The usefulness of this method is, however, strongly questioned even on laboratory scale.

The serious limitations of direct frequency response testing for identification of chemical processes led investigators to develop alternate methods that would eliminate the difficulties encountered with this technique. A method that seems, at least theoretically, to take advantage of some of the limitations inherent in direct frequency response testing is random testing followed by statistical analysis of the resulting responses [20,46,51,54,55]. The limitations due to noise would be expected to be eliminated since random disturbances are used to force the system [19]. In fact, already existing noise could be used to advantage. Also a limited number of tests is required for the identification. Even though the tests are long this requirement would not seem to be a major limitation, since the system is not "strongly" disturbed. Particularly complex function generators are not required, which simplifies the method from the experimental viewpoint.

The random testing technique was investigated theoretically by Tierney and Homan [51] utilizing digital computer simulations and by Gallier [19,20] experimentally. Several serious limitations of the method have been reported. In particular the length of the required test records is excessive if high accuracy of results is desired. This requirement

is due to the high standard deviations of the correlation functions that are obtained from finite records. Furthermore the computational requirements are quite severe, and particularly cumbersome is the enormous amount of data handling. Tierney reports that about 2500 data points were required to obtain reliable estimates of the correlation function! Another limitation which has been recognized [20,51] is the difficulty of generating an input noise function which has a sufficiently wide, nearly-flat, power spectrum in order to obtain adequate frequency response representations over an acceptable frequency range.

Nevertheless, even with these problems it is not the analytical and computational problems that make the method inadequate. Chemical processes suffer from inherent noise in the output variables even when they are operated at steady state conditions. The noise sources are usually unknown. They can exist as random disturbances which are inherent in input variables as well as in the process parameters (which usually include some function of flow, temperature or concentration). In addition, instrument noise may sometimes be present which results in an apparent observation which is not inherent in the process. The important point is that this noise is not simply related to a specific forcing variable and its source is usually unknown. Thus, any speculation that observed output, steady state noise aids in the random identification procedure is basically incorrect. If a certain input variable is randomly pulsed, the observed output

is corrupted with the steady state noise which is not correlated with the forced input.

The signal to noise ratio will give additional problems. The amount of power that can be introduced into the system through a random forcing function is quite limited, particularly so when the output variable in question has weak coupling with the forcing variable. Thus the signal to noise ratio problem may be even more severe for random testing than for other test methods. When the problem of signal to noise ratio becomes severe it would seem reasonable to attempt to introduce more energy into the system through the forcing variable. An improvement in this respect could be the introduction of a single, powerful pulse rather than a random function. Since frequency domain analysis and representation are required, the usefulness of the method is further limited.

Pulse testing and Fourier transformation has been proposed and developed to overcome some of the serious difficulties encountered with direct sinusoidal testing. This method, which is a refinement of the transient response testing, was first introduced for analysis of chemical processes by Hougen [24,25]. With this method the forcing variable is excited by a single pulse and the input and response are recorded. Both are transformed numerically into the Fourier domain and the ratio of the transform of the response to the transform of the input gives the frequency response function for the particular input-output pair tested. The frequency response is then plotted on a Bode diagram which is a diagram of magnitude

ratio and phase lag versus frequency. Since the power spectral density of a single pulse theoretically covers the whole frequency range, it follows that a single test is adequate to evaluate the complete frequency response function.

Hougen and his associates have extensively investigated this technique [13,14,23,24,25,37,44] which has gained fairly wide acceptance [9,49,52]. Stermole [49] has studied the dynamic response of heat exchangers and Watjen [52] has investigated the dynamics of a pulsed-plate extraction column.

Naturally this method has some remarkable advantages over sinusoidal testing. Basically, one test replaces a whole series of sinusoidal tests. The computation is not too involved, and above all no special function generators for plant testing are required.

From a practical standpoint some of the above statements have to be qualified. First the quality of the calculated transfer function is strongly dependent on the quality of the spectral density function of the applied pulse. This problem has been studied by Dreike [12] who demonstrated the sensitivity of the identification to pulse shapes and durations. If the nature of the system is not known in advance, some experimentation has to be expected before a proper pulse can be applied. Thus the statement that one test is sufficient is not generally true in practice. Furthermore, at least in theory, the response has to be recorded for an infinite length of time or equivalently until steady state is again reached. In practice this requirement can take a rather

long period of time. Dreike, Hougen and Mesmer [13] have studied the problem of truncation error due to utilization of limited (finite) record segments, i.e., utilization of response records which end before returning to steady state. This truncation error and the effect of improper input pulses (which have only a limited range of reliable frequency content) manifest themselves in the resulting frequency response function in the form of a "cut-off frequency". Beyond the cut-off frequency, cycling occurs. As a result the uncertainty or scattering of the results becomes excessive, and the frequency response function becomes unreliable.

One might argue that to eliminate this problem the response should be recorded until steady state is "practically" reached. For a chemical process the time involved may be impractical, particularly if the process has long time constants, because of possible drifts in steady state operating conditions. Thus the problem of finite cutoff frequency is unavoidable. This problem often drastically limits the useful range of the frequency response function.

Noise introduces additional difficulties. Tail ends of the response curves, where the amplitudes are low, may lose their significance completely. Also other parts of the response curve may contain noticeable uncertainty. The need to rely upon these uncertain sections of the responses is a serious limitation. Since the response curve has to be numerically Fourier transformed, curve fitting of the forcing pulse and the response function is essential. The quality of

the approximating functions is a decisive factor in determining the quality of the resultant frequency response function. Some of this problem can be solved by using smoothing techniques, thereby improving the results over certain ranges. The amount of computation required is certainly not trivial and requires a digital computer [14]. However, computational length should not be considered a major limitation.

Summarizing, all three methods discussed above may have their particular advantages and drawbacks which will affect the selection of a method for application to a specific problem.

When dynamic plant testing is carried out for the purpose of model identification, one may or may not have a preconceived idea of the model topology. However, there is essentially no knowledge about the model parameters since their evaluation is the prime objective of the tests. Under these circumstances it may turn out to be very difficult to evaluate the parameters numerically from Bode plot representations, particularly when noise is present. Systems describable by a first order lag (i.e. having one pole and no zero) are easily identifiable from the frequency domain representations, while systems of higher complexity may lead to considerable difficulties. Dudnikov and Golant [15,21] have proposed a method for explicit determination of the system's parameters from frequency domain representations. It appears, however, that their method does not provide significant improvement.

Generally the useful range of the magnitude ratio plot does not exceed one to two decades on the logarithmic frequency scale when actual equipment and plant tests are performed, as can be seen from various experimental test results [47,48,49,52]. Here useful range means the frequencies on the Bode plot to the right of the first break point or, in the resonant case, to the right of a peak. Thus the identifiable natural frequencies are limited to this range of at most two decades. Also the identifiable order of the system is subject to these limitations. The phase lag versus frequency plot is usually less reliable than is the magnitude ratio plot, therefore offering little additional information. If, for example, the system is of second order (i.e. exhibits two poles) then the pole corresponding to the high frequency will not be identifiable if it is outside the "useful range". In this case the system would be identified as first order. On the other hand, if the high frequency pole is in the "useful range" it may be too close to the other pole to enable reliable numerical evaluation of either of the two, particularly if there is an uncertainty associated with the frequency response function (or the Bode plot). Irregularities in the Bode plot, which frequently exist to some extent, may drastically affect the possibility of numerical evaluation of the parameters. Further, dispersion of the experimental points in the range beyond the cutoff frequency often introduces uncertainty as to the order of the system. Naturally, for higher order systems the problems are more severe. When the

system contains complex poles evaluation of the damping ratio is especially difficult even when noise effects are minimized [12].

In the foregoing sections the limitations of present day system identification methods have been discussed. In spite of the fact that frequency domain methods are very useful as a design and analysis tool in various control problems, and therefore frequency domain representations of the system dynamics are often desired, it appears that frequency domain techniques offer no real advantage for chemical process identification. On the contrary, they seem to possess some serious limitations and disadvantages. It was the purpose of this research to develop an identification method that overcomes these difficulties.

The most significant limitations of any model identification procedure result from uncertainties associated with noise in the response data. Minimization of the noise effects and of error propagation in the numerical analysis is essential if the technique is to be successful. In a chemical plant, fluctuations in the output variables are normally present even when operated under steady state conditions. This noise can originate from two sources: (1) Instrument noise, which can be caused by sensing devices, transducing elements or under certain circumstances even by the recorders themselves, and (2) plant noise which results from external disturbances such as fluctuations in flow rates, temperatures, concentrations and the like. These external disturbances can

be in the input variables (from a dynamic viewpoint), or in plant parameters which are assumed to be constant. The point of significance is that the specific noise sources are unknown and therefore the noise must be considered as inherent in the process. Thus, error or uncertainty is associated with any response curve of the system. Since the noise usually possesses a statistically constant magnitude, it follows that the relative significance of the noise depends on the strength of the test signal. In other words, it is the signal to noise ratio which determines the degree of uncertainty associated with the test data. There is usually a practical limit to the attainable response amplitude, and the corresponding signal to noise ratio can be considered as the minimum uncertainty level.

When the system responses are noisy, there is a certain amount of error theoretically associated with the corresponding model. Stated differently, perturbation of the model within given limits does not yield responses that differ significantly enough to be distinguishable from each other when the noise level is considered. This discrepancy in the identifiable model is a theoretical uncertainty which does not take into account any error that is introduced by the numerical manipulations of a specific identification technique. It is, therefore, a lower bound for the error in the model. The amount of error introduced by a particular identification method depends on data handling and numerical procedures used. The confidence limits of the obtained model

are thus a function of the sensitivity of a particular identification technique.

The problem of model reliability is normally ignored by systems engineers who deal with process identification. This information is, however, of considerable importance in the design of advanced control systems which is based on extensive knowledge of the process dynamics. Such a measure of reliability can be given in the form of confidence limits for the evaluated system parameters. The determination of such confidence limits involves an extensive study of error propagation through the numerical procedures based on the error in the original response data.

One of the major goals of this work was to devise a modeling technique that simultaneously evaluates the system parameters along with their associated uncertainties. Any ambiguity associated with the identified model is removed and provision is made for a quantitative measure of its reliability. Knowledge of the error associated with the experimentally determined models does not eliminate the difficulties inherent with system identification in the presence of noise. However this knowledge is of significant value when these models are used for design or analysis purposes.

In the present work a technique for linear system identification, based exclusively on time domain analysis, is proposed. This technique, which is based on some fundamental mathematical relations of linear differential equations, offers many unique advantages over the previously discussed

methods. In particular, the system parameters are explicitly determined thereby overcoming difficulties associated with frequency domain techniques. In addition, implementation of this technique yields error bounds associated with each of the system parameters.

The identification procedure is based on analysis of transient output records of the process to be identified. These output records are the responses of the process relaxing from an excited state. Normally this can be achieved through pulsation of process inputs by arbitrary pulses. For identification of an n^{th} order model, n linearly independent response functions are required. In general the linear independence of the responses can be secured by pulsing different process inputs, or combinations of inputs, for each response.

It is implicit in the above discussion that the order of the system to be identified is either known or strongly suspected. Once the tests have been carried out and the responses obtained the correct order of the system can be determined through proper interpretation of the error analysis associated with the identification.

In Chapter II general background principles of linear systems are discussed, and the basic technique is derived. The error propagation analysis is presented in Chapter III, and the principles of the error minimization technique are discussed. Chapter IV presents a discussion of the error propagation through integration of the response functions.

A method is derived for estimation of the error in the integrals (based on the noise in the response function).^{*} In Chapter V the proposed identification technique is formalized and is presented in its final schematic working form. Chapters VI and VII present and discuss the results of the experimental investigation of the method.

^{*}The discussion in Chapter IV is not essential for the understanding of the principles of the identification technique. It points out, however, certain limitations and is very important for proper utilization of the method.

CHAPTER II

THEORY

Linear mathematical models constitute an adequate description of a broad class of chemical processes over certain ranges of operating conditions. These processes are generally multivariable, i.e. possess a number of inputs and outputs, and can be described by a set of first order differential equations relating the variables. Such a set of equations can be expressed in a generalized vector-matrix form, where the process inputs and outputs are considered to be vector quantities. By means of the vectorial description each of the output variables is related to all of the other variables. While this description is convenient from a mathematical viewpoint, it is frequently sufficient to have the relationship between a single output variable and the input variables only. This description is commonly called a scalar model. The identification of the vectorial formulation directly is undesirable due to the high sensitivity of the numerical manipulation of experimental data. This difficulty is minimized in identification of the scalar model because of the reduced number of parameters to be determined.

It is convenient to identify the scalar model through

an equivalent canonical description often called the state space formulation. Through this transformation process the scalar model is represented by a hypothetical vector model that permits direct evaluation of the system's parameters from the process' responses. For general time varying systems the identification procedure requires differentiation of the response functions, a numerical process which is extremely sensitive to errors. For constant parameter systems the canonical representation offers an unique mathematical advantage. By proper manipulation of the equations, the differentiation procedure is replaced by integration of the responses. Since integration is inherently a smoothing process, there is an immediate contribution to the stability of the identification results.

General Considerations

A general chemical process* is schematically described by Figure 1.

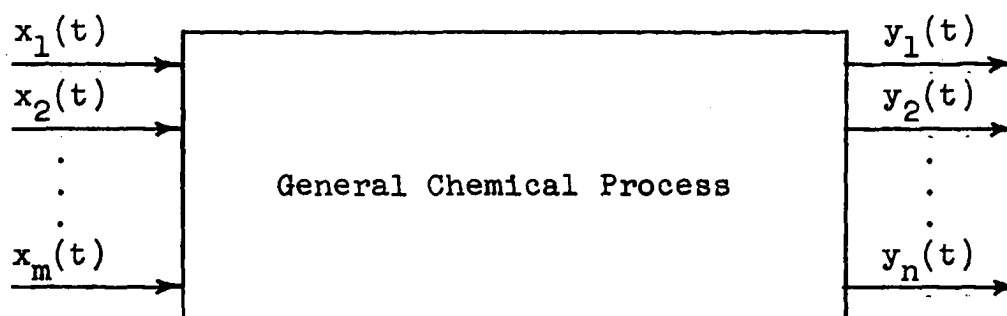


Figure 1. General Chemical Process

*By process, or synonymously, plant or system, in the following discussions, will be meant a mathematical model of a dynamic system.

The variables $y_1(t)$, $i = 1, \dots, n$, are output or dependent variables, and $x_j(t)$, $j = 1, \dots, m$, are the forcing variables or inputs.

The output variables $y_1(t)$ need not necessarily all be physical outputs; some of them can be considered internal functions of the process itself. The input variables are frequently classified, from a control standpoint, as to whether they are controllable or uncontrollable inputs. For example, changes in the ambient temperature may be considered an uncontrollable input to a heat-exchanger whereas variations in coolant flow-rate are generally considered as controllable. From the identification aspect a similar distinction should be made, i.e. forceable and unforceable inputs.

A general process can be linear, linear time varying or nonlinear and is describable by its input-output relationship as follows:

$$\dot{y}(t) = f[x(t), y(t)] \quad (2-1)$$

where $y(t) = y_1(t), \dots, y_n(t)$ is an n vector, $x(t) = x_1(t), \dots, x_m(t)$ is an m vector, and $f[]$ describes the functional relationship between input and output. If the process is linear (or linear time varying), it can be formulated as

$$\dot{\tilde{y}}(t) = A(t)\tilde{y}(t) + B(t)\tilde{x}(t) \quad (2-2)$$

where $\tilde{y}(t)$ and $\tilde{x}(t)$ are column vectors of order n and m respectively, $A(t)$ is a $(n \times n)$ matrix with time dependent elements and $B(t)$ is a $(n \times m)$ matrix with time dependent

elements. (·) stands for differentiation with respect to time. The matrices $A(t)$ and $B(t)$ are independent of $\bar{y}(t)$ and $\bar{x}(t)$.

In expanded form Equation (2-2) can be written as a set of n first order differential equations of the form:

$$\begin{aligned} \frac{dy_1(t)}{dt} &= a_{11}(t)y_1(t) + \dots + a_{1n}(t)y_n(t) + b_{11}(t)x_1(t) \\ &\quad + \dots + b_{1m}(t)x_m(t) \\ &\vdots \\ \frac{dy_n(t)}{dt} &= a_{n1}(t)y_1(t) + \dots + a_{nn}(t)y_n(t) + b_{n1}(t)x_1(t) \\ &\quad + \dots + b_{nm}(t)x_m(t) \end{aligned} \tag{2-3}$$

Relationships of the form (2-3) can normally be obtained from fundamental energy and mass balances on the system.

Equations of the general type such as (2-2) or (2-3) are often obtained as linearized models where the existing nonlinear terms have been discarded or properly truncated [50]. The input-output relationships of a linear system described by Equation (2-2) are completely characterized by the matrices $A(t)$ and $B(t)$. The structure of these matrices describes the topology* of the system.

If the parameters $a_{ij}(t)$ and $b_{ik}(t)$ in (2-3) are continuous time functions and possess the required number of

*In this context topology means the structure of interactions in the plant.

finite continuous derivatives,* then it is possible to express one of the output variables $y_1(t)$ in terms of an n^{th} order linear ordinary differential equation of the form:

$$\frac{d^n y_1(t)}{dt^n} + a_1(t) \frac{dy_1(t)}{dt^{n-1}} + \dots + a_n(t) y_1(t) = \sum_{j=1}^m \left[g_{j1}(t) \frac{d^{n-1} x_j(t)}{dt^{n-1}} + \dots + g_{jn}(t) x_j(t) \right] \quad (2-4)$$

where the coefficients $a_1(t)$ are combinations of $a_{1j}(t)$ and their time derivatives, and the coefficients $g_{j1}(t)$ are combinations of $a_{1j}(t)$ and $b_{1k}(t)$ and their time derivatives. Some of the coefficients $g_{j1}(t)$ may be zero and therefore the order of the right hand side operators may be lower than $n-1$.

To illustrate the transformation of Equation (2-2) into form (2-4) consider the constant parameter, second order system:

$$\begin{vmatrix} \dot{y}_1(t) \\ \dot{y}_2(t) \end{vmatrix} = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} \cdot \begin{vmatrix} y_1(t) \\ y_2(t) \end{vmatrix} + \begin{vmatrix} b_{11} & 0 \\ 0 & b_{22} \end{vmatrix} \cdot \begin{vmatrix} x_1(t) \\ x_2(t) \end{vmatrix}$$

which in terms of $y_1(t)$ transforms into

$$\ddot{y}_1(t) + a_1 \dot{y}_1(t) + a_2 y_1(t) = g_{11} \dot{x}_1 + g_{12} x_1 + g_{22} x_2 \quad .$$

where:

$$\begin{aligned} a_1 &= -(a_{11} + a_{22}) \\ a_2 &= a_{11} a_{22} - a_{12} a_{21} \\ g_{11} &= b_{11} \end{aligned}$$

*A condition usually fulfilled in real processes.

$$g_{12} = -a_{22} b_{11}$$

$$g_{22} = a_{12} b_{22}$$

The usual identification methods evaluate the zeros and poles of systems of the form (2-4). It can be seen from the above simple second order system that the coefficients a_{1j} may become (and usually are) very sensitive to variations in a_1 . Therefore, reconstruction of the matrices A and B in Equation (2-2) cannot be reliably executed if uncertainties in the parameters a_1 are present. Thus, models of the form (2-2) are difficult, if not impossible to obtain from experimental test results.

Equation (2-4) can be written in operator form as

$$L^{(n)}[y_1(t)] = \sum_{j=1}^m M_j^{(k_j)} [x_j(t)] \quad , \quad (2-5)$$

where $L^{(n)}[]$ and $M_j^{(k_j)}[]$ are linear differential operators, with time dependent parameters, of order n and k_j respectively, where $k_j \leq n-1$. If the right hand side of Equation (2-5) is lumped as a general time function $f(t)$, Equation (2-5) transforms to

$$L^{(n)}[y_1(t)] = f(t) \quad . \quad (2-6)$$

For simplicity of notation the subscript 1 will be dropped in the following discussions. If the system described by (2-6) is unforced, ($f(t) = 0$), the homogenous equation can be written:

$$L^{(n)}[y(t)] = \frac{d^n y(t)}{dt^n} + a_1(t) \frac{d^{n-1} y(t)}{dt^{n-1}} + \dots + a_n(t)y(t) = 0 \quad (2-7)$$

There exists a set of n linearly independent solutions* to (2-7), $\varphi_j(t)$, $j = 1, \dots, n$, such that any solution $\Psi(t)$ to (2-7) is given as a linear combination of $\varphi_j(t)$:

$$\Psi(t) = \sum_{j=1}^n C_j \varphi_j(t) \quad (2-8)$$

where C_j are constants. The set $\varphi_j(t)$, $j = 1, \dots, n$, is said to be a fundamental set of solutions to (2-7), or alternatively, a basis to the n dimensional solution space. Given such a set of solutions, the linear system (2-7) is uniquely described [10].

To evaluate the solution (or response) of system (2-6) to any arbitrary time function $f(t)$, the unit impulse response will be introduced. The unit impulse function $\delta(t)$ is mathematically defined as

$$\begin{aligned} \delta(t) &= 0 & \text{for } t \neq 0 \\ \int_{-\infty}^{\infty} \delta(t) dt &= 1 \end{aligned} \quad (2-9)$$

Consider the solution to the system

$$L^{(n)}[y(t)] = f(t) = \delta(t - \lambda), \quad (2-10)$$

where λ is the time at which the impulse function is applied to the system. If the system is a general (time varying) linear system, then the response will be a function of both the observation time, t , and the time at which the impulse

* n solutions $\varphi_j(t)$ are linearly independent if there exists no nontrivial set of constants $C_1, \dots, C_n$ such that

$$C_1 \varphi_1(t) + C_2 \varphi_2(t) + \dots + C_n \varphi_n(t) = 0$$

is applied, λ . The response to the unit impulse function is denoted by $W(t, \lambda)$ and is frequently called the weighting function. Since $\delta(t - \lambda)$ is zero for all $t \neq \lambda$, the impulse response is the homogeneous (zero-forcing) response of the system for $t > \lambda$, satisfying a particular set of initial conditions at $t = \lambda$. Thus, the weighting function is a particular solution $\Psi(t)$ to the system (2-7). It will later be shown how it can be obtained from the fundamental set of solutions $\phi_j(t)$. For $t < \lambda$, $W(t, \lambda)$ is of course identically equal to zero and so are all its derivatives. Thus:

$$L^{(n)}[W(t, \lambda)] = 0 \quad \text{for } t \neq \lambda \quad (2-11)$$

If Equation (2-10) is considered,

$$L^{(n)}[W(t, \lambda)] = \delta(t - \lambda) . \quad (2-12)$$

Integrating Equation (2-12) and applying the second condition of Equation (2-9):

$$\int_{-\infty}^{\infty} L^{(n)}[W(t, \lambda)] dt = 1 \quad (2-13)$$

Substitution of Equation (2-11) into (2-13)

$$\lim_{\epsilon \rightarrow 0} \int_{\lambda - \epsilon}^{\lambda + \epsilon} L^{(n)}[W(t, \lambda)] dt = 1 \quad (2-14)$$

or since $W(t, \lambda) = 0$ for $t < \lambda$:

$$\left\{ \frac{d^{n-1}}{dt^{n-1}} [W(t, \lambda)] + a_1(t) \frac{d^{n-2}}{dt^{n-2}} [W(t, \lambda)] + \dots + \right. \\ \left. a_{n-1}(t) W(t, \lambda) \right\}_{t=\lambda} + a_n(t) \lim_{\epsilon \rightarrow 0} \int_{\lambda-\epsilon}^{\lambda+\epsilon} W(t, \lambda) dt = 1 \quad (2-15)$$

Since $W(t, \lambda) = 0$ for $t < \lambda$ it follows that $W(t, \lambda)$ and its first $n-2$ derivatives must be zero at $t = \lambda$.^{*} Thus, the weighting function possesses the following properties:

$$\begin{aligned} (a) \quad \frac{\partial^\alpha}{\partial t^\alpha} W(t, \lambda) \Big|_{t=\lambda} &= 0 \quad \text{for } \alpha = 0, 1, \dots, n-2, n \\ (b) \quad \frac{\partial^{n-1}}{\partial t^{n-1}} W(t, \lambda) \Big|_{t=\lambda} &= 1 \\ (c) \quad L^{(n)} [W(t, \lambda)] &= 0 \quad \text{for } t \neq \lambda \\ (d) \quad W(t, \lambda) &\text{ is unique!} \end{aligned} \quad (2-16)$$

The response of system (2-6) to an arbitrary forcing function $f(t)$ can be expressed in terms of the system weighting function by the following equation:

$$y(t) = \int_{-\infty}^t W(t, \lambda) f(\lambda) d\lambda \quad (2-17)$$

The upper limit of integration indicates that only those values of $f(t)$, for which $\lambda < t$, are significant in determining $y(t)$. For a physical realizable system, $W(t, \lambda) = 0$, for $t < \lambda$, since otherwise the system would be anticipatory, i.e. responsive to future values of the forcing functions.

^{*}If this condition were not fulfilled it would mean that either $W(t, \lambda)$ or at least one of its first $n-2$ derivatives experiences a finite discontinuity at $t = \lambda$ in which case all the higher derivatives (including the $(n-1)^{\text{st}}$) would be undefined and (2-15) would not be satisfied.

That Equation (2-17) holds can be shown by direct substitution of $y(t)$ from (2-17) into (2-6). Thus, taking the first derivative of (2-17)

$$\frac{dy(t)}{dt} = \int_{-\infty}^t \frac{\partial W(t, \lambda)}{\partial t} f(\lambda) d\lambda + W(t, \lambda) f(\lambda) \Big|_{t=\lambda} . \quad (2-18)$$

The last term vanishes since by (2-16) $W(t, \lambda) = 0$ for $t = \lambda$.

Similarly:

$$\frac{\partial^k y(t)}{\partial t^k} = \int_{-\infty}^t \frac{\partial^k W(t, \lambda)}{\partial t^k} f(\lambda) d\lambda \quad \text{for } k = 0, 1, \dots, n-1 \quad (2-19)$$

For the n^{th} derivative:

$$\frac{d^n y(t)}{dt^n} = \int_{-\infty}^t \frac{\partial^n W(t, \lambda)}{\partial t^n} f(\lambda) d\lambda + \frac{\partial^{n-1} W(t, \lambda)}{\partial t^{n-1}} f(\lambda) \Big|_{t=\lambda} \quad (2-20)$$

and by utilizing condition (b) of (2-16) one obtains:

$$\frac{d^n y(t)}{dt^n} = \int_{-\infty}^t \frac{\partial^n W(t, \lambda)}{\partial t^n} f(\lambda) d\lambda + f(t) \quad (2-21)$$

Substituting the above expression for the derivatives into the expanded form of (2-6):

$$\int_{-\infty}^t \left\{ \frac{\partial^n W(t, \lambda)}{\partial t^n} + a_1(t) \frac{\partial^{n-1} W(t, \lambda)}{\partial t^{n-1}} + \dots + a_n(t) W(t, \lambda) \right\} f(\lambda) d\lambda + f(t) = f(t) \quad (2-22)$$

from which it follows that:

$$\int_{-\infty}^t L^{(n)}[W(t, \lambda)] f(\lambda) d\lambda = 0 \quad (2-23)$$

which is immediately verified by condition (c), Equation (2-16).

If $y(t) = 0$ prior to $t = \tau$ then

$$y(t) = \int_{\tau}^t W(t, \lambda) f(\lambda) d\lambda \quad . \quad (2-24)$$

If a transient is present at $t = \tau$ and $f(t) = 0$, for $t < \tau$, then

$$y(t) = \sum_{j=1}^n C_j \varphi_j(t) + \int_{\tau}^t W(t, \lambda) f(\lambda) d\lambda \quad , \quad (2-25)$$

where $\varphi_j(t)$, $j = 1, \dots, n$, are solutions to (2-7) and C_j , $j = 1, \dots, n$, are constants determined by the initial conditions (i.e. the value of $y(t)$ and its derivatives at time $t = \tau$) such that

$$y(\tau) = \sum_{j=1}^n C_j \varphi_j(\tau) \quad . \quad (2-26)$$

The utility of the weighting function as a means for system characterization is now evident: It can be conveniently used for evaluation of system response to an arbitrary forcing function, and it is of particular convenience for time invariant linear systems, as will be discussed later. It can be shown [57] that the weighting function is given by the relation

$$W(t, \lambda) = U(t - \lambda) \sum_{j=1}^n \varphi_j(t) \varphi_j^*(\lambda) \quad (2-27)$$

where $\varphi_j(t)$ are any n linearly independent solutions to the homogeneous Equation (2-7), and $\varphi_j^*(t)$ are n linearly independent solutions to the adjoint Equation of (2-7)

$$\begin{aligned} L^{(n)*}[y(t)] = & (-1) \frac{d^n y(t)}{dt^n} + (-1) \frac{d^{n-1}}{dt^{n-1}} [a_1(t)y(t)] + \dots \\ & + a_n(t)y(t) = 0 \quad . \end{aligned} \quad (2-28)$$

In Equation (2-27) $U(t - \lambda)$ is the unit step function applied at $t = \lambda$ where it is understood that:

$$U(t - \lambda) = \begin{cases} 0 & \text{for } t < \lambda \\ 1 & \text{for } t \geq \lambda \end{cases}$$

The weighting function can also be expressed in terms of the $\varphi_j(t)$ alone [10] as follows:

$$W(t, \lambda) = U(t - \lambda) \sum_{j=1}^n \frac{W_j(t) W_{rj}(\lambda)}{W_r(\varphi_1, \dots, \varphi_n)(\lambda)} \quad (2-29)$$

where $W_r(\varphi_1, \dots, \varphi_n)(\lambda)$ is the Wronskian of $(\varphi_1, \dots, \varphi_n)$ evaluated at $t = \lambda$, i.e.

$$W_r(\varphi_1, \dots, \varphi_n)(\lambda) = \det. \begin{vmatrix} \varphi_1(t) & \dots & \varphi_n(t) \\ \dot{\varphi}_1(t) & \dots & \dot{\varphi}_n(t) \\ \vdots & \dots & \vdots \\ \varphi_1^{(n-1)}(t) & \dots & \varphi_n^{(n-1)}(t) \end{vmatrix}_{t=\lambda} \quad (2-30)$$

$W_{rj}(\lambda)$ is the determinant obtained from $W_r(\varphi_1, \dots, \varphi_n)(\lambda)$ by replacing the j^{th} column $(\varphi_j, \dot{\varphi}_j, \dots, \varphi_j^{(n-1)})$ by $(0, \dots, 0, 1)$.

An equivalent but slightly different form for $W(t, \lambda)$ has been proposed by Miller [36]

$$W(t, \lambda) = \frac{(-1)^{n-1} U(t - \lambda)}{W_r(\varphi_1, \dots, \varphi_n)(\lambda)} \det. \begin{vmatrix} \varphi_1(t) & \dots & \varphi_n(t) \\ \varphi_1(\lambda) & \dots & \varphi_n(\lambda) \\ \dot{\varphi}_1(\lambda) & \dots & \dot{\varphi}_n(\lambda) \\ \vdots & \dots & \vdots \\ \varphi^{(n-2)}(\lambda) & \dots & \varphi^{(n-2)}(\lambda) \end{vmatrix} \quad (2-31)$$

If a fundamental set of solutions to a linear system of the

form (2-7) is given, then the weighting function can be evaluated by means of Equation (2-29) or (2-31). The solution of system (2-6) to any forcing function $f(t)$ can thus be computed by means of Equation (2-17) or (2-25).

Furthermore, it can be shown [10,57] that given the set of solutions $\phi_j(t)$, the homogeneous operator of Equation (2-7) can be constructed as follows:

$$L^{(n)}[y(t)] = \frac{d^n y(t)}{dt^n} + \dots + a_n(t)y(t)$$

$$= \frac{\det \begin{vmatrix} \phi_1(t) & \dots & \phi_n(t) & y(t) \\ \dot{\phi}_1(t) & \dots & \dot{\phi}_n(t) & \dot{y}(t) \\ \vdots & \dots & \vdots & \vdots \\ \phi_1^{(n)}(t) & \dots & \phi_n^{(n)}(t) & y^{(n)}(t) \end{vmatrix}}{W_r(\phi_1, \dots, \phi_n)(t)} \quad (2-32)$$

At this stage reconsider Equation (2-5) and its equivalent form Equation (2-6)

$$L^{(n)}[y(t)] = \sum_{j=1}^m M_j^{(k_j)}[x_j(t)] = f(t) \quad (2-5)$$

Equation (2-25) can be rewritten as

$$y(t) = \sum_{j=1}^n c_j \phi_j(t) + \int_{\tau}^t W(t, \lambda) \sum_{j=1}^m M_j^{(k_j)}[x_j(\lambda)] d\lambda$$

$$= \sum_{j=1}^n c_j \phi_j(t) + \sum_{j=1}^m \int_{\tau}^t W(t, \lambda) M_j^{(k_j)}[x_j(\lambda)] d\lambda \quad (2-33)$$

Equation (2-33) expresses the response of the system to any input function $x_j(t)$. Equation (2-33) can be expressed in an equivalent, but more convenient, form as

$$y(t) = \sum_{j=1}^n c_j \varphi_j(t) + \sum_{j=1}^m \int_{\tau}^t W_j(t, \lambda) x_j(\lambda) d\lambda \quad (2-34)$$

where $W_j(t, \lambda)$ in Equation (2-34) is related to $W(t, \lambda)$ in (2-33) by the expression

$$W_j(t, \lambda) = (-1)^{k_j} \frac{d^{k_j}}{dt^{k_j}} [g_j(t) W(t, \lambda)] + \dots + g_j^{(k_j)}(t) W(t, \lambda) \quad (2-35)$$

In operator form (2-35) can be written as:

$$W_j(t, \lambda) = M_j^{(k_j)*} [W(t, \lambda)] \quad (2-35a)$$

where $M_j^{(k_j)*} []$ is the adjoint operator of $M_j^{(k_j)} []$. Equations (2-34) and (2-35) can be verified directly through integration by parts of Equation (2-33) and comparison to (2-34).

Consider again the linear system, Equation (2-6):

$$L^{(n)}[y(t)] = \frac{d^n y(t)}{dt^n} + a_1(t) \frac{d^{n-1} y(t)}{dt^{n-1}} + \dots + a_n(t) y(t) = f(t) \quad (2-6)$$

Associated with Equation (2-6) there is a system of equations

$$\begin{aligned} y(t) &= z_1(t) \\ \frac{dz_1(t)}{dt} &= z_2(t) \\ &\vdots \\ \frac{dz_{n-1}(t)}{dt} &= z_n(t) \\ \frac{dz_n(t)}{dt} &= -a_n(t) z_1(t) - a_{n-1}(t) z_2(t) - \dots - a_1(t) z_n(t) \\ &\quad + \bar{f}_n(t) \end{aligned} \quad (2-36)$$

where $\bar{f}_n(t) = f(t)$. This relationship can be expressed in vector-matrix form as

$$\dot{\mathbf{z}}(t) = \frac{d\mathbf{z}}{dt} = \bar{\mathbf{A}}(t)\mathbf{z}(t) + \bar{\mathbf{F}}(t) \quad (2-37)$$

where $\mathbf{z}(t)$ is a column vector, $\mathbf{z}(t) = \text{col}(z_1(t), \dots, z_n(t))$, $\bar{\mathbf{F}}(t)$ is a column vector, $\bar{\mathbf{F}}(t) = \text{col}(0, \dots, 0, \bar{f}_n(t))$ and $\bar{\mathbf{A}}(t)$ is the $n \times n$ matrix of coefficients:

$$\bar{\mathbf{A}}(t) = \begin{vmatrix} 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ -a_n(t) & -a_{n-1}(t) & -a_{n-2}(t) & \dots & -a_1(t) \end{vmatrix} \quad (2-38)$$

That the system represented by Equation (2-36) is equivalent to Equation (2-6) with regard to the relationship of $y(t)$ to $z_1(t)$ can be easily verified by direct substitution. Equation (2-36) is sometimes called the state-space equation of the system represented by Equation (2-6) [58]. The equivalence relation between Equations (2-36) and (2-6) warrants some discussion. In Equation (2-36) the variables $z_i(t)$ (for $i = 2, \dots, n$) are artificial and have no physical significance in the sense of input output variables of the original system. The only variable which has physical meaning is $z_1(t)$. However, both systems have the same fundamental solution space and moreover, both $y(t)$ and $z_1(t)$ are related to $f(t)$ or $\bar{f}_n(t)$ by the same mathematical relationship. Mathematical manipulation of the system of Equations (2-36) does not alter this relationship as long as the physical (or mathematical) meaning

of $z_1(t)$ and $\bar{F}_n(t)$ is preserved. What has been done, in effect, is the formulation of a hypothetical structure for a scalar system.

Of more importance here is the relationship between the original system, Equation (2-3) and its "equivalent" form Equation (2-36). Equation (2-36) is obtainable from (2-3) by a linear transformation, and from a mathematical viewpoint the two are similar systems. However the topological structure has been completely altered and all the output variables except $y_1(t)$ have been "lost" in the process of transformation. The individual output variables $y_1(t)$ and $z_1(t)$ do, however, still bear the same relationship to the forcing, or input, variables. Thus, if only the dynamics of $y_1(t)$ is of interest the topological transformation is entirely immaterial.

The homogeneous part of (2-37) is:

$$\dot{\bar{z}}(t) = \bar{A}(t)\bar{z}(t) \quad (2-39)$$

The fundamental solution matrix (i.e. the matrix obtained by adjoining n linearly independent solution vectors $\Psi_j(t) = \text{col}(\varphi_{j1}, \dots, \varphi_{jn})$ of Equation (2-39) is given by

$$\Phi(t) = \begin{vmatrix} \varphi_1(t) & \dots & \varphi_n(t) \\ \dot{\varphi}_1(t) & \dots & \dot{\varphi}_n(t) \\ \vdots & \dots & \vdots \\ \varphi_1^{(n-1)}(t) & \dots & \varphi_n^{(n-1)}(t) \end{vmatrix}, \quad (2-40)$$

where $\varphi_j(t)$, $j = 1, \dots, n$, are n linearly independent solutions of Equation (2-7). If $\Phi(t)u(t) = \underline{y}(t)$ is a solution

to Equation (2-37), where $u(t)$ is a matrix function yet undetermined, then

$$\dot{\underline{y}}(t) = \dot{\Phi}(t)u(t) + \Phi(t)\dot{u}(t) \quad (2-41)$$

Substitution of Equation (2-41) into (2-37) yields

$$\dot{\underline{y}}(t) = \dot{\Phi}(t)u(t) + \Phi(t)\dot{u}(t) = \bar{A}(t)\Phi(t)u(t) + \bar{F}(t) \quad (2-42)$$

and since $\dot{\Phi}(t) = \bar{A}(t)\Phi(t)$ (as $\Phi(t)$ satisfies Equation (2-39))

$$\dot{\underline{y}}(t) = \bar{A}(t)\Phi(t)u(t) + \Phi(t)\dot{u}(t) = \bar{A}(t)\Phi(t)u(t) + \bar{F}(t) \quad (2-43)$$

or
$$\Phi(t)\dot{u}(t) = \bar{F}(t) \quad (2-44)$$

Since $\Phi(t)$ is a nonsingular matrix [10]

$$u(t) = \int_{\tau}^t \Phi^{-1}(\lambda)\bar{F}(\lambda)d\lambda + C \quad (2-45)$$

and

$$\underline{y}(t) = \Phi(t) \cdot C + \int_{\tau}^t \Phi(t)\Phi^{-1}(\lambda)\bar{F}(\lambda)d\lambda \quad (2-46)$$

The constant matrix C is determined by the initial conditions such that:

$$\underline{y}(\tau) = \Phi(\tau)C$$

Comparing Equation (2-46) with (2-25) it can be seen that:

$$\bar{W}(t, \lambda) = \Phi(t)\Phi^{-1}(\lambda) \quad (2-47)$$

where $\bar{W}(t, \lambda)$ is the matrix weighting function for Equation (2-37). $\bar{W}(t, \lambda)$ is related to $W(t, \lambda)$ of Equation (2-25).

Since $z(t)$ in Equation (2-37) equals $y(t)$ in Equation (2-6) and since $\bar{F}(t) = \text{col}(0, \dots, 0, \bar{F}_n(t))$, where $\bar{F}_n(t) = f(t)$, it follows that the weighting function $W(t, \lambda)$ is given as:

$$W(t, \lambda) = [\bar{W}(t, \lambda)]_{1n} = [\bar{G}(t) \bar{G}^{-1}(\lambda)]_{1n} \quad (2-48)$$

where the subscript (1n) in the above equation stands for the element in the first row and n^{th} column. Equation (2-48) can be verified by expansion of Equation (2-46) subject to the given constraints.

Systems with Constant Parameters

In the preceding sections the properties of a general class of linear (time varying) systems have been discussed. Although of considerable importance in adaptive and programmed control schemes, general (nonperiodic) time varying systems are of little, if any, significance in the process industries. Time variations in parameters are normally not present (except for high frequency noise), and if such exist the variation is slow and insignificant from a control aspect. Examples of such slow variations are catalyst deactivation in catalytic reactors, changes of ambient temperature in heat transfer equipment, etc. Periodic controller resetting usually takes care of such variations. When fast time variations in parameters are expected, then cascaded or adaptive control schemes are usually applied.

Identification of nonperiodic, time varying systems is virtually impossible since the identified model is valid only for the observation period. Thus, no useful information about the system can be obtained.

Systems with periodically varying parameters are

uncommon in the chemical industry, but interest in such systems is growing. Investigation of the dynamics of a pulsed arc reactor is currently underway at Massachusetts Institute of Technology, and studies of the performance of pulsed processes (such as heat exchangers, distillation columns, etc.) have been reported with quite impressive results. Mathematical relationships for identification of linear systems with periodic time varying parameters are presented in Appendix B.

Although some nonlinear effects are present in most chemical processes, a constant parameter, linear models often adequately describe the overall dynamic behavior. Most nonlinear effects in chemical processes can be described by linear models with time delays and are accurate within the uncertainty limits of experimental test data. Since generalized techniques for identifying nonlinear models are unavailable, it is frequently desirable to find an approximate "linearized model" which, in turn, is used as a basis for improvement and proper introduction of nonlinear terms. The forthcoming discussions will center on identification of linear models with constant parameters.

Invariant linear systems can be considered as a special subcase of the general time varying, linear systems; thus, the preceding mathematical relationships are valid for the latter. Significant simplifications in the theory are, however, possible and these will be discussed in the present section.

If the linear system is invariant then Equation (2-6)

can be written as

$$L^{(n)}[y(t)] = \frac{d^n y(t)}{dt^n} + a_1 \frac{d^{n-1}}{dt^{n-1}} y(t) + \dots + a_n y(t) = f(t) \quad (2-49)$$

where $a_1, 1 = 1, \dots, n$, are constants. Associated with Equation (2-49) is a system similar to (2-37):

$$\dot{\bar{z}}(t) = \bar{A}\bar{z}(t) + \bar{F}(t) \quad (2-50)$$

where $\bar{A}$ is the following constant matrix:

$$\bar{A} = \begin{vmatrix} 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & 1 \\ -a_n & -a_{n-1} & -a_{n-2} & \dots & -a_1 \end{vmatrix} \quad (2-51)$$

and $\bar{F}(t)$ and $\bar{z}(t)$ have the same meaning as in Equation (2-37).

The homogeneous fundamental solution matrix $\Phi(t)$ of Equation (2-50) is the same as (2-40) and it is a simple matter to show that:

$$\Phi(t) = e^{t\bar{A}}C \quad (2-52)$$

where $\bar{A}$ is the matrix of Equation (2-51) and C a constant, nonsingular matrix determined by the initial conditions ($\Phi(0) = C$). From Equation (2-47) it follows that the weighting function for Equation (2-50) is:

$$\begin{aligned} \bar{W}(t, \lambda) &= \Phi(t)\Phi^{-1}(\lambda) = [e^{t\bar{A}}C][e^{\lambda\bar{A}}C]^{-1} = e^{t\bar{A}} \cdot C^{-1}e^{-\lambda\bar{A}} \\ &= e^{(t-\lambda)\bar{A}} = \bar{W}(t - \lambda) \end{aligned} \quad (2-53)$$

Thus, the weighting function for a constant parameter, linear

system is a function of the elapsed time between the time the input was applied λ and the time of observation t . Equation (2-17) then becomes

$$y(t) = \int_{-\infty}^t W(t - \lambda) f(\lambda) d\lambda = \int_{-\infty}^t e^{(t-\lambda)A} f(\lambda) d\lambda \quad (2-54)$$

Let Q be a nonsingular constant matrix such that

$$QAQ^{-1} = J \quad (2-55)$$

where J is the Jordan canonical form of A

$$J = \begin{vmatrix} J_0 & 0 & 0 & \dots & 0 \\ 0 & J_1 & 0 & \dots & 0 \\ \vdots & \vdots & \vdots & \dots & \vdots \\ \vdots & \vdots & \vdots & \dots & \vdots \\ 0 & 0 & 0 & \dots & J_s \end{vmatrix} \quad (2-56)$$

and where J_0 is a diagonal matrix: $J_0 = \text{diag}(\rho_1, \dots, \rho_q)$

and the J_i are $(r_i \times r_i)$ matrices of the form

$$J_i = \begin{vmatrix} \rho_{q+1} & 1 & 0 & \dots & 0 & 0 \\ 0 & \rho_{q+1} & 1 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \dots & \vdots & \vdots \\ \vdots & \vdots & \vdots & \dots & \vdots & \vdots \\ \vdots & \vdots & \vdots & \dots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & \rho_{q+1} & 1 \\ 0 & 0 & 0 & \dots & 0 & \rho_{q+1} \end{vmatrix} \quad i = 1, \dots, s \quad (2-57)$$

The elements $\rho_1, \dots, \rho_q$ are the distinct roots of the characteristic equation of $\bar{A}$ ($\det.(I\theta - \bar{A}) = 0$) and ρ_{q+1} are the characteristic roots of $\bar{A}$ of multiplicity r_i , ($r_i > 1$).

It follows that

$$e^{tJ} = \begin{vmatrix} e^{tJ_0} & 0 & \dots & 0 \\ 0 & e^{tJ_1} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & e^{tJ_s} \end{vmatrix} \quad (2-58)$$

where

$$e^{tJ_0} = \begin{vmatrix} e^{t\rho_1} & 0 & \dots & 0 \\ 0 & e^{t\rho_2} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & e^{t\rho_q} \end{vmatrix} \quad (2-59)$$

and

$$e^{tJ_1} = e^{t\rho_{q+1}} \begin{vmatrix} 1 & t & \dots & \frac{t^{r_1-1}}{(r_1-1)!} \\ 0 & 1 & \dots & \frac{t^{r_1-2}}{(r_1-2)!} \\ \vdots & \vdots & \ddots & \vdots \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{vmatrix} \quad 1 > 0 \quad (2-60)$$

It can be shown [40] that

$$e^{(t-\lambda)QAQ^{-1}} = e^{(t-\lambda)J} = Qe^{(t-\lambda)\bar{A}}Q^{-1}, \quad (2-61)$$

from which it follows that

$$\bar{W}(t, \lambda) = \Phi(t)\Phi^{-1}(\lambda) = e^{(t-\lambda)A} = Q^{-1}e^{(t-\lambda)J}Q, \quad (2-62)$$

or

$$\Phi(t) = [Q^{-1}e^{(t-\lambda)J}Q]\Phi(\lambda) = e^{(t-\lambda)A}\Phi(\lambda). \quad (2-63)$$

Equation (2-63) will later be used as the basis for the present identification technique.

Derivation of the Identification Technique

In the preceding sections general mathematical relationships have been discussed, which constitute the foundation for the proposed identification technique. Although the technique has not yet been explicitly outlined, the general philosophy has been implied. The analysis is carried out entirely in the time domain without resorting to frequency domain analysis. The system weighting function in its analytical form or an equivalent formulation is the ultimate goal of the identification.

In the present section the mathematical procedures of the identification are presented, and in the following section a preliminary evaluation of the expected advantages and limitations is carried out, including a discussion of some of the alternative routes that can be followed.

Consider the general linear (time varying) process described in Figure 2:

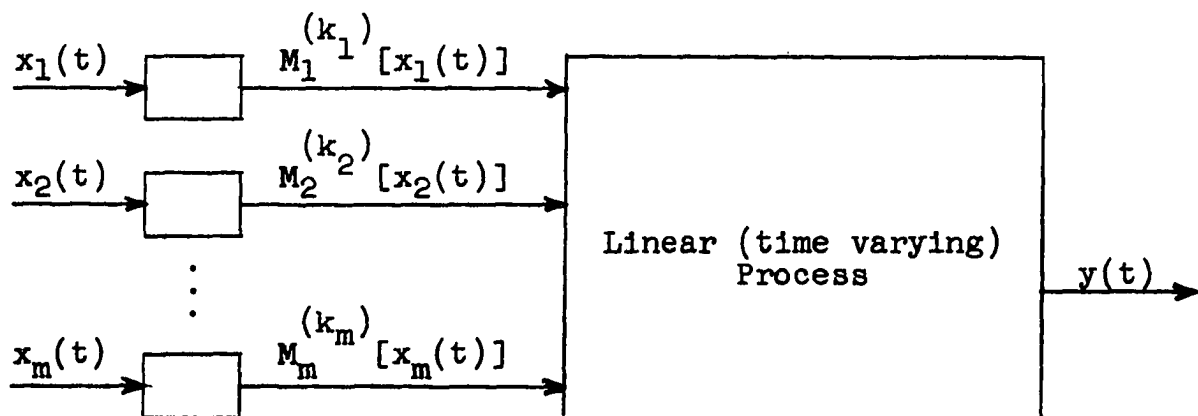


Figure 2. Scalar Linear Process

where $x_1(t), \dots, x_k(t)$ are process inputs and $y(t)$ is the process output, all of which are assumed to be measurable. The process can be described by the scalar n^{th} order linear differential equation

$$L^{(n)}[y(t)] = \sum_{j=1}^m M_j^{(k_j)}[x_j(t)] \quad , \quad (2-5)$$

where $y(t)$ is the output variable of interest (out of the n existing output variables $y_1(t)$), with respect to which the system dynamics is to be evaluated. No one of the process parameters is assumed to be known. The order n of the system is considered to be known although this knowledge is not an essential requirement.*

The system of Equation (2-5) can also be represented in the equivalent form (2-6)

$$L^{(n)}[y(t)] = f(t) \quad (2-6)$$

where $f(t)$, as before, represents a general arbitrary forcing function. Schematically Equation (2-6) is represented in Figure 3:

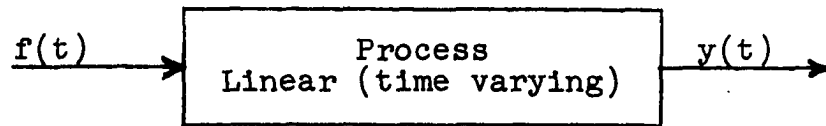


Figure 3. Scalar Linear Process with General Input Function

Although Equations (2-5) and (2-6) are mathematically equivalent, they differ from each other in concept. Equation (2-6) describes the system response to an "overall" disturbance

*Frequently the formal system equations are given from theoretical considerations. If no prior knowledge of the system topology is available, a model of "desired" order can be assumed which is later verified or invalidated.

(or forcing function), but it does not reveal anything about the coupling or linkage between the "real" input variables and the particular output variable under consideration. It can, therefore, be considered only as a partial model and will be hereafter called The General Model of the System. Equation (2-5) describes the total dynamic behaviour of the system giving the exact relationship between $y(t)$ and any of the inputs $x_j(t)$. It will be hereafter called The Particular Model of the System. The weighting function $W(t,\lambda)$ associated with Equation (2-6) will be called the General Weighting Function of the system, and the weighting function associated with (2-5) $W_j(t,\lambda)$, which relates a particular input variable $x_j(t)$ to the output $y(t)$, will be called the Particular Weighting Function $W_j(t,\lambda)$.

Although the general model is only a partial description of the system it may often be of interest by itself, particularly when stability and general response characteristics are to be investigated.

Determination of the General Model

It has been indicated in previous sections (Equations (2-29) and (2-31)) that the general weighting function of a linear system of order n can be evaluated numerically from a set of n linearly independent solutions of the homogeneous equation. In practical terms this means that n linearly independent response curves* of the unforced plant are needed.

* n response curves are linearly independent if there

Such a set can be obtained experimentally by forcing the system n times to linearly independent initial conditions* at which time the input is removed and the system relaxes autonomously to its steady state.

Given a record of the n linearly independent responses, the Wronskian can be obtained by numerical differentiation of the response curves. The general weighting function $W(t, \lambda)$ can thus be constructed by means of Equation (2-29) or (2-31). Also, the homogeneous operator $L^{(n)}[y]$ can be determined by Equation (2-32).

A slightly different approach for determining the general weighting function is possible by making use of the state-space formulation and the linear system described by Equation (2-37). A fundamental solution matrix $\Phi(t)$ can be constructed, Equation (2-40), and the weighting function evaluated by means of Equations (2-47) and (2-48).

Determination of the Particular Model

Associated with each of the input variables $x_e(t)$ there is a particular weighting function $W_e(t, \lambda)$ which relates $x_e(t)$ with its corresponding output response $y(t)$. Equation

exist no constants $c_1, \dots, c_n$ such that:

$$c_1[y_1(t), \dot{y}_1(t), \dots, y_1^{(n-1)}(t)] + c_2[y_2(t), \dot{y}_2(t), \dots, y_2^{(n-1)}(t)] + \dots + c_n[y_n(t), \dot{y}_n(t), \dots, y_n^{(n-1)}(t)] = 0$$

*By the existence and uniqueness theorems of differential equations [10] if the initial conditions are linearly independent, the solutions are independent over the whole time domain.

(2-34) relates the response $y(t)$ to all the input variables $x_j(t)$ as follows:

$$y(t) = \sum_{j=1}^n C_j \phi_j(t) + \sum_{j=1}^m \int_{\tau}^t W_j(t, \lambda) x_j(\lambda) d\lambda \quad (2-34)$$

If all the $x_j(t)$'s except $x_e(t)$ are kept zero then Equation (2-34) can be written as:

$$y(t) = \sum_{j=1}^n C_j \phi_j(t) + \int_{\tau}^t W_e(t, \lambda) x_e(\lambda) d\lambda \quad (2-64)$$

where $W_e(t, \lambda)$ is the particular weighting function to be determined.

The term $\sum_{j=1}^n C_j \phi_j(t)$ in (2-64) is the transient response which depends on the initial conditions, i.e. the state of the system at $t = \tau$. If the initial conditions of the system are zero prior to application of the input $x_e(t)$ then Equation (2-64) becomes:

$$y(t) = \int_{\tau}^t W_e(t, \lambda) x_e(\lambda) d\lambda. \quad (2-65)$$

$W_e(t, \lambda)$ has been previously shown to be given by Equation (2-35) as:

$$W_e(t, \lambda) = (-1)^{k_e} \frac{d^{k_e}}{dt^{k_e}} [g_{e0}(t)W(t, \lambda)] + \dots + g_{ek_e}(t)W(t, \lambda) \quad (2-35)$$

where $W(t, \lambda)$ is the general weighting function of the system, and the differential operator on the right hand side is the adjoint of the operator $M_e^{(k_e)} []$ in Equation (2-5). $W(t, \lambda)$ is determined by the methods described in the previous section;

thus, in order to determine $W_e(t, \lambda)$ the operator $M_e^{(k_e)} []$ remains to be evaluated.

The relation between $x_e(t)$ and $y(t)$ can also be expressed in a form equivalent to Equation (2-23) as

$$y(t) = \int_{\tau}^t W(t, \lambda) M_e^{(k_e)} [x_e(\lambda)] d\lambda \quad (2-66)$$

or alternatively

$$y(t) = \int_{\tau}^t W(t, \lambda) g_{e0}(t) \frac{d^{k_e}}{dt^{k_e}} [x_e(\lambda)] d\lambda + \dots \\ + \int_{\tau}^t W(t, \lambda) g_{ek_e}(t) x_e(\lambda) d\lambda \quad (2-67)$$

The coefficients $g_{ej}(t)$ can be determined numerically by the following procedure: The system is forced k_e times with linearly independent forcing functions which are recorded together with their corresponding responses. By numerical solution of k_e relations of the type represented by Equation (2-67) the coefficients $g_{ej}(t)$ are evaluated.

If the process is stationary (constant parameter) the evaluation of the operator $M_e^{(k_e)} []$ is considerably simplified. Equation (2-67) can then be written as

$$y(t) = g_{e0} \int_{\tau}^t W(t, \lambda) \frac{d^{k_e}}{dt^{k_e}} [x_e(\lambda)] d\lambda + \dots + g_{ek_e} \int_{\tau}^t W(t, \lambda) x(\lambda) d\lambda \quad (2-68)$$

The testing procedure described above is applied in a similar manner, but the equations can be solved explicitly for the parameters. Utilization of special forcing functions such

as a step function, a ramp and sinusoidal input further simplifies the evaluation.

Success in the evaluation of the particular model hinges almost entirely on the accuracy of the general weighting function $W(t, \lambda)$. Any error in $W(t, \lambda)$ will propagate and possibly even be magnified in the process of evaluating the particular weighting function. On the other hand, if a satisfactory general model can be obtained, it would seem that no particular difficulties should be encountered with evaluation of the particular model. All further discussion will be limited to identification of the General Model.

Identification of the General Model for

Constant Parameter Systems

The techniques described previously for identification of the general model of linear time varying systems also hold, of course, for constant parameter systems. They are, however, subject to rather serious limitations because of the noise problems involved and particularly because of the sensitivity of the operation of differentiation of numerical data. Furthermore, the weighting function is obtained numerically, i.e. a plot of $W(t, \lambda)$ versus time, which in general is not useful in application. To obtain its analytical (functional) description is not a trivial operation, particularly when $W(t, \lambda)$ is corrupted by noise effects. Noise in the weighting function plot will possibly completely obscure the results. A different technique for identification of

constant parameter systems, which is not subject to the above limitations, is based on the specialized equations for constant parameter systems presented previously.

Consider the homogeneous system associated with Equation (2-49)

$$L^{(n)}[y(t)] = \frac{d^n y(t)}{dt^n} + a_1 \frac{d^{n-1} y(t)}{dt^{n-1}} + \dots + a_n y(t) = 0 \quad (2-69)$$

The general weighting function for the associated system, Equation (2-50) is given by Equation (2-53) as

$$\bar{W}(t - \lambda) = \Phi(t) \Phi^{-1}(\lambda) = e^{(t-\lambda)\bar{A}} \quad (2-53)$$

such that $[\bar{W}(t-\lambda)]_{1n}$ is the weighting function in question.

$\Phi(t)$ is the fundamental solution matrix as given by Equation (2-40)

$$\Phi(t) = \begin{vmatrix} \phi_1(t) & \dots & \phi_n(t) \\ \dot{\phi}_1(t) & \dots & \dot{\phi}_n(t) \\ \vdots & \dots & \vdots \\ \phi_1^{(n-1)}(t) & \dots & \phi_n^{(n-1)}(t) \end{vmatrix} \quad (2-40)$$

where $\phi_j(t)$ are the linearly independent, unforced (homogeneous) responses of the system. $\bar{A}$ is the equivalent coefficient matrix given by Equation (2-51)

$$\bar{A} = \begin{vmatrix} 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \dots & \vdots \\ \vdots & \vdots & \vdots & \dots & \vdots \\ -a_n & -a_{n-1} & -a_{n-2} & \dots & -a_1 \end{vmatrix} \quad (2-51)$$

Equation (2-53) could be used for direct determination of $\bar{W}(t - \lambda)$ from the response curves. Also, $\bar{W}(t - \lambda)$ would be obtained numerically as a function of its argument. A more promising approach would seem to be the evaluation of the matrix $\bar{A}$ (or its Jordan canonical form). Given $\bar{A}$ numerically, the analytical form of the weighting function, transfer function or differential equation is easily constructed. Noise problems, however, could be a severe limitation if the determination of $\bar{A}$ depends on only two points of the response curves; moreover, little of the available information would be utilized. If one could obtain "average" results by utilization of the whole response curves, better results could be expected. This idea constitutes the basis for the forthcoming discussion.

Postmultiplication of Equation (2-53) by $\Phi(\lambda)$ yields

$$\Phi(t) = e^{(t-\lambda)\bar{A}} \Phi(\lambda) \quad (2-70)$$

Let $\delta = t - \lambda$ (where $\delta = \text{constant}$), then

$$\Phi(t) = e^{\delta\bar{A}} \Phi(\lambda) \quad (2-70a)$$

Since δ is considered constant it follows that λ becomes a linear function of t . In addition the matrix function $e^{\delta\bar{A}}$ becomes a constant matrix. Integration of Equation (2-70a) on both sides with respect to t from $t = t_1$ to $t = t_1 + \delta$ yields:

$$\int_{t_1}^{t_1+\delta} \Phi(t) dt = e^{\delta\bar{A}} \int_{t_1}^{t_1+\delta} \Phi(\lambda) dt \quad (2-71)$$

Substituting λ for t on the right side of Equation (2-71) and

noticing that $dt - d\lambda = 0$ and $dt = d\lambda$ and

$$\begin{aligned} \text{for } t = t_1, \quad \lambda = t_1 - \delta \\ \text{and for } t = t_1 + \delta, \quad \lambda = t_1 \end{aligned}$$

$$\text{one obtains } \int_{t_1}^{t_1+\delta} \Phi(t) dt = e^{\delta \bar{A}} \int_{t_1-\delta}^{t_1} \Phi(\lambda) d\lambda \quad (2-72)$$

The above integration procedure can be repeated, leading finally to the result

$$\begin{aligned} \int_{t_{n-1}}^{t_{n-1}+\delta} dt_{n-2} \dots \int_{t_1}^{t_1+\delta} dt \Phi(t) = e^{\delta \bar{A}} \int_{t_{n-1}}^{t_{n-1}+\delta} dt_{n-2} \dots \\ \int_{t_2}^{t_2+\delta} dt_1 \int_{t_1-\delta}^{t_1} d\lambda \Phi(\lambda) \end{aligned} \quad (2-73)$$

or in compact notation

$$\overline{\Phi}_{n-1}(t_{n-1}) = e^{\delta \bar{A}} \overline{\Phi}_{n-1}(t_{n-1} - \delta) \quad (2-74)$$

where

$$\overline{\Phi}_{n-1}(t_{n-1}) = \int_{t_{n-1}}^{t_{n-1}+\delta} dt_{n-2} \dots \int_{t_1}^{t_1+\delta} dt \Phi(t) \quad (2-75)$$

and

$$\overline{\Phi}_{n-1}(t_{n-1} - \delta) = \int_{t_{n-1}}^{t_{n-1}+\delta} dt_{n-2} \dots \int_{t_2}^{t_2+\delta} dt_1 \int_{t_1-\delta}^{t_1} d\lambda \Phi(\lambda) \quad (2-76)$$

The matrix $\Phi(t)$ can be written as:

$$\Phi(t) = \begin{vmatrix} \varphi_1(t) & \dots & \varphi_n(t) \\ \dot{\varphi}_1(t) & \dots & \dot{\varphi}_n(t) \\ \vdots & \dots & \vdots \\ \varphi_1^{(n-1)}(t) & \dots & \varphi_n^{(n-1)}(t) \end{vmatrix} = [\varphi_{ij}(t)]$$

where $\varphi_{ij}(t)$ stands for the element of the i^{th} row and the j^{th} column, and:

$$\varphi_{ij}(t) = \varphi_j^{(i-1)}(t)$$

where the term on the right represents the $(i-1)^{\text{th}}$ derivative with respect to time of the j^{th} solution.

It follows from the above that

$$\overline{\Phi}_{n-1}(t_{n-1}) = [\overline{\varphi}_{1j}(t_{n-1})] \quad (2-77)$$

where

$$\overline{\varphi}_{1j}(t_{n-1}) = \int_{t_{n-1}}^{t_{n-1}+\delta} dt_{n-2} \dots \int_{t_1}^{t_1+\delta} dt \varphi_j^{(i-1)}(t) . \quad (2-78)$$

Similarly

$$\overline{\Phi}_{n-1}(t_{n-1} - \delta) = [\overline{\varphi}_{1j}(t_{n-1} - \delta)]$$

where

$$\overline{\varphi}_{1j}(t_{n-1} - \delta) = \int_{t_{n-1}}^{t_{n-1}+\delta} dt_{n-2} \dots \int_{t_2}^{t_2+\delta} dt_1 \int_{t_1-\delta}^{t_1} d\lambda \varphi_j^{(i-1)}(\lambda) . \quad (2-80)$$

The advantages of integrating the solution matrix can

now be recognized. The repeated integration of $\Phi(t)$ that yields Equation (2-74) eliminates the need for numerical differentiation of the response curves. Rather, integration of the data is performed, thus smoothing the results and cancelling out random noise effects. Also, all the available information over a preselected interval $n\delta$ is utilized in the evaluation, thereby obtaining average values for $e^{\delta\bar{A}}$ if drifts or nonlinear effects are present.

Since t_{n-1} and δ in Equation (2-74) are preselected constants, the matrices $\Phi_{n-1}(t_{n-1})$ and $\Phi_{n-1}(t_{n-1} - \delta)$ are constant (numerical) matrices that can be directly determined from the response curves and their integrals. Given the above matrices, Equation (2-74) can be rewritten to evaluate $e^{\delta\bar{A}}$:

$$\Phi_{n-1}(t_{n-1}) \Phi_{n-1}(t_{n-1} - \delta)^{-1} = e^{\delta\bar{A}} = H \quad (2-81)$$

Given H numerically, $\bar{A}$ can be evaluated by the following procedure. H is transformed into its Jordan canonical form, from which the characteristic values of $\bar{A}$ are determined from Equations (2-56)-(2-60).

From the characteristic values $\rho_1, \dots, \rho_n$ of $\bar{A}$, which are actually the natural frequencies (the poles) of the system, the coefficients of the differential system, Equation (2-69), can be determined. Given the differential equation, the transfer function and the weighting function can easily be evaluated in their analytical form.

The foregoing discussion constitutes the foundation

for the proposed identification technique. In the following chapters the discussion will center exclusively around the identification of the General Model of constant parameter systems.

Perturbation Models and Steady State Operation

The foregoing discussion centered around a linear system describable by Equation (2-5):

$$L^{(n)}[y(t)] = \sum_{j=1}^m M_j^{(k_j)}[x_j(t)] \quad (2-5)$$

When the forcing functions are removed, i.e., $x_j(t) = 0$, $j = 1, \dots, m$, the homogeneous equation is obtained:

$$L^{(n)}[y(t)] = 0 \quad (2-6)$$

In reality, systems possess a finite steady state operating level, and when the forcing functions (or inputs) are removed, the plant ceases to function. Consider a plant with output $\bar{y}(t)$ and a set of inputs $\bar{x}_j(t)$. At steady state all derivatives of the plant variables are zero and thus:

$$a_n \bar{y} = \sum_{j=1}^m g_{jk_j} \bar{x}_j \quad (2-82)$$

Equation (2-82) describes the steady state mass, energy and entropy balances of the system.

One can consider the plant variables to be composed of a steady state part, and a perturbation part which describes the deviation from steady state

$$\begin{aligned} \bar{y}(t) &= y_{ss} + y_p(t) \\ \bar{x}_j(t) &= x_{jss} + x_{jp}(t) \quad j = 1, \dots, m \end{aligned} \quad (2-83)$$

Equation (2-5) then becomes

$$L^{(n)}[y_{ss} + y_p(t)] = \sum_{j=1}^m M_j^{(k_j)} [x_{jss} + x_{jp}(t)] \quad (2-84)$$

But since y_{ss} and x_{jss} are constants, Equation (2-84) degenerates to

$$L^{(n)}[y_p(t)] + a_n y_{ss} = \sum_{j=1}^m M_j^{(k_j)} [x_{jp}(t)] \sum_{j=1}^m g_{jk_j} x_{jss} \quad (2-85)$$

Substitution of Equation (2-82) into (2-85) yields

$$L^{(n)}[y_p(t)] = \sum_{j=1}^m M_j^{(k_j)} [x_{jp}(t)] \quad (2-86)$$

which is the same as Equation (2-5) above. Thus Equation (2-5) can be considered as a perturbation model of the system. In order to reconstruct the total model of the system the steady state conditions have to be known.

When a process is to be identified, the total variables are usually measured rather than the perturbations. Since the identification is performed with respect to the perturbation model, the steady state operating conditions have to be measured independently and subtracted.

Preliminary Evaluation of the Technique

In the present section a preliminary discussion of the identification method is presented which is based only on interpretation of the theoretical principles. Some of the advantages and possible limitations of the proposed method are discussed and evaluated. Part of the argumentation in this section is intuitive and is intended primarily to

introduce the reader to some of the problems which have been studied. A detailed analysis of these problems is presented in later chapters.

Certain mathematical procedures which are theoretically exact are very sensitive from a numerical calculation viewpoint and may become doubtful or even worthless when uncertainty in the data or truncation are involved. In contrast to numerical differentiation which suffers from this limitation, numerical integration is considered to be a smoothing process. However, integration does not really "smooth out" random noise but in reality changes its relative significance. When the noise is truly random, i.e., the mean value of it is zero, then integration may reduce its effect or even eliminate it entirely. The important point is that no information is lost. When the noise in the response is biased, i.e., it is not randomly distributed around the true value, then integration may act as a noise amplifier. The relative magnitude of error may then increase even to the point where all useful information is lost. Fortunately the bias problem can be greatly reduced by proper calibration of the measuring devices. The noise frequency also affects the resultant integral curves, and it may be expected that the higher the frequency the better will be the results of integration.

It was previously shown that by proper integration of the fundamental solution matrix the need for data differentiation is eliminated. For this technique $n-1$ integrations are required for this purpose for an n^{th} order model. There

is however no theoretical limit for the permitted number of integrations. If integration results in suppression of noise, it might be desirable to perform more than $n-1$ integrations. The answer to this problem is, naturally, dependent on the general error propagation into the integrals and cannot be given a priori. Chapter IV deals with estimation of error propagation into integrals of response functions and provides quantitative answers to many of the questions raised above.

In performing the integrations, the limits t_{n-1} and δ , Equation (2-74), were arbitrarily selected, thereby leaving two degrees of freedom in constructing the fundamental matrices. Naturally it would be desirable to select t_{n-1} and δ such as to minimize the error in the identified model. It is at this stage premature to discuss this problem any further since it strongly depends on the propagation of error. It will therefore be postponed for later chapters.

In conventional methods for process identification the system poles and zeros for a certain input-output pair are evaluated simultaneously from the frequency response function (usually from Bode plots). Under certain circumstances this procedure may be a drawback since the results of identification are dependent on the dynamic relationship between a particular input-output pair. In the present method the poles and zeros are identified in two separate procedures, and therefore the identification of the system poles (or General Model) is independent of any particular input-output relationship. The forcing variables can be

selected at will to obtain the best response curves. Usually it is desirable to force those inputs that are best coupled with the output under consideration. Since relaxed (unforced) responses are used to identify the system poles, it is even possible to pulse plant parameters. Of course, pulsation of several variables simultaneously in order to force more energy into the system, and thereby obtaining stronger response signals, is also possible. The flexibility in testing may be of great advantage in overcoming noise problems.

The foregoing discussion also hints at an answer to the problem of obtaining n , linearly independent response curves. Theoretically linearly independent responses can be obtained by forcing certain input variables with pulses that have different power spectral densities. Thereby different natural frequencies are forced differently each time. In practice the responses may not be sufficiently different (linearly independent) when noise is present. The answer to the problem is to force different variables (or combination of variables) to obtain the desired responses. Experimental evaluation of the problem is presented later.

CHAPTER III

ERROR PROPAGATION ANALYSIS

Common methods for system identification do not offer an error propagation analysis for estimation of the reliability of the identified model. The importance of the knowledge of the uncertainty which is in the experimentally determined model has been discussed previously. In this work, particular emphasis has been placed on this problem.

The error propagation analysis presented in this chapter serves a two-fold purpose:

- (1) it provides an estimation of the model reliability;
- (2) it serves as a tool for the actual execution of the identification procedure.

In its second capacity, the error analysis procedure insures a minimization of the propagation of the error that is inherent in the response functions into the identified parameters. This error minimization is achieved by a systematic search through the response functions for the location where the contribution of the individual parameters (natural frequencies) is most significant. Thus the error associated with high natural frequencies would be minimized when the initial segments of the response function is utilized.

Conversely the error of the lower natural frequencies is minimized when wider segments of the response function are used. In addition to the latter use of the error propagation analysis, it is also valuable in determining the actual order of system when it is not known before the identification is carried out.

In a mathematical sense, the word error is generally used to mean a deviation between an observed quantity and the true value which it represents. When the true value is known, the error can be explicitly defined. When only a range of values within which the true value lies is known, the error loses its exact meaning. In this case one can only speak of an uncertainty which possesses given bounds. Obviously, every observation must lie within the permitted range for the true value. Under these circumstances one does not know generally whether a particular observation really possesses an error. All that can be said is that a certain amount of uncertainty is associated with the observation. When a large number of similar observations is made, a distribution of values will be obtained which allows the following quantities to be defined: 1) the mean value of the observations, which is statistically the best representation for the true value, 2) the expected deviation (or the expected magnitude of the error) which, loosely speaking, is the most likely value of the error, and 3) the extreme error which represents the largest obtainable value for the error. Error propagation analysis can be based on the expected error or on the

extreme error. When it is based on the expected error, an expected value for the error has to be determined at each stage of the computation. Similarly, when the error propagation analysis is based on extreme error, the error must be maximized at each stage of the computation. The present work is concerned with the latter approach.

There are two different kinds of primary error sources which will affect the results of any numerical computation and which differ in character, origin and significance. One such error is the uncertainty of the data. This error is largely responsible for the uncertainty associated with the final results. It does, however, not have any influence on the numerical values of the results since the computations are based only on the nominal values of the data. Another possible error source is roundoff and truncation error which affects the numerical values of the results. Roundoff error may become significant when the degree of accuracy of each computational step is low or when a large number of consecutive computations is performed.

In the present analysis the effects of roundoff error are totally insignificant compared to the error due to data uncertainty. This latter error is larger than the former by several orders of magnitude. With this as the basis roundoff and truncation error will therefore be neglected in the present analysis.

Before proceeding with the detailed error analysis it is desirable to summarize the computational steps involved

in determining the system parameters from the response functions. The computational steps are:

- 1) Evaluation of the matrices $\overline{\Phi}_{n-1}(t_{n-1})$ and $\overline{\Phi}_{n-1}(t_{n-1} - \delta)$ of Equations (2-77) and (2-79) from the response curves and their integrals.
- 2) Inversion of the matrix $\overline{\Phi}_{n-1}(t_{n-1} - \delta)$.
- 3) Postmultiplication of $\overline{\Phi}_{n-1}(t_{n-1})$ by $(\overline{\Phi}_{n-1}(t_{n-1} - \delta))^{-1}$
- 4) Evaluation of the characteristic roots of $e^{\delta \bar{A}}$.
- 5) Construction of the Jordan canonical form of $e^{\delta \bar{A}}$.
- 6) Evaluation of the characteristic roots (natural frequencies) of the system from the Jordan canonical form, Equations (2-59) and (2-60).

The first step of the computation involves the selection of integration limits t_{n-1} and δ . This problem will be discussed at later stages and for the time being it is sufficient to assume t_{n-1} and δ are given values. The construction of the matrices $\overline{\Phi}_{n-1}(t_{n-1})$ and $\overline{\Phi}_{n-1}(t_{n-1} - \delta)$ is straightforward and follows directly from their definitions. There is no particular difficulty associated with the evaluation of these matrices although the error bounds of the data have to be evaluated at this stage. In Chapter IV a method for estimation of error propagation into the integrals of the response curves is presented.

When a matrix has multiple characteristic roots, the Jordan form is its canonical representation. From a numerical viewpoint, however, it is justified to assume that every $n \times n$ matrix is similar to a diagonal matrix regardless of the

multiplicity of the eigenvalues. The reason for this statement is that if there are multiple roots, they can be made distinct by infinitesimal perturbations of the matrix elements. The result is, of course, an infinitesimal change in the values of the eigenvalues, which when considered from a practical viewpoint is negligible. In the present analysis the assumption of distinct eigenvalues will be made since it leads to much simpler numerical manipulations.

The error propagation analysis will be carried out individually for each computational step and a general notation will be used for convenience.

Error Bounds of the Inverse of $\Phi_{n-1}(t_{n-1} - \delta)$

The matrix $\Phi_{n-1}(t_{n-1} - \delta)$ represents an approximate value* for the integrated fundamental matrix. For convenience of notation denote $\Phi_{n-1}(t_{n-1} - \delta)$ by $\bar{\Phi}$. Then if the true value of the integrated fundamental matrix is represented by T , the error matrix E will be

$$E = T - \bar{\Phi}. \quad (3-1)$$

The matrix E is composed of unknown elements e_{ij} which may be positive, negative or zero, but which are not greater in magnitude than some bounds η_{ij} , which are the extreme error estimates for the individual elements of the response curves and their integrals. It is desired to find the bounds of

*Since the correct value is unknown due to the existing uncertainty.

variation of the elements of the inverse of T knowing the value of $\bar{\Phi}$ and the bounds of E . The problem can be stated alternatively as how does the inverse of $\bar{\Phi}$ deviate from the inverse of T (as a function of the elements of $\bar{\Phi}$ and E). Of course $\bar{\Phi}$ and E are assumed to be nonsingular matrices. It will be assumed that $\bar{\Phi}^{-1}$ can be obtained by exact methods, that is if $\bar{\Phi}$ is given, $\bar{\Phi}^{-1}$ can be evaluated with absolute accuracy or at least correct to a desired number of decimal places.

From Equation (3-1) it follows that

$$T^{-1} = (\bar{\Phi} + E)^{-1} = [(I + E\bar{\Phi}^{-1})\bar{\Phi}]^{-1} = \bar{\Phi}^{-1}(I + E\bar{\Phi}^{-1})^{-1}, \quad (3-2)$$

where I is the identity matrix. If the elements of E are small enough so that any norm of $E\bar{\Phi}^{-1}$ is less than 1, it can be shown [16,53] that T^{-1} can be expressed in terms of the convergent series

$$T^{-1} = \bar{\Phi}^{-1}[I - E\bar{\Phi}^{-1} + (E\bar{\Phi}^{-1})^2 - (E\bar{\Phi}^{-1})^3 + \dots] \quad (3-3)$$

Thus

$$\begin{aligned} \bar{\Phi}\bar{\Phi}^{-1}[I - E\bar{\Phi}^{-1} + (E\bar{\Phi}^{-1})^2 - (E\bar{\Phi}^{-1})^3 + \dots] = \\ I - E\bar{\Phi}^{-1} + (E\bar{\Phi}^{-1})^2 - (E\bar{\Phi}^{-1})^3 + \dots = (I + E\bar{\Phi}^{-1}) \end{aligned} \quad (3-4)$$

The accumulation of inherent error in the inverse is given by the difference between T^{-1} and $\bar{\Phi}^{-1}$. If the discrepancy matrix (or the error of the inverse) is denoted by D , then

$$\begin{aligned}
D &= T^{-1} - \bar{\Phi}^{-1} = \bar{\Phi}^{-1} [I - E\bar{\Phi}^{-1} + (E\bar{\Phi}^{-1})^2 - (E\bar{\Phi}^{-1})^3 + \dots] - \bar{\Phi}^{-1} \\
&= -\bar{\Phi}^{-1} E\bar{\Phi}^{-1} [I - E\bar{\Phi}^{-1} + (E\bar{\Phi}^{-1})^2 - (E\bar{\Phi}^{-1})^3 + \dots] \\
&= -\bar{\Phi}^{-1} E\bar{\Phi}^{-1} [I + E\bar{\Phi}^{-1}]^{-1}. \quad (3-5)
\end{aligned}$$

Neglecting the higher order terms in Equation (3-5), the first order discrepancy (approximation) matrix D_1 is obtained as

$$D_1 = -\bar{\Phi}^{-1} E\bar{\Phi}^{-1}. \quad (3-6)$$

This equation can symbolically be written as

$$d(\bar{\Phi}^{-1}) = -\bar{\Phi}^{-1} (d\bar{\Phi})\bar{\Phi}^{-1}. \quad (3-6a)$$

(It is interesting to notice that the first order discrepancy matrix can be obtained formally by differentiating the equation $I = \bar{\Phi} \cdot \bar{\Phi}^{-1}$. Thus $0 = d\bar{\Phi} \cdot \bar{\Phi}^{-1} + \bar{\Phi} d(\bar{\Phi}^{-1})$, from which Equation (3-6a) follows immediately.)

The first order discrepancy matrix D_1 of Equation (3-6) is of much simpler form than the total discrepancy D , Equation (3-5), and should, therefore, be used whenever applicable. To estimate the goodness of D_1 it will be necessary to estimate the upper bounds of the discrepancy matrix D .

Define $N(P)$, the norm of a real matrix P , as the quantity which satisfies the conditions:

- a) $N(P) \geq 0$
 - b) $N(P) = 0$ if and only if $P = 0$
 - c) $N(cP) = |c|N(P)$, where c is a scalar.
- (3-7)

- d) $N(P_1 \pm P_2) \leq N(P_1) + N(P_2)$
e) $N(P_1 P_2) \leq N(P_1) N(P_2)$ (3-7)
f) $N(P) = 1$ if P is a fundamental unit continued
matrix (with all elements zero except
for one element which is unity).

Several norms that satisfy the above conditions are available. Two norms used in the present work are

$$N_T(P) = n \cdot \max_{i,j} |p_{ij}| \quad (3-8)$$

where n is the order of the matrix P , and

$$N_h(P) = \left[\sum_{i,j} p_{ij}^2 \right]^{\frac{1}{2}} \quad (3-9)$$

Application of conditions (3-7) to (3-5) gives with $N(E\bar{\Phi}^{-1}) < 1$.

$$N(D) \leq \frac{N(\bar{\Phi}^{-1})N(E\bar{\Phi}^{-1})}{1 - N(E\bar{\Phi}^{-1})}, \quad (3-10)$$

and if $N(E)N(\bar{\Phi}^{-1}) < 1$ it follows that

$$N(D) \leq \frac{N(E)[N(\bar{\Phi}^{-1})]^2}{1 - N(E)N(\bar{\Phi}^{-1})} \quad (3-11)$$

Equation (3-11) represents an upper bound for the discrepancy matrix D and similarly the numerator of this equation represents an upper bound for the first order discrepancy matrix D_1 . Thus one may write:

$$\text{Bound } (D) = \frac{\text{Bound } (D_1)}{1 - N(E)N(\bar{\Phi}^{-1})}. \quad (3-12)$$

The replacement of the inequality sign in Equation (3-11) by the equality sign in (3-12) is permitted since it results in an even more conservative estimate.

From Equation (3-12) it is clear that the first order discrepancy estimate is very unsatisfactory when $N(E)N(\bar{\Phi}^{-1}) \rightarrow 1$ and is of doubtful value unless $N(E)N(\bar{\Phi}^{-1}) < 0.1$. If $N(E)N(\bar{\Phi}^{-1}) < 0.1$, the first order discrepancy bounds will underestimate the true discrepancy by 11 per cent or less.

In the foregoing discussion it was shown how the discrepancy of the inverse matrix can be evaluated. It has yet to be shown how the extreme discrepancies are determined. The discussion of this subject will first focus on the first order discrepancy approximation, and will then be extended to the general case.

The first order discrepancy matrix was given by Equation (3-6) as $-\bar{\Phi}^{-1} E \bar{\Phi}^{-1}$. Formal expansion of this product in terms of the $k\ell$ th element gives

$$d_{1:k\ell} = -\bar{\Phi}_k^{-1} E \bar{\Phi}_\ell^{-1} = \sum_i \sum_j \bar{\Phi}_{ki}^{-1} e_{ij} \bar{\Phi}_{j\ell}^{-1}, \quad (3-13)$$

where by $\bar{\Phi}_{ki}^{-1}$ is meant the ki th element of $\bar{\Phi}^{-1}$. Since one is free to select the signs of the elements e_{ij} so as to obtain the extreme $d_{1:k\ell}$, it is possible to select consistently the sign of each element to be equal to, or opposite to, the sign of $\bar{\Phi}_{ki}^{-1} \bar{\Phi}_{j\ell}^{-1}$.

Define:

$$\theta_{ij} = \begin{cases} 1 & \text{When } \bar{\varphi}_{ki}^{-1} \text{ has the same sign as } \bar{\varphi}_{jl}^{-1} \\ -1 & \text{When } \bar{\varphi}_{ki}^{-1} \text{ has the opposite sign to} \\ & \text{that of } \bar{\varphi}_{jl}^{-1} \\ 0 & \text{When either } \bar{\varphi}_{ki}^{-1} \text{ or } \bar{\varphi}_{jl}^{-1} \text{ is zero.} \end{cases} \quad (3-14)$$

Then $e_{ij} = \theta_{ij} |e_{ij}|$.

The kl^{th} element of the extreme first order discrepancy matrix can then be written as:

$$d_{1:kl} = -\sum_i \sum_j \bar{\varphi}_{ki}^{-1} \theta_{ij} |e_{ij}| \bar{\varphi}_{jl}^{-1}. \quad (3-15)$$

If $|e_{ij}|$ take on values as large as the maximum bounds η_{ij} , Equation (3-15) becomes

$$d_{1:kl} = -\sum_i \sum_j \bar{\varphi}_{ki}^{-1} \theta_{ij} \eta_{ij} \bar{\varphi}_{jl}^{-1} \quad (3-16)$$

Theoretically it is necessary to fix the sign of each e_{ij} in order to obtain the extreme discrepancy for some particular element. Sometimes these signs may maximize the discrepancy associated with some other elements of D_1 , or even all the elements, but this possibility is not generally true. To illustrate this point consider the simple second order matrix

$$\bar{\Phi}^{-1} = \begin{vmatrix} 2 & 5 \\ -3 & 1 \end{vmatrix}.$$

The error matrix has all elements equal to 1 in

absolute value. It is a simple matter to show that the extreme value for $d_{1:11}$ is 35. It is obtained when the matrix Θ is

$$\Theta = \begin{vmatrix} 1 & -1 \\ 1 & -1 \end{vmatrix}$$

and that the extreme value for $d_{1:12}$ is 42 which is obtained when Θ is

$$\Theta = \begin{vmatrix} 1 & 1 \\ 1 & 1 \end{vmatrix} .$$

It is thus obvious that in general not all elements of D_1 will attain their extreme value simultaneously. Since in practice the elements of Θ can possess only one set of values, there arises the question as to what should be the criterion for selection of Θ . For the case where all the elements of D_1 attain their extreme value simultaneously the answer to this question is quite obvious. For the more general case the situation is not so simple. Recalling that the ultimate purpose of the error propagation analysis is the evaluation of the extreme attainable errors of the system's poles, it becomes evident that the answer should depend on the error behavior in the subsequent computational steps. As will be seen later (cf. Equation (3-33)), the error of the characteristic roots of the system is proportional to the norm of the error of the product of $\overline{\Phi_{n-1}(t_{n-1})}$ and $\overline{\Phi_{n-1}(t_{n-1} - \delta)}^{-1}$. Therefore it is necessary to select Θ such as to maximize the norm of this error term which will later

be shown to be given by

$$D_h = \bar{\Phi}_1 D + E_1 \bar{\Phi}^{-1} \quad (3-22)$$

where D_h is the error of the product $\overline{\Phi_{n-1}(t_{n-1})} \cdot \overline{\Phi_{n-1}(t_{n-1} - \delta)}^{-1}$, $\bar{\Phi}_1$ stands for $\overline{\Phi_{n-1}(t_{n-1})}$ and E_1 represents the error associated with $\bar{\Phi}_1$. D and $\bar{\Phi}^{-1}$ have the same meaning as before.

Utilization of the properties of Equation (3-7) gives:

$$\begin{aligned} N(D_h) &= N(\bar{\Phi}_1 D + E_1 \bar{\Phi}^{-1}) \leq N(\bar{\Phi}_1 D) + N(E_1 \bar{\Phi}^{-1}) \\ &\leq N(\bar{\Phi}_1)N(D) + N(E_1 \bar{\Phi}^{-1}). \end{aligned} \quad (3-17)$$

Thus it can be concluded that the matrix Θ must be selected such that the norm of the discrepancy matrix D be maximized. Normally this condition also implies the maximization of the norm of D_1 . The criterion is a rather loose one since it leaves the choice of the norm to be arbitrary, and there is no guarantee that the use of different norms will yield consistent results. It is, however, the best general criterion that can be established from theoretical considerations, and it can be expected that in general any norm will be adequate. In the present work the norm $N_T(D_1)$ of Equation (3-8) was selected for convenience. Thus the selection of Θ will be based on the criterion

$$\max_{\Theta} N_T(D_1) = \max_{\Theta} [n \cdot \max_{1,j} |d_{1j}|] \Rightarrow \max_{\Theta} \cdot \max_{1,j} |d_{1j}| \quad (3-18)$$

The maximization of the first order discrepancy matrix can now be summarized as follows: Each element $d_{1:k\ell}$ is maximized individually according to (3-14) and (3-16).

The matrix Θ that corresponds to the element $d_{1:kl}$ that is largest in absolute value is then used to recalculate the elements of D_1 .

The calculation of the extreme discrepancy matrix D can be accomplished by means of Equation (3-5) using the error matrix E with the corresponding signs that have been evaluated for the first order approximation D_1 . Thus

$$D = D_1 [I - E\bar{\Phi}^{-1} + (E\bar{\Phi}^{-1})^2 - (E\bar{\Phi}^{-1})^3 + \dots] = D_1 [I + E\bar{\Phi}^{-1}]^{-1}, \quad (3-19)$$

and the calculation actually involves the evaluation of a correction matrix $[I + E\bar{\Phi}^{-1}]^{-1}$. In practice it may be more convenient to compute the convergent series directly.

There may arise the question as to whether the maximization of the first order discrepancy matrix and the total discrepancy matrix occur simultaneously (with the same sign matrix Θ). Obviously if this condition is not true, Equation (3-19) will not yield the extreme total discrepancy. Only if any one of the elements of D is larger than the corresponding elements in $\bar{\Phi}^{-1}$, will D and D_1 not necessarily be maximized simultaneously [16]. However, if such is the case, the inverse $\bar{\Phi}^{-1}$ is usually totally unreliable and no further consideration of the problem is justified.

Error Bounds of $e^{\delta\bar{A}}$

Given the matrix $\bar{\Phi}^{-1} = \bar{\Phi}_{n-1}(\bar{t}_{n-1} - \delta)^{-1}$ with bounds of error D and the matrix $\bar{\Phi}_1 = \bar{\Phi}_{n-1}(\bar{t}_{n-1})$ with bounds of error E_1 . It is desired to find the extreme error bounds

for the product

$$H = \bar{\Phi}_1 \bar{\Phi}^{-1} = e^{\delta \bar{A}}. \quad (3-20)$$

If the error associated with H is expressed by D_h one may write

$$H + D_h = (\bar{\Phi}_1 + E_1)(\bar{\Phi}^{-1} + D) = \bar{\Phi}_1 \bar{\Phi}^{-1} + E_1 \bar{\Phi}^{-1} + \bar{\Phi}_1 D + E_1 D. \quad (3-21)$$

The error matrix D has been evaluated in the previous section and is fixed in magnitude and signs of its elements. The error matrix E_1 is composed of unknown elements $e_{1:1j}$ which are bounded in magnitude by η_{1j} , where the elements η_{1j} are identical to the error bounds of the matrix $\bar{\Phi}$ in the previous section.

If the error matrices E_1 and D are relatively small compared to the other components of Equation (3-21), their product can be neglected and after subtraction of (3-20) from (3-21)

$$D_h = \bar{\Phi}_1 D + E_1 \bar{\Phi}^{-1} \quad (3-22)$$

is obtained. Formal expansion of (3-22) in terms of the elements yields

$$d_{h:1j} = \sum_k (\bar{\varphi}_{1:1k} d_{kj} + e_{1:1k} \bar{\varphi}_{kj}^{-1}) \quad (3-23)$$

where as before $\bar{\varphi}_{kj}^{-1}$ stands for the kj^{th} element of $\bar{\Phi}^{-1}$.

As will be seen in the next section (Equation (3-33)), the errors of the characteristic roots of H are proportional to the norm of the error matrix D_h . Therefore in order to

find the extreme errors it is necessary to adjust the signs of $e_{1:ik}$ such as to maximize the norm of D_h . Maximization of any norm of D_h will be achieved if every element $d_{h:ij}$ attains its largest possible magnitude. However as can be easily seen, not all the elements of D_h are necessarily maximized simultaneously.

Let $e_{1:ik} = \theta_{1:ik} \cdot |e_{1:ik}|$ and define $\theta_{1:ik}$ such that

$$\theta_{1:ik} = \begin{cases} 1 & \text{if the product } \bar{\varphi}_{1:ik} d_{kj} \bar{\varphi}_{kj}^{-1} \text{ is positive} \\ -1 & \text{if the product } \bar{\varphi}_{1:ik} d_{kj} \bar{\varphi}_{kj}^{-1} \text{ is negative} \\ 0 & \text{if } \bar{\varphi}_{kj}^{-1} \text{ is zero.} \end{cases} \quad (3-24)$$

Then

$$d_{h:ij} = \sum_k (\bar{\varphi}_{1:ik} d_{kj} + \theta_{1:ik} |e_{1:ik}| \bar{\varphi}_{kj}^{-1}) \quad (3-25)$$

and if $|e_{1:ik}|$ attain their maximum value η_{1k} (3-25) becomes:

$$d_{h:ij} = \sum_k (\bar{\varphi}_{1:ik} d_{kj} + \theta_{1:ik} \eta_{1k} \bar{\varphi}_{kj}^{-1}) \quad (3-26)$$

Maximization of any element $d_{h:ij}$ determines the corresponding i^{th} row of Θ_1 which, in turn, fixes the i^{th} row of the matrix D_h . Sometimes several or even all the elements of a row of D_h are maximized simultaneously but this situation is not in general true. Furthermore, all the other rows of Θ_1 remain arbitrary and unaffected by this process.

In the present work the norm $N_T(D_h)$ which was defined by Equation (3-8) was selected as the criterion for

maximization of the error. Thus it is necessary to evaluate the matrix Θ_1 such that the largest element of D_h is maximized. To find the largest possible element of D_h it is necessary to determine the attainable extreme of each element $d_{h:ij}$ individually (and independently of any other elements). Since, however, the maximization of the elements is only row dependent, that is, only the elements of one row of Θ_1 are fixed by fixing an element in the corresponding row of D_h , the calculation process can be carried out row by row. Thus the largest element $d_{h:ik}$ of the i^{th} row of D_h is found, and the i^{th} row of Θ_1 corresponding to this element is used to recalculate all the other elements of this row in D_h . The process is then repeated for all other rows of the discrepancy matrix D_h .

Error Bounds of the Eigenvalues of $e^{\delta\bar{A}}$

The error bounds for the matrix function $e^{\delta\bar{A}}$ have been found in the previous section. In this section the error bounds for the eigenvalues of $e^{\delta\bar{A}}$ will be determined.

Given the real matrix $e^{\delta\bar{A}} = H$ which is assumed to possess distinct eigenvalues* $v_1, \dots, v_n$ and eigenvectors $v_1, \dots, v_n$, then by definition of the eigenvectors

$$Hv_1 = v_1v_1. \quad (3-27)$$

Also let $\underline{v}_1, \dots, \underline{v}_n$ be the eigenvectors of H^T (the transpose of H). The error bounds of the elements of H are expressed

*Cf. discussion on page 62.

by the matrix D_h . If D_h is a small perturbation of H then a typical eigenvector of $H + D_h$ can be expressed as $v_1 + \sum_{i \neq j} \epsilon_{ij} v_j$ where ϵ_{ij} are small constants. The eigenvalues that correspond to this eigenvector is expressed as $v_1 + \Delta v_1$. An equation similar to (3-27) can then be written to give

$$(H + D_h)(v_1 + \sum_{i \neq j} \epsilon_{ij} v_j) = (v_1 + \Delta v_1)(v_1 + \sum_{i \neq j} \epsilon_{ij} v_j) \quad (3-28)$$

Expansion by terms of Equation (3-28) and utilization of (3-27) yields

$$\begin{aligned} \sum_{i \neq j} \epsilon_{ij} H v_j + D_h v_1 + D_h \sum_{i \neq j} \epsilon_{ij} v_j &= \Delta v_1 v_1 + v_1 \sum_{i \neq j} \epsilon_{ij} v_j \\ &+ \Delta v_1 \sum_{i \neq j} \epsilon_{ij} v_j. \end{aligned} \quad (3-29)$$

The last terms on both sides of Equation (3-29) are products of small perturbation terms and can be neglected. Thus

$$\sum_{i \neq j} \epsilon_{ij} H v_j + D_h v_1 = \Delta v_1 v_1 + \sum_{i \neq j} \epsilon_{ij} v_1 v_j. \quad (3-30)$$

Premultiplication of Equation (3-30) by $\underline{v}_1^T$ (the transpose of $\underline{v}_1$) yields

$$\underline{v}_1^T D_h v_1 = \underline{v}_1^T \Delta v_1 v_1, \quad (3-31)$$

since the product $\underline{v}_1^T \cdot v_j = 0$ for $i \neq j$ [17].

Finally the error of the eigenvalues is given by

$$\Delta v_1 = \frac{\underline{v}_1^T D_h v_1}{\underline{v}_1^T v_1} = \frac{(D_h v_1, \underline{v}_1)}{(v_1, \underline{v}_1)} \quad (3-32)$$

where $(v_1, \underline{v}_1) = v_{11} v_{11} + \dots + v_{1n} v_{1n}$ is the inner product of v_1 and $\underline{v}_1$, and similarly for $(D_h v_1, \underline{v}_1)$.

Equation (3-32) is a first order approximation for the error of the eigenvalues. A higher order approximation would result in considerable complication of the equations. Fortunately, it is unnecessary for the purposes of the present error estimates.

It is interesting to study the upper bounds for the error in Equation (3-32). Application of the definitions for the norm of a matrix (Equation (3-7)) and definitions for the norm of a vector [17] to (3-32) yields:

$$|\Delta v_1| \leq \frac{N(D_h) \cdot \|v_1\| \cdot \|\underline{v}_1\|}{|(v_1, \underline{v}_1)|} = \underline{C}_1 N(D_h), \quad (3-33)$$

where

$$\underline{C}_1 = \frac{\|v_1\| \cdot \|\underline{v}_1\|}{|(v_1, \underline{v}_1)|} = \frac{1}{|\cos \beta|}, \quad (3-33a)$$

$|\Delta v_1|$ and $|(v_1, \underline{v}_1)|$ are absolute values of the corresponding quantities, and $\|v_1\|$ and $\|\underline{v}_1\|$ are norms of the vectors v_1 and $\underline{v}_1$ respectively defined as $\|v_1\| = (v_{11}^2 + \dots + v_{1n}^2)^{\frac{1}{2}}$ etc. The quantity $\underline{C}_1$ is frequently called the distortion coefficient for the matrix D_h corresponding to the eigenvalue v_1 . The angle β is the angle between the vectors v_1 and $\underline{v}_1$.

Thus the error in the eigenvalues is proportional to the norm of the error matrix D_h . It should be noticed that all the quantities, except for D_h , on the right hand side of Equation (3-33) are related only to the nominal values of the data and not to the error terms.

If the matrix H has complex eigenvalues, Equation

(3-32) has to be modified for finding the discrepancies of both the real and imaginary parts. Let

$$\left. \begin{aligned} v_1 &= \text{Re}(v_1) + j\text{Im}(v_1) \\ v_1 &= \text{Re}(v_1) + j\text{Im}(v_1) \\ \underline{v}_1 &= \text{Re}(\underline{v}_1) + j\text{Im}(\underline{v}_1) \end{aligned} \right\} \quad (3-34)$$

where Re and Im stand for the real and imaginary parts respectively and $j = \sqrt{-1}$. Substitution of the quantities of (3-34) in (3-32) yields

$$\begin{aligned} [\Delta\text{Re}(v_1) + j\Delta\text{Im}(v_1)][\text{Re}(v_1) + j\text{Im}(v_1)], [\text{Re}(\underline{v}_1) + j\text{Im}(\underline{v}_1)] = \\ = D_h[\text{Re}(v_1) + j\text{Im}(v_1)], [\text{Re}(\underline{v}_1) + j\text{Im}(\underline{v}_1)]. \end{aligned} \quad (3-35)$$

After expanding Equation (3-35) and equating the real and imaginary terms

$$\begin{aligned} \Delta\text{Re}(v_1)[\text{Re}(v_1), \text{Re}(\underline{v}_1) - \text{Im}(v_1), \text{Im}(\underline{v}_1)] - \Delta\text{Im}(v_1)[\text{Re}(v_1), \text{Im}(v_1) \\ + \text{Im}(v_1), \text{Re}(\underline{v}_1)] = D_h\text{Re}(v_1), \text{Re}(\underline{v}_1) - D_h\text{Im}(v_1), \text{Im}(\underline{v}_1) \end{aligned} \quad (3-36a)$$

and

$$\begin{aligned} \Delta\text{Re}(v_1)[\text{Re}(v_1), \text{Im}(\underline{v}_1) + \text{Im}(v_1), \text{Re}(\underline{v}_1)] + \Delta\text{Im}(v_1)[\text{Re}(v_1), \text{Re}(\underline{v}_1) \\ - \text{Im}(v_1), \text{Im}(\underline{v}_1)] = D_h\text{Re}(v_1), \text{Im}(\underline{v}_1) + D_h\text{Im}(v_1), \text{Re}(\underline{v}_1) \end{aligned} \quad (3-36b)$$

are obtained. Simultaneous solution of Equations (3-36a) and (3-36b) for $\Delta\text{Re}(v_1)$ and $\Delta\text{Im}(v_1)$ yields the final result:

$$\Delta\text{Re}(v_1) = \frac{\underline{p} \cdot \underline{q} + \underline{r} \cdot \underline{s}}{\underline{p}^2 + \underline{r}^2} \quad (3-37a)$$

and

$$\Delta \text{Im}(v_1) = \frac{\bar{p} \cdot \bar{s} - \bar{r} \cdot \bar{q}}{\bar{p}^2 + \bar{r}^2} \quad (3-37b)$$

where

$$\begin{aligned} \bar{p} &= \text{Re}(\underline{v}_1), \text{Re}(\underline{v}_1) - \text{Im}(\underline{v}_1), \text{Im}(\underline{v}_1) \\ \bar{q} &= D_h \text{Re}(\underline{v}_1), \text{Re}(\underline{v}_1) - D_h \text{Im}(\underline{v}_1), \text{Im}(\underline{v}_1) \\ \bar{r} &= \text{Re}(\underline{v}_1), \text{Im}(\underline{v}_1) + \text{Im}(\underline{v}_1), \text{Re}(\underline{v}_1) \\ \bar{s} &= D_h \text{Re}(\underline{v}_1), \text{Im}(\underline{v}_1) + D_h \text{Im}(\underline{v}_1), \text{Re}(\underline{v}_1) \end{aligned} \quad (3-38)$$

For real eigenvalues Equation (3-37a) reduces to (3-32) as expected.

Error Bounds of the System's Poles

Since the eigenvalues v_1 and consequently the poles ρ_1 are distinct, it follows from Equations (2-63) and (2-81) that

$$v_1 = e^{\delta \rho_1} \quad i = 1, \dots, n. \quad (3-39)$$

The eigenvalues v_1 may be real or complex and Equation (3-39) can in general be written as

$$\text{Re}(v_1) + j\text{Im}(v_1) = e^{\delta [\text{Re}(\rho_1) + j\text{Im}(\rho_1)]}. \quad (3-40)$$

When the poles are real, Equation (3-40) reduces to (3-39). Multiplication of Equation (3-40) by its complex conjugate equation yields

$$\text{Re}(v_1)^2 + \text{Im}(v_1)^2 = e^{2\delta \text{Re}(\rho_1)}, \quad (3-41)$$

from which the real part of the pole $\text{Re}(\rho_1)$ is obtained as

$$\text{Re}(\rho_1) = \frac{1}{2\delta} \log_e [\text{Re}(v_1)^2 + \text{Im}(v_1)^2]. \quad (3-42)$$

Differentiation of Equation (3-41) and replacement of the

differentials by difference terms yields

$$2[\text{Re}(v_1)\Delta\text{Re}(v_1) + \text{Im}(v_1)\Delta\text{Im}(v_1)] = 2\delta e^{2\delta\text{Re}(\rho_1)}\Delta\text{Re}(\rho_1), \quad (3-43)$$

from which the error bound associated with the real part of the pole is obtained as

$$\Delta\text{Re}(\rho_1) = \frac{\text{Re}(v_1)\Delta\text{Re}(v_1) + \text{Im}(v_1)\Delta\text{Im}(v_1)}{\delta e^{2\delta\text{Re}(\rho_1)}}. \quad (3-44)$$

After substitution of Equation (3-41) into (3-44) the final result for $\Delta\text{Re}(\rho_1)$ is obtained:

$$\Delta\text{Re}(\rho_1) = \frac{\text{Re}(v_1)\Delta\text{Re}(v_1) + \text{Im}(v_1)\Delta\text{Im}(v_1)}{\delta[\text{Re}(v_1)^2 + \text{Im}(v_1)^2]}. \quad (3-45)$$

Since the signs of $\Delta\text{Re}(v_1)$ and $\Delta\text{Im}(v_1)$ are arbitrary, the extreme error of $\text{Re}(\rho_1)$ can be written as

$$\max[\Delta\text{Re}(\rho_1)] = \frac{|\text{Re}(v_1)\Delta\text{Re}(v_1)| + |\text{Im}(v_1)\Delta\text{Im}(v_1)|}{\delta[\text{Re}(v_1)^2 + \text{Im}(v_1)^2]} \quad (3-46)$$

When the poles are real Equation (3-44) becomes

$$\Delta\text{Re}(\rho_1) = \Delta\rho_1 = \frac{\Delta v_1}{\delta e^{2\delta\rho_1}} \quad (3-47)$$

which is of considerable significance as will be seen later.

To evaluate the complex part of the pole, Equation (3-40) can be rewritten as

$$\text{Re}(v_1) + j\text{Im}(v_1) = e^{\delta\text{Re}(\rho_1)} [\cos\{\delta\text{Im}(\rho_1)\} + j \sin\{\delta\text{Im}(\rho_1)\}] \quad (3-48)$$

from which

$$\operatorname{Re}(v_1) = e^{\delta \operatorname{Re}(\rho_1)} \cos [\delta \operatorname{Im}(\rho_1)] . \quad (3-49)$$

Thus $\operatorname{Im}(\rho_1)$ is given by

$$\operatorname{Im}(\rho_1) = \frac{1}{\delta} \cos^{-1} \frac{\operatorname{Re}(v_1)}{\sqrt{\operatorname{Re}(v_1)^2 + \operatorname{Im}(v_1)^2}} \quad (3-50)$$

The error associated with the complex part of the pole is obtained by differentiation of Equation (3-49), and after some algebra is given by

$$\Delta \operatorname{Im}(\rho_1) = \frac{\Delta \operatorname{Re}(v_1) + \delta \operatorname{Re}(v_1) \Delta \operatorname{Re}(\rho_1)}{\delta \operatorname{Im}(v_1)} . \quad (3-51)$$

The extreme error of $\operatorname{Im}(\rho_1)$ is obtained by adjustment of the signs and can be expressed as

$$\max [\Delta \operatorname{Im}(\rho_1)] = \frac{|\Delta \operatorname{Re}(v_1)| + |\delta \operatorname{Re}(v_1) \Delta \operatorname{Re}(\rho_1)|}{\delta \operatorname{Im}(v_1)} \quad (3-52)$$

Discussion

The mathematical expressions that describe the propagation of error through the numerical computations are rather complex, and digital computing facilities are necessary for their utilization.

Based on the error propagation analysis several general observations concerning the identification technique can be made. Equation (3-33) provides a good starting point for this discussion.

$$|\Delta v_1| \leq \underline{C}_1 N(D_h) = \frac{\|\underline{v}_1\| \cdot \|\underline{v}_1\|}{|(\underline{v}_1, \underline{v}_1)|} N(D_h) . \quad (3-33)$$

There are two separate factors that contribute to

the error of an eigenvalue v_1 : (1) the distortion coefficient C_1 which is a function only of the nominal values of the data but not of the initial error estimates, and (2) the norm of the error matrix D_h which is a function of both the data and the error estimates.

The distortion coefficient constitutes a sensitivity measure for the identification of the eigenvalue v_1 . It may be considered a weighting factor for $N(D_h)$ with respect to v_1 and therefore it provides a quantitative measure of how well a particular pole is represented in the response data. The norm $N(D_h)$ provides a general measure for the uncertainty associated with the identification but does not act as a discriminating factor between the individual eigenvalues. If, for example, the response vectors are weakly independent* and consequently the matrix $\Phi_{n-1}^T(t_{n-1} - \delta)$ is poorly conditioned, (i.e. the elements of the inverse are very sensitive to errors in the original matrix), the norm $N(D_h)$ will be large, and a large error will be indicated for all the eigenvalues. Thus the error estimate establishes a criterion for acceptance or rejection of a certain set of response functions. It will be seen later that this characteristic is of particular importance since it indicates the possibility for utilizing the error criterion in the determination of the system's order.

*Strictly speaking the response vectors can only be either linearly dependent or independent. What is meant here by weakly independent is that due to noise effects the independence of the responses cannot be clearly established.

It was previously stated that the parameters t_{n-1} and δ are free to be selected arbitrarily. It is quite evident from the error propagation analysis that the selection of both δ and t_{n-1} effects the error of the final results. Although t_{n-1} does not enter the calculations explicitly at any stage, it is associated with the fundamental matrices $\Phi_{n-1}(t_{n-1})$ and $\Phi_{n-1}(t_{n-1} - \delta)$: it determines what section of the response record is to be used for carrying out the identification. Since in general the noise level is independent of the signal level, it is quite obvious that moving further out on the response record results in increasing the relative uncertainty of the data and therefore in obtaining a less reliable identification. Consequently, since the best possible identification is desired, none of the initial information of the response functions should be discarded. It therefore follows from Equation (2-80) that t_{n-1} must be selected to be equal to δ . The above conclusion was also reached experimentally as will be shown later.

The parameter δ determines the width of the record segment over which the identification is to be performed. Its effect on the quality of the results is quite significant. Consider again Equation (3-47) which describes the error associated with real poles:

$$\Delta \rho_1 = \frac{\Delta v_1}{\delta e^{\rho_1}} \quad (3-47)$$

The relative error of ρ_1 is given by

$$\frac{\Delta \rho_1}{\rho_1} = \frac{\Delta v_1}{\delta \rho_1 e} \quad (3-53)$$

If Δv_1 is considered to be independent* of δ , Equation (3-53) can be written as

$$\frac{\Delta \rho_1}{\rho_1} \simeq \frac{K}{\delta \rho_1 e} \quad (3-54)$$

where K is a proportionality constant. A plot of Equation (3-54) as a function of $\delta \rho_1$ for $K=1$ is presented in Figure 4. It can be seen from this plot that the relative error associated with the pole is a strong function of δ . It is minimized when $-\delta \rho_1 = 1$ or when $\delta = -\frac{1}{\rho_1} = T_1$, where T_1 is the 1th time constant of the system. It can thus be stated that the error associated with a particular time constant is minimized when δ is equal to this time constant. This last observation is, of course, true for all the time constants of the system, and the error for each time constant will be minimized at a particular value of δ . Thus there does not, in general, exist one unique δ which should be selected for the identification, but rather, several values of δ which minimize the individual errors associated with the various time constants. Figure 5 shows the relative error, Equation (3-53), associated with different values of ρ_1 as a function of δ .

*This independence is generally not strictly true but it can be expected that Δv_1 will be a weak function of δ , particularly so when the relative error in the original response data is small.

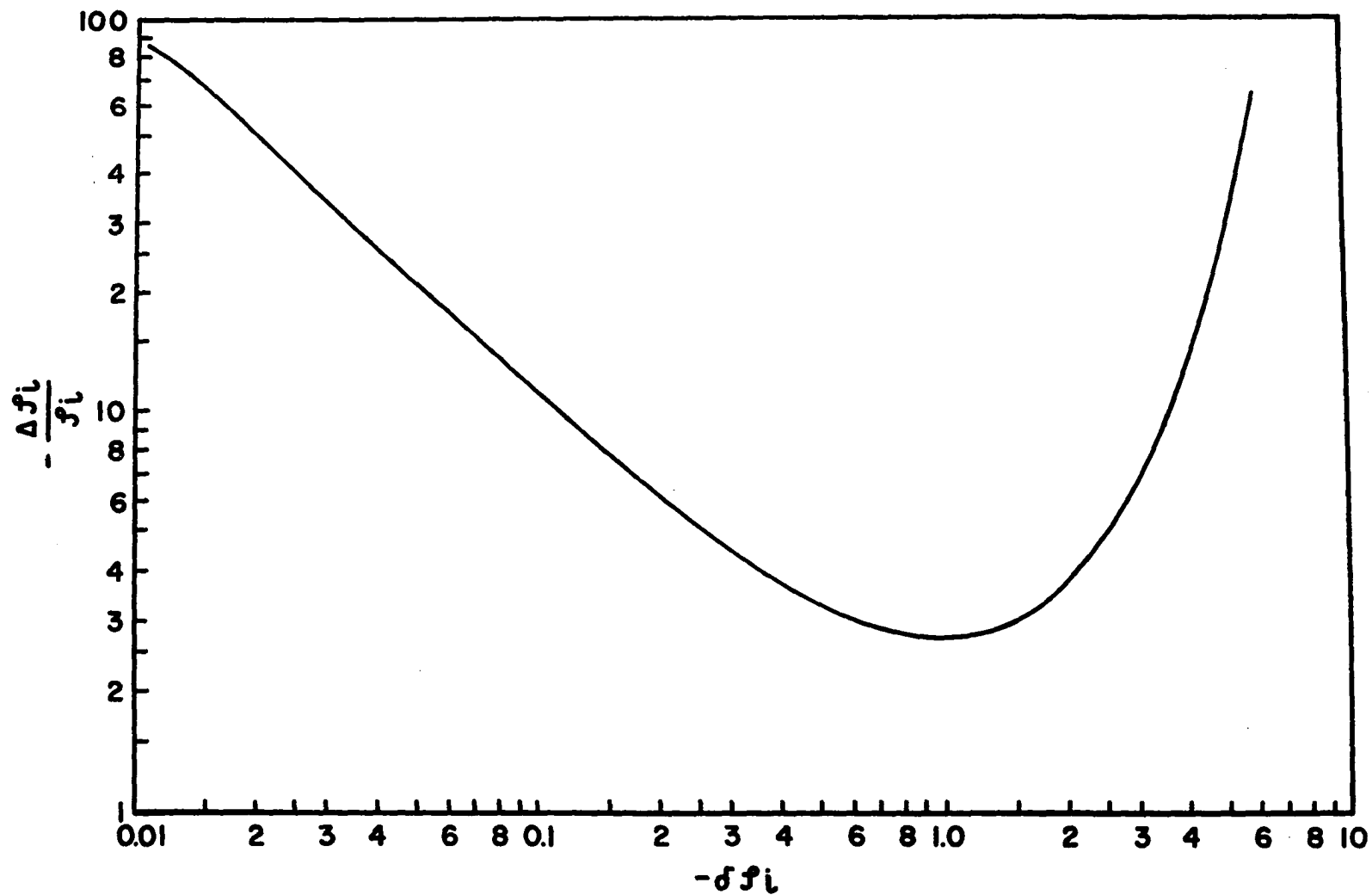


FIGURE 4. ERROR PROPAGATION FOR REAL POLES AS A FUNCTION OF $(-\delta f_i)$ —NORMALIZED PLOT OF EQUATION (3-54)

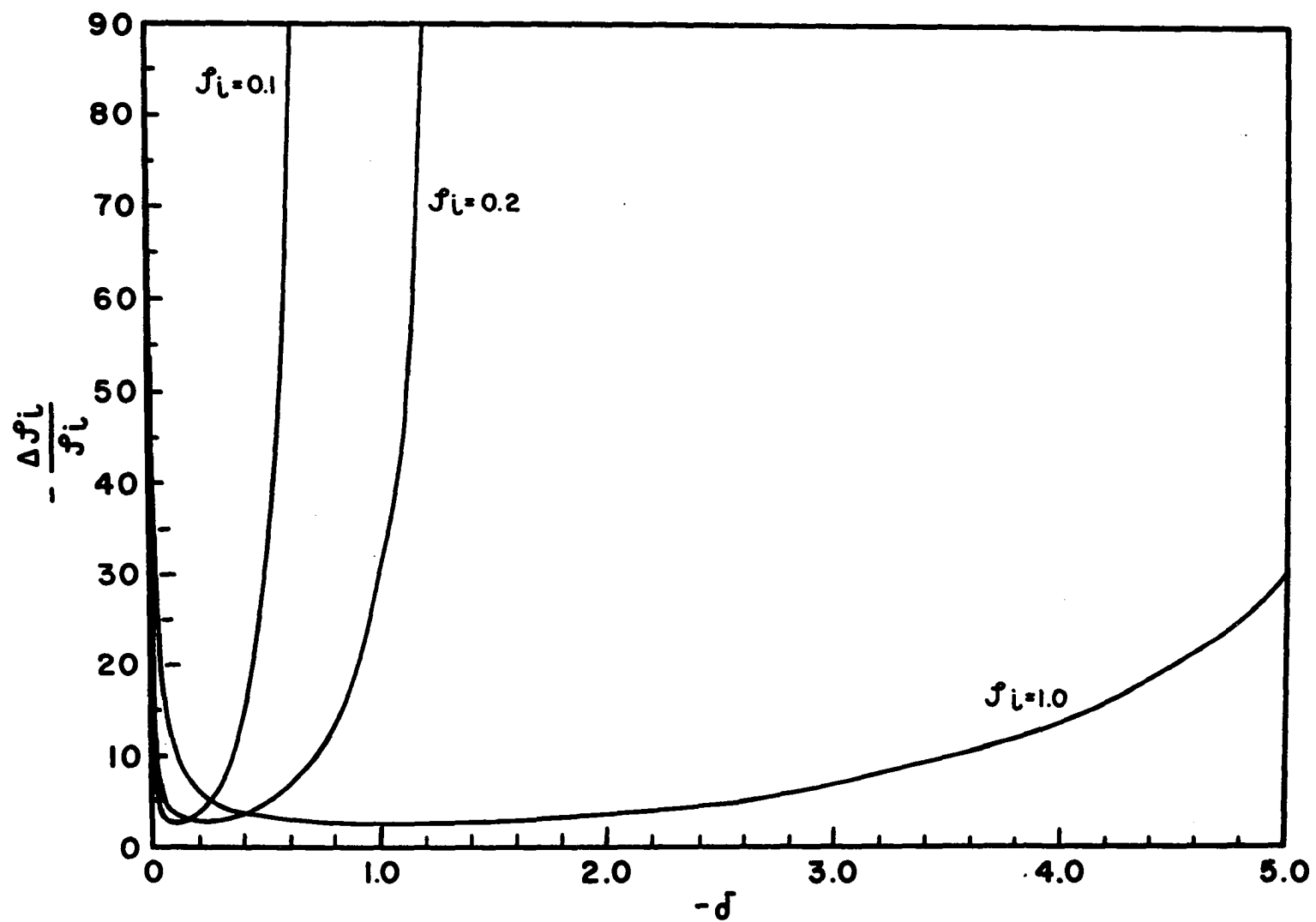


FIGURE 5. ERROR PROPAGATION FOR REAL POLES AS FUNCTION OF δ

The above observations are not strictly correct since Δv_1 is not a constant, and it is not entirely independent of δ . Therefore the error may not be minimized exactly at the time constants and slight deviations from this value may be expected. Of greater importance is the dependence of both Δv_1 and the identified value for ρ_1 on the actual response records. If the responses are noisy, i.e. relatively large error is associated with them, then both Δv_1 and ρ_1 will fluctuate as δ is changed. As a result the functions represented in Figures 4 and 5 will be noisy. The general pattern of these plots will, however, be preserved.

When the poles are complex, the error associated with the real part will behave similarly to the error in the case of real poles. The relative error associated with the imaginary part can be obtained by proper substitution of Equation (3-40) into (3-51) to yield

$$\frac{\Delta \text{Im}(\rho_1)}{\text{Im}(\rho_1)} = \frac{|\Delta \text{Re}(v_1)| + |\delta \text{Re}(v_1) \Delta \text{Re}(\rho_1)|}{\text{Im}(\rho_1) e^{\delta \text{Re}(\rho_1)} \sin [\delta \text{Im}(\rho_1)]} \quad (3-55)$$

It is quite obvious that due to the sinusoidal term in the denominator of Equation (3-55), the relative error associated with the imaginary part will possess a periodic oscillation and will approach infinity at periods when $\delta \text{Im}(\rho_1) = n\pi$. Equation (3-55) has not been plotted here, but it will be discussed further in Chapter VI.

The foregoing discussion strongly indicates as to how the identification procedure should be carried out and

formalized in practice. This subject is discussed in detail in Chapter V.

CHAPTER IV

ERROR ESTIMATION FOR INTEGRALS OF RESPONSE CURVES

The error propagation analysis presented previously requires a knowledge of the uncertainty associated with the initial response data. As the basic information required for the proposed identification procedure is the response functions and their integrals, it is necessary to obtain estimates for the propagation of error when the response records are integrated.

A rigorous statistical analysis of this problem presents certain difficulties. These difficulties arise from the fact that finite records of random functions are available while the analysis is strictly applicable only to infinite records. Certain simplifying assumptions have been made in order to apply the theory to the analysis of the available records. While a conventional approach utilizing the direct evaluation of the correlation and spectral density functions of the stationary random process that is superimposed on the time response is theoretically possible, it would not be justified in this case. The mathematical analysis is too complex and costly compared with the amount of information gained.

The problems associated with the rigorous evaluation of the error propagation are discussed and certain significant

points of this identification method are elucidated. A simplified method is proposed for obtaining an approximate estimate of the error propagation through the integration process. This method is based on zero crossing distribution of Gaussian random functions.

Errors in the response curves can occur for many different reasons. Generally the uncertainty is present in the form of random noise superposed on the true function such that "on the average" its effect on the function is zero. Sometimes a consistent error is present and may be caused by offsets or time lags in measuring devices. When the first type of uncertainty is present, the propagation of the error into the integrals of the responses can be evaluated. The second type of error presents more serious problems since accurate knowledge of its existence is usually unavailable and consequently erroneous identification will be a possible result. If a measure of the consistent error is given it can be easily handled by subtraction of its effect prior to the analysis. Another source of error which usually has to be taken into account is the presence of reading error introduced during interpretation of the response records. This error is normally randomly distributed and it is indirectly included in the present analysis. Since its effect is small compared to uncertainties from other sources, it will not require separate consideration.

In the present work only the error propagation due to the presence of random noise is discussed. Since the propagation of error is a function only of the statistical properties of the noise and is independent of the nature of

the response function with which it is associated, it is sufficient to study the behavior of the random function alone without specific regard to the response.

In theory, the noise can assume any statistical distribution. However it will in general be Gaussian or will have a distribution that can be closely approximated by a Gaussian process. The following assumptions will be made:

- (1) The noise is a stationary* and ergodic** random process.
- (2) The noise is described by a Gaussian distribution.
- (3) The expectation (or the mean value) is zero, i.e.:

$$E[\epsilon(t)] = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^T \epsilon(\tau) d\tau = 0 \quad (4-1)$$

Two different cases will be considered for the necessary analysis:

- (1) Continuous analysis, in which case the responses are continuously recorded and simultaneously integrated. Thus the integration is considered to be exact.
- (2) Discrete analysis, in which case the responses are sampled at constant time intervals. The integration is then performed numerically utilizing sampled data.

*A random process is defined as stationary if every translation in time preserves the probability and statistical properties.

**A stationary random process possesses the ergodic property if its ensemble averages and statistical properties are the same as its time averages and statistical properties.

Continuous Analysis

Suppose the noise is assumed to be a Gaussian random process with zero mean and denoted by $x(t)$. It is desired to evaluate the statistical properties of the integral of $x(t)$ given by

$$y(t) = \int_0^t x(\tau) d\tau \quad (4-2)$$

from the statistical properties of $x(t)$. $y(t)$ is of course a random process, and it is Gaussian since $x(t)$ is Gaussian [27]. Since $x(t)$ is not a deterministic function of time particular values cannot be assigned to it for predetermined values of time. The limits of integration in (4-2), therefore, have only a statistical correspondence with $x(t)$. In order to evaluate the statistical properties of $y(t)$ it is therefore necessary to perform averaging processes over all possible lower limits of integration. With this requirement in mind Equation (4-2) is rewritten in the form

$$y(t, \gamma) = \int_{\gamma}^{t+\gamma} x(\tau) d\tau, \quad (4-3)$$

where t is considered a parameter.

The statistical properties of $y(t)$ of interest in the present study are the mean value and the mean square value. The mean value of $y(t)$ is given by

$$E[y(t)] = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^T d\gamma \int_{\gamma}^{t+\gamma} x(\tau) d\tau. \quad (4-4)$$

By changing the order of integration [27], (4-4) becomes

$$\begin{aligned}
E[y(t)] &= \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^T d\gamma \int_{\gamma}^{\gamma+t} x(\tau) d\tau \\
&= \lim_{T \rightarrow \infty} \frac{1}{2T} \left[\int_{-T}^{-T+t} x(\tau) d\tau \int_{-T}^{\tau} d\gamma + \int_{-T+t}^T x(\tau) d\tau \int_{\tau-t}^{\tau} d\gamma \right. \\
&\quad \left. + \int_T^{\gamma+t} x(\tau) d\tau \int_{\tau-t}^T d\gamma \right] \\
&= \lim_{T \rightarrow \infty} \frac{1}{2T} \left[\int_{-T}^{-T+t} (\tau + T) x(\tau) d\tau + t \int_{-T+t}^T x(\tau) d\tau \right. \\
&\quad \left. + \int_T^{\gamma+t} (T - \tau + t) x(\tau) d\tau \right] = 0. \quad (4-5)
\end{aligned}$$

The last step in the computation of Equation (4-5) follows from the fact that $E[x(t)] = 0$ and can be directly verified.

The fact that the expected value for $y(t)$ is zero does not mean that a particular record of $y(t)$ is randomly distributed around zero. The correct interpretation of Equation (4-5) is that statistically $y(t)$ has an equal chance of being above or below its mean value when evaluated over all possible records; a specific function $y(t)$ may assume any arbitrary value.

In order to evaluate the mean square value of $y(t)$ it is first desirable to determine its autocorrelation function. Let

$$y(t) = \int_{\gamma}^{\gamma+t} x(\tau_1) d\tau_1 = - \int_0^t x(\gamma - \tau_1) d\tau_1,$$

and

$$y(t + \tau) = \int_{\gamma}^{\gamma+t+\tau} x(\tau_2) d\tau_2 = - \int_0^{t+\tau} x(\gamma - \tau_2) d\tau_2.$$

Then,

$$\begin{aligned} y(t) \cdot y(t + \tau) &= \int_0^t x(\gamma - \tau_1) d\tau_1 \int_0^{t+\tau} x(\gamma - \tau_2) d\tau_2 \\ &= \int_0^t d\tau_1 \int_0^{t+\tau} x(\gamma - \tau_1) x(\gamma - \tau_2) d\tau_2. \quad (4-6) \end{aligned}$$

The autocorrelation function $R_{yy}(\tau)$ of $y(t)$ is given by the expectation (mean value) of $y(t) \cdot y(t + \tau)$ and is evaluated by averaging (4-6) over all initial limits of integration γ :

$$\begin{aligned} R_{yy}(\tau) &= E[y(t) \cdot y(t + \tau)] \\ &= \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^T d\gamma \left[\int_0^t d\tau_1 \int_0^{t+\tau} x(\gamma - \tau_1) x(\gamma - \tau_2) d\tau_2 \right] \\ &= \int_0^t d\tau_1 \int_0^{t+\tau} d\tau_2 \left[\lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^T x(\gamma - \tau_1) x(\gamma - \tau_2) d\gamma \right] \\ &= \int_0^t d\tau_1 \int_0^{t+\tau} R_{xx}(\tau_2 - \tau_1) d\tau_2, \quad (4-7) \end{aligned}$$

where $R_{xx}(\tau)$ is the autocorrelation function of $x(t)$. The mean square value of $y(t)$ is given by $R_{yy}(0)$ as

$$E[y(t)^2] = R_{yy}(0) = \int_0^t d\tau_1 \int_0^t R_{xx}(\tau_2 - \tau_1) d\tau_2. \quad (4-8)$$

The double integral in Equation (4-8) can be rewritten by changing the variable of integration to yield

$$E[y(t)^2] = \int_0^t d\tau_1 \int_0^t R_{xx}(\tau_2 - \tau_1) d\tau_2 = \int_0^t d\tau_1 \int_{-\tau_1}^{t-\tau_1} R_{xx}(\tau_2) d\tau_2 . \quad (4-9)$$

The repeated integral in Equation (4-9) can, as before, be evaluated by changing the order of integration to give

$$\begin{aligned} E[y(t)^2] &= \int_0^t R_{xx}(\tau_2) d\tau_2 \int_0^{t-\tau_2} d\tau_1 + \int_{-t}^0 R_{xx}(\tau_2) d\tau_2 \int_{-\tau_2}^t d\tau_1 \\ &= 2 \int_0^t (t - \tau_2) R_{xx}(\tau_2) d\tau_2 . \end{aligned} \quad (4-10)$$

The last step in Equation (4-10) follows from the fact that $R_{xx}(\tau)$ is an even function of τ [27]. Equation (4-10) can then be rewritten as

$$E[y(t)^2] = 2t \int_0^t R_{xx}(\tau) d\tau - 2 \int_0^t \tau R_{xx}(\tau) d\tau . \quad (4-11)$$

As t becomes large the second term on the right becomes a constant since $R_{xx}(\tau)$ approaches zero for large values of τ .* The mean square value then approximately becomes (for large τ)

$$E[y(t)^2] \cong 2t \int_0^t R_{xx}(\tau) d\tau . \quad (4-12)$$

Thus the mean square value for $y(t)$ approaches infinity as

* $x(t)$ is a stationary Gaussian random process.

t approaches infinity. The random process $y(t)$ is not a stationary random process since its mean square value is a function of t . This result may be somewhat surprising since it does not comply with normal intuition. It can, however, be reasoned qualitatively and the physical interpretation may uncover some interesting points. The Gaussian random process $x(t)$ is defined only by its statistical parameters, i.e., the mean value and correlation function. The mean value provides information about the average amplitude but implies nothing as to possible amplitude fluctuations. Similarly the correlation function determines the mean amplitude fluctuation (in an absolute value sense) and the average probability for sign changes. It does not provide any deterministic information of these values. The function $x(t)$ may therefore possess any arbitrarily large values and may change signs at arbitrarily low frequencies without violating the statistical requirements.* As a result, $x(t)$ may accumulate finite or infinite deviation from its mean value over any finite interval. When the integration is begun at a finite instant of time γ , $y(t)$ may possess an initial uncorrectable error which increases with t . The function $y(t)$ will generally drift away from zero at an arbitrary rate. Equation (4-12) determines the average drifting rate. The fact that the mean value of $y(t)$ is zero means that $y(t)$ possesses equal probabilities to drift above or below zero.

*An exception to this situation will be discussed later.

The mean square value of $y(t)$ will be bounded and the process be stationary if and only if

$$\int_0^{\infty} R_{xx}(\tau) d\tau = 0. \quad (4-13)$$

This condition coincides with the requirement that the power spectral density function $G_{xx}(\omega)$ of $x(t)^*$ possesses no component of zero frequency (DC component) since

$$G_{xx}(0) = \frac{2}{\pi} \int_0^{\infty} R_{xx}(\tau) d\tau. \quad (4-14)$$

This result is quite interesting and deserves some further discussion. In order for $x(t)$ to possess an integral which is a stationary function, it must have no power associated with the zero frequency component. Loosely, this condition means that if $x(t)$ is visualized as a Fourier series with random coefficients and frequencies, the frequency ω_n must never assume the value zero. Thus $x(t)$ must change signs at finite intervals, which conforms with intuitive notions of random functions.

The fact that integrals of random Gaussian noise functions are generally not stationary presents some difficulties in the theoretical analysis if a "general" noise model is to be assumed. In the following discussion it will be

*The power spectral density function is defined as the Fourier transform of the correlation function [3]

$$G_{xx}(\omega) = \frac{1}{\pi} \int_{-\infty}^{\infty} R_{xx}(\tau) e^{-j\omega\tau} d\tau = \frac{2}{\pi} \int_0^{\infty} R_{xx}(\tau) \cos \omega\tau d\tau$$

and can be interpreted as a function describing the power associated with each value of frequency.

shown that from practical considerations it is justified to assume a "particular" noise model that satisfies the requirements of Equations (4-13) and (4-14).

The principal difficulty associated with the analysis arises from the necessity of evaluating the statistical properties of the random noise from finite records. In practice it is necessary to assume that

$$\frac{1}{t_0} \int_0^{t_0} x(\tau) d\tau = 0, \quad (4-15)$$

where t_0 is an arbitrary interval of time over which the mean value is evaluated. Equation (4-15) obviously violates, in a strict sense, the definition of the mean value of Equation (4-5) and the interpretation associated with it. If t is an arbitrary interval of time smaller than t_0 , it follows that

$$\max_t \left[\int_0^t x(\tau) d\tau \right] = |K|, \quad (4-16)$$

where K is some finite number. Practically this limit means that, regardless of record length, the maximum accumulative error is bounded. Thus the conclusion that the integral of a random function is not stationary is meaningless in the present analysis. The preceding argument is very important and deserves some further clarification. In reality, as in theory, it is possible that the integral of $x(t)$ will experience a drift from zero. The point is that one has no specific knowledge of the existence of this drift. To illustrate, consider the typical noisy response curve shown

in Figure 6. The dotted smooth curve represents the true, noise free, response function. This curve cannot be predicted on the basis of the given experimental data since it does not constitute the "best" fit. The requirement that the mean value over a finite record length be zero is in itself insufficient to describe the function since there is an infinite number of curves satisfying this condition. A criterion commonly used for curve fitting is an orthogonal polynomial fit based on minimization of the mean square error. The solid smooth curve represents such a fit. By forcing a "best" fit on the response curve, one simultaneously forces a noise model on the random function. In an intuitive sense the above arguments can be associated with the notion that stationary noise fluctuates at a certain average frequency around its mean value. Likewise one tends to associate some finite, "reasonable", maximum, zero-crossing interval with the intuitive concept of a noise function. Of course the term reasonable leaves much to individual judgment.

From the preceding discussion it is clear that, due to practical limitations, one inevitably restricts the permissive noise models to a particular and very limited class. Since the noise is evaluated from finite records, the integral of the random function is artificially bounded, (Equation (4-16)), and over the tested region it is therefore forced to be stationary. It seems therefore justified to assume that also over the infinite record the random function satisfies the requirement that $G_{xx}(0) = 0$.

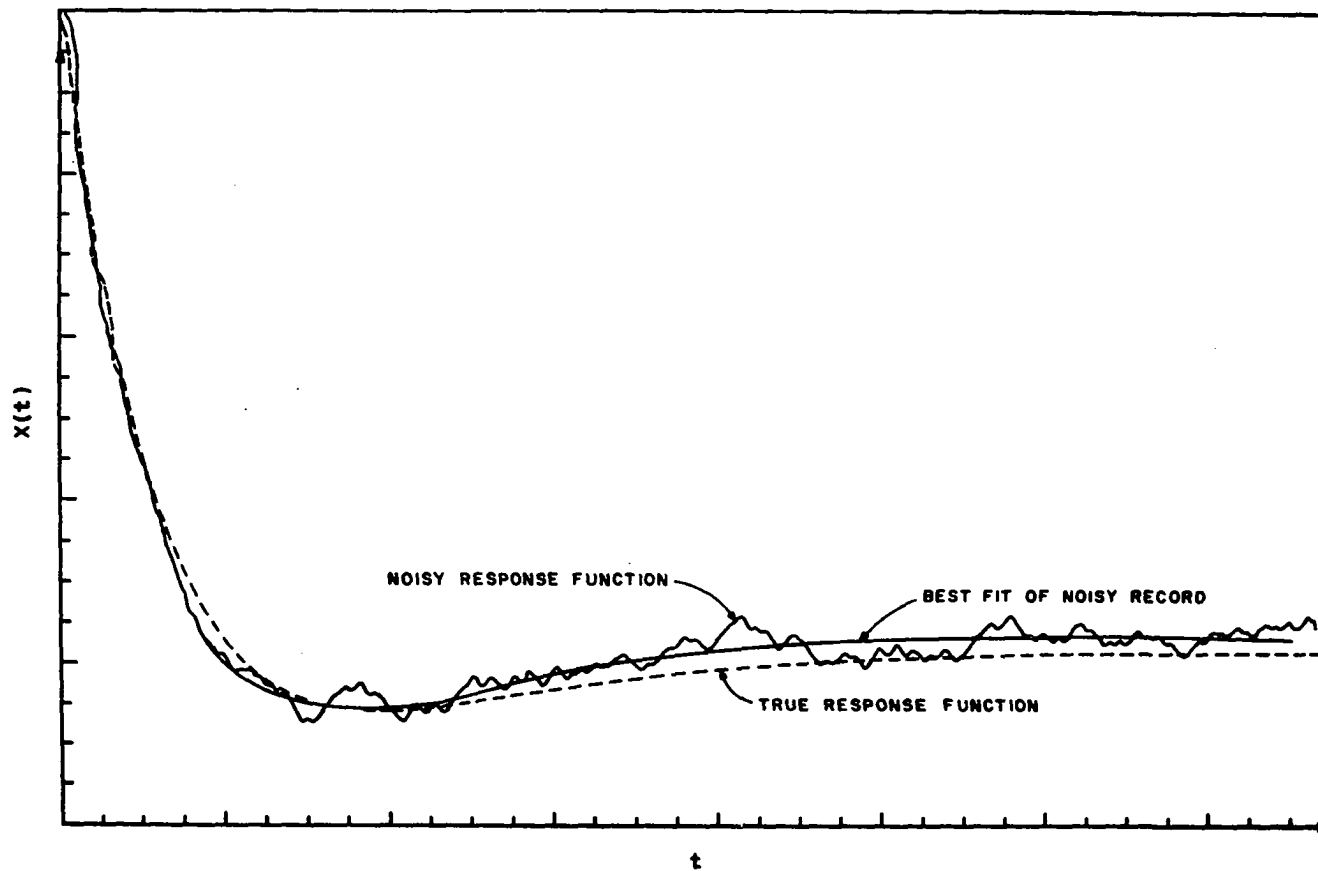


FIGURE 6. TYPICAL NOISY RESPONSE FUNCTION

Zero-Crossing Problem

In the present work the evaluation of the statistical properties of the noise is carried out only to obtain an estimate for the propagation of error into the integrals of the response curves. This calculation is obviously only of secondary importance for the identification procedure, and it is therefore not justified to invest excessive computational effort in obtaining a noise model. Evaluation of the correlation function and spectral density directly would require a large amount of computation and is therefore impractical. A method for evaluating approximative noise models that proved to be quite successful in estimating the propagation of error into the integrals, and also to be simple in application, is based on the relationship between zero-crossing distributions of random time functions and the respective correlation functions.

The problem of evaluating the axis-crossing distribution of random time functions has been carried out by several investigators to varying degrees of completion [29,30,32,33, 42,43]. To obtain the desired results, an approach proposed by McFadden [32,33] will be used.

McFadden considers an "infinitely clipped" noise function, that is, a random process cut such that only the zero-crossing intervals remain invariant but all other information is lost. After clipping, the function is defined as follows: Let $\xi(t)$ be clipped to form $x(t)$ where $x(t)$ is both stationary and ergodic. The function $x(t)$ may assume

only the values ± 1 and either value has equal probability. The autocorrelation of $x(t)$ is denoted by $r(\tau)$. For a Gaussian random process $\xi(t)$, with normalized autocorrelation $\rho(\tau)$, $x(t)$ (after infinite clipping) is defined as

$$\begin{aligned} x(t) &= 1 \text{ for } \xi(t) \geq 0 \\ x(t) &= -1 \text{ for } \xi(t) < 0 \end{aligned} \quad (4-17)$$

Then the autocorrelation function of $x(t)$ is related to $\rho(t)$ [28] by

$$r(\tau) = \frac{2}{\pi} \arcsin \rho(\tau) , \quad (4-18)$$

where $\rho(\tau)$ is the autocorrelation function of $\xi(t)$.

The probability density $P(\tau)$ is defined for axis crossing intervals in the following manner: Given a zero at time t , $P(\tau)d\tau$ is the conditional probability that the next zero lies between $t + \tau$ and $t + \tau + d\tau$. Then

$$\int_0^{\infty} P(\tau)d\tau = 1 \quad (4-19)$$

and

$$\gamma = E[\tau] = \int_0^{\infty} \tau P(\tau)d\tau , \quad (4-20)$$

where $\gamma = E[\tau]$ is the mean axis-crossing interval. The expected number of zero-crossings per unit time is denoted by β and defined by $\beta = 1/\gamma$.

If $x_1 = x(t_1)$ and $x_2 = x(t_2)$ then the probability that x_1 and x_2 have different signs can be evaluated as follows. Let the probability that x_1 and x_2 are both positive be P_{++} and similarly define P_{--} , P_{-+} and P_{+-} . Then

$$P_{++} + P_{+-} + P_{-+} + P_{--} = 1$$

$$E[x_1] = 1 \cdot P_{++} + (-1) \cdot P_{--} + 1 \cdot P_{+-} + (-1) \cdot P_{-+}$$

$$E[x_2] = 1 \cdot P_{++} + (-1) \cdot P_{--} + (-1) \cdot P_{+-} + 1 \cdot P_{-+}$$

$$E[x_1 \cdot x_2] = 1 \cdot P_{++} + 1 \cdot P_{--} + (-1) \cdot P_{+-} + (-1) \cdot P_{-+}$$

where $E[x_1]$ and $E[x_2]$ are the expected values of x_1 and x_2 respectively, and $E[x_1 x_2]$ is the expected value of the product of x_1 and x_2 . Clearly $E[x_1] = E[x_2] = 0$ since both have equal likelihood of being negative and positive, and $E[x_1 x_2] = r_{12}$. The solution of the four equations yields

$$P_{+-} = P_{-+} = \frac{1}{4} (1 - r_{12}) , \quad (4-21)$$

and the probability that x_1 and x_2 are of different signs is

$$P_{+-} + P_{-+} = \frac{1}{2} (1 - r_{12}) . \quad (4-22)$$

Let $t_1 = t$ and $t_2 = t + \Delta t$ where Δt is small enough that the probability of finding more than one zero between t and $t + \Delta t$ is negligible. Then the probability that $x(t_1)$ and $x(t_2)$ are of opposite signs becomes the expected number of zeros per unit time multiplied by Δt .

Thus

$$\beta \Delta t = \frac{1}{2} [1 - r(\Delta t)] . \quad (4-23)$$

For sufficiently small Δt , $r(\Delta t)$ can be expressed by the truncated Taylor series $r(\Delta t) = 1 + r'(0)\Delta t$, and finally

$$\beta \Delta t = \frac{1}{2} (1 - r'(0)\Delta t - 1) = -\frac{1}{2} r'(0)\Delta t ,$$

or

$$\beta = -\frac{1}{2} r'(0) . \quad (4-24)$$

For clipped Gaussian noise β can be expressed as a function of the correlation function $\rho(\tau)$ by

$$\begin{aligned} \beta &= -\frac{1}{2} r'(0) = \lim_{\tau \rightarrow 0} \left\{ -\frac{1}{2} \frac{d}{d\tau} \left[\frac{2}{\pi} \arcsin \rho(\tau) \right] \right\} \\ &= \frac{1}{\pi} \lim_{\tau \rightarrow 0} \left[-\frac{\rho'(\tau)}{\sqrt{1 - \rho^2(\tau)}} \right] . \end{aligned} \quad (4-25)$$

If $\rho'(0) = 0$ then by L'Hospital's rule (since $\rho(0) = 1$):

$$\beta = \frac{1}{\pi} \lim_{\tau \rightarrow 0} \frac{\rho''(\tau) \sqrt{1 - \rho^2(\tau)}}{\rho(\tau) \rho'(\tau)} , \quad (4-26)$$

from which it follows that

$$\lim_{\tau \rightarrow 0} \left[\frac{\rho'(\tau)}{\sqrt{1 - \rho^2(\tau)}} \right]^2 = \lim_{\tau \rightarrow 0} [-\rho''(\tau)] .$$

Finally

$$\beta = \frac{1}{\pi} [-\rho''(0)]^{1/2} . \quad (4-27)$$

The above result is exact, it was previously derived from a different approach by Rice [43].

Evaluation of the probability distribution function of zero-crossing intervals $P(\tau)$ is more difficult than the mean interval, and to date only approximate solutions to the problem have been found. Rice [43], who was the first to study the problem, developed a first approximation to the solution for Gaussian random processes. Following his lead several attempts have been made to obtain better estimates

[29,30,33]. Experimental studies have also been performed and results were obtained which agree closely with some of the theoretical formulas [42]. For reasons which will be discussed later, it was decided to use, for the present purpose, an approximation proposed by McFadden [32] which coincides with Rice's solution.

Let $P_n(\tau)$ be the probability of the sum of $n+1$ successive axis-crossing intervals. (Thus $P_0(\tau)$ is the probability density for one zero-crossing interval). Also, let $p(n, \tau)$ be the probability that a given time interval of length τ contains exactly n zeros. Let us now find a general relationship between $P_n(\tau)$ and $p(n, \tau)$.

Consider the following events E_1 ; E_2 ; E_3 ; E_4 and E_5 :

E_1 : There are exactly n points in the interval
($t, t + \tau$)

E_2 : There are exactly n points in the interval
($t, t + \tau + d\tau$)

E_3 : There are exactly n points in ($t, t + \tau$) and
there is one point in the adjacent interval
($t + \tau, t + \tau + d\tau$).

The probability of more than one point in an infinitesimal interval is neglected.

E_4 : There are not more than n points in ($t, t + \tau$)
and there is one in ($t + \tau, t + \tau + d\tau$).

E_5 : There are not more than $(n-1)$ points in ($t, t + \tau$)
and there is one in ($t + \tau, t + \tau + d\tau$).

E_2 can be divided into two alternatives:

$E_2^{(1)}$: There are exactly n points in $(t, t + \tau)$ and none in $(t + \tau, t + \tau + d\tau)$

$E_2^{(2)}$: There are exactly $n-1$ points in $(t, t + \tau)$ and there is one in $(t + \tau, t + \tau + d\tau)$.

Then

$$\Pr\{E_1\} = \Pr\{E_2^{(1)}\} + \Pr\{E_3\}$$

$$\Pr\{E_2\} = \Pr\{E_2^{(1)}\} + \Pr\{E_2^{(2)}\}$$

Subtraction of the above two equations yields

$$\Pr\{E_1\} - \Pr\{E_2\} = \Pr\{E_3\} - \Pr\{E_2^{(2)}\}.$$

Since

$$\Pr\{E_1\} = p(n, \tau),$$

and

$$\Pr\{E_2\} = p(n, \tau + d\tau)$$

the difference (neglecting higher order terms) is

$$\Pr\{E_2\} - \Pr\{E_1\} = p'(n, \tau)d\tau.$$

Now

$$\Pr\{E_3\} = \Pr\{E_4\} - \Pr\{E_5\},$$

and

$$\Pr\{E_4\} = \beta d\tau \int_{\tau}^{\infty} P_n(l) dl$$

$$\Pr\{E_5\} = \beta d\tau \int_{\tau}^{\infty} P_{n-1}(l) dl.$$

The above two equations may be interpreted as follows: $\beta d\tau$ is the probability of finding one point in $(t+\tau, t+\tau+d\tau)$.

$\int_{\tau}^{\infty} P_n(l) dl$ is the conditional probability that the $(n+1)^{st}$ point prior to the one in $(t+\tau, t+\tau+d\tau)$ occurs before the time t , given a point in $(t+\tau, t+\tau+d\tau)$. Then $\beta d\tau \int_{\tau}^{\infty} P_n(l) dl$

is the joint probability of finding one point in $(t+\tau, t+\tau+d\tau)$ and no more than n points in $(t, t+\tau)$. Similarly for

$\beta d\tau \int_{\tau}^{\infty} P_{n-1}(t) dt$. Then

$$\Pr\{E_3\} = \beta d\tau \int_{\tau}^{\infty} [P_n(t) - P_{n-1}(t)] dt .$$

By similar reasoning

$$\Pr\{E_2^{(2)}\} = \beta d\tau \int_{\tau}^{\infty} [P_{n-1}(t) - P_{n-2}(t)] dt ,$$

and finally

$$p'(n, \tau) = -\beta \int_{\tau}^{\infty} [P_n(t) - 2P_{n-1}(t) + P_{n-2}(t)] dt . \quad (4-28)$$

Differentiation of Equation (4-28) with respect to τ yields

$$p''(n, \tau) = \beta [P_n(\tau) - 2P_{n-1}(\tau) + P_{n-2}(\tau)] . \quad (4-29)$$

For $n = 0$ and 1 :

$$p''(0, \tau) = \beta P_0(\tau)$$

$$p''(1, \tau) = \beta [P_1(\tau) - 2P_0(\tau)] .$$

Now consider the interval $(t, t+\tau)$. If the number of zeros in the interval is even then, and only then, $x(t)$ and $x(t+\tau)$ have the same sign. In this case the autocorrelation $r(\tau)$ or $E[x(t)x(t+\tau)]$ is given by the infinite series

$$r(\tau) = \sum_{n=0}^{\infty} (-1)^n p(n, \tau) . \quad (4-30)$$

Differentiation of Equation (4-30) twice with respect to τ gives

$$r''(\tau) = \sum_{n=0}^{\infty} (-1)^n p''(n, \tau) . \quad (4-31)$$

Substituting for $p''(n, \tau)$ in Equation (4-31) yields

$$\begin{aligned}
 r''(\tau) &= \beta P_0(\tau) - \beta[P_1(\tau) - 2P_0(\tau)] + \sum_{n=2}^{\infty} (-1)^n p''(n, \tau) \\
 &= 3\beta P_0(\tau) - \beta P_1(\tau) + \sum_{n=2}^{\infty} (-1)^n \beta [P_n(\tau) - 2P_{n-1}(\tau) \\
 &\quad + P_{n-2}(\tau)] \quad (4-32)
 \end{aligned}$$

or

$$\frac{r''(\tau)}{4\beta} = \sum_{n=0}^{\infty} (-1)^n P_n(\tau) \quad (4-33)$$

If $P_n(\tau)$ is neglected for $n \geq 1$ for small τ then

$$P_0(\tau) \cong \frac{r''(\tau)}{4\beta} \quad (4-34)$$

For Gaussian processes Equation (4-18) can be substituted in (4-34) to yield

$$P_0(\tau) \cong \frac{1}{2\pi\beta} \frac{d^2}{d\tau^2} [\arcsin \rho(\tau)] \quad (4-35)$$

Equation (4-35) is the desired result for the probability distribution of zero-crossing intervals. Clearly Equation (4-35) is a good approximation for $P(\tau)$ for small values of τ since in this range the probability of having more than one zero in the interval τ is negligible. For larger values of τ , Equation (4-35) becomes inaccurate as an approximation for $P(\tau)$, and its region of reliability depends in general on the properties of $r(\tau)$. A more detailed discussion of this problem will be presented later.

Evaluation of Noise Model

Because of the fact that Equation (4-35) is only a first approximation to $P(\tau)$ which becomes unreliable for large values of τ , it cannot be used for direct evaluation of $\rho(\tau)$. Furthermore, since the purpose of the present analysis is only to obtain an estimate of error propagation in integration it would be undesirable to introduce excessive numerical calculations.

If a formal noise model could be proposed which would be completely determined by a limited number of parameters, and if, furthermore, these parameters could be determined from zero-crossing information over the range of small values of τ , then Equation (4-35) could become very useful. A general representative noise model cannot be proposed at this point due to lack of knowledge about the system to be represented. With this limitation in mind, only a typical noise model which represents a fairly wide class of random processes can be discussed.

Bendat [3] has shown that a great many natural random phenomena can be characterized by an Exponential-Cosine autocorrelation function

$$R(\tau) = \underline{A}e^{-K|\tau|}\cos c\tau, \quad (4-36)$$

where

$$\underline{A} > 0, K > 0, c > 0.$$

The corresponding normalized correlation function is

$$\rho(\tau) = e^{-K|\tau|}\cos c\tau. \quad (4-37)$$

Thus, in this case $\rho(\tau)$ is completely determined by the two constants K and c . A typical plot of Equation (4-37) is given in Figure 7. The corresponding spectral density functions are shown in Figure 8 for $K^2 > 3c^2$ and for $K^2 < 3c^2$. The disadvantage of this model is the fact that it possesses a finite power component at $\omega = 0$. As has been previously discussed this characteristic would cause drifts in the integral functions which are in practice undetectable. To overcome this problem it was decided to approximate the exponential cosine process by a "band pass" model, i.e., a noise model which has a constant spectral density between two finite frequencies and zero elsewhere. This model of course represents an idealized situation and is physically incorrect. It is, however, a reasonable assumption since the high frequency components are not expected to contribute much to the integration error and the low frequency components are "absorbed" by the process and are undetectable.

The band pass approximation is represented schematically in Figure 9. The solid curve represents the exponential cosine model, and the dashed line is the spectral density of the idealized situation.

The spectral density function for the band pass model is given by

$$G_{xx}(\omega) = \begin{cases} \frac{K}{2} & \text{for } \omega_0 \leq \omega \leq m\omega_0 \\ 0 & \text{elsewhere} \end{cases} \quad (4-38)$$

The correlation function of the noise is then

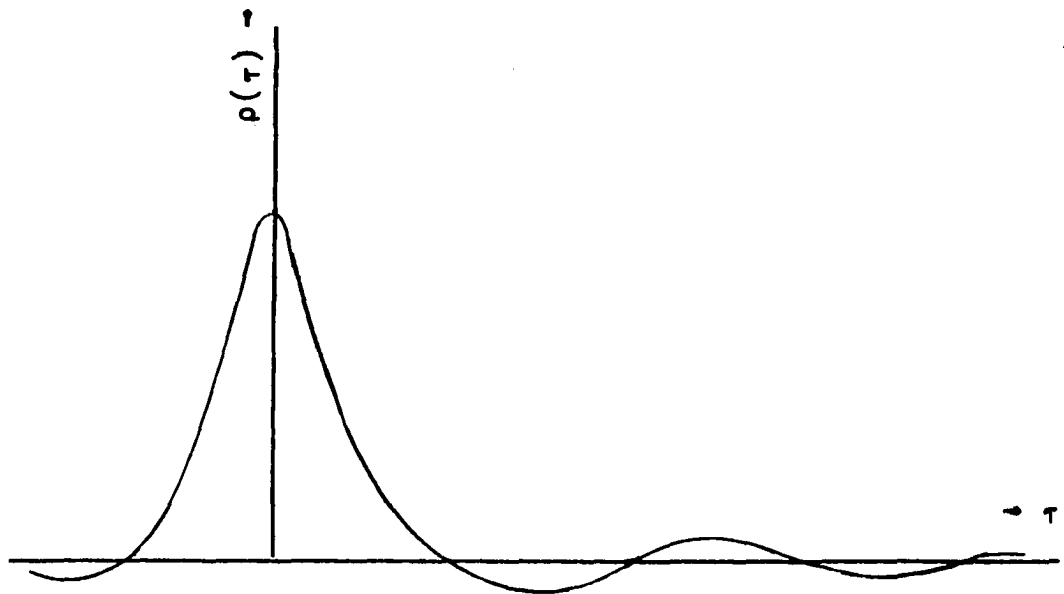


Figure 7. Typical Cosine-Exponential Autocorrelation Function

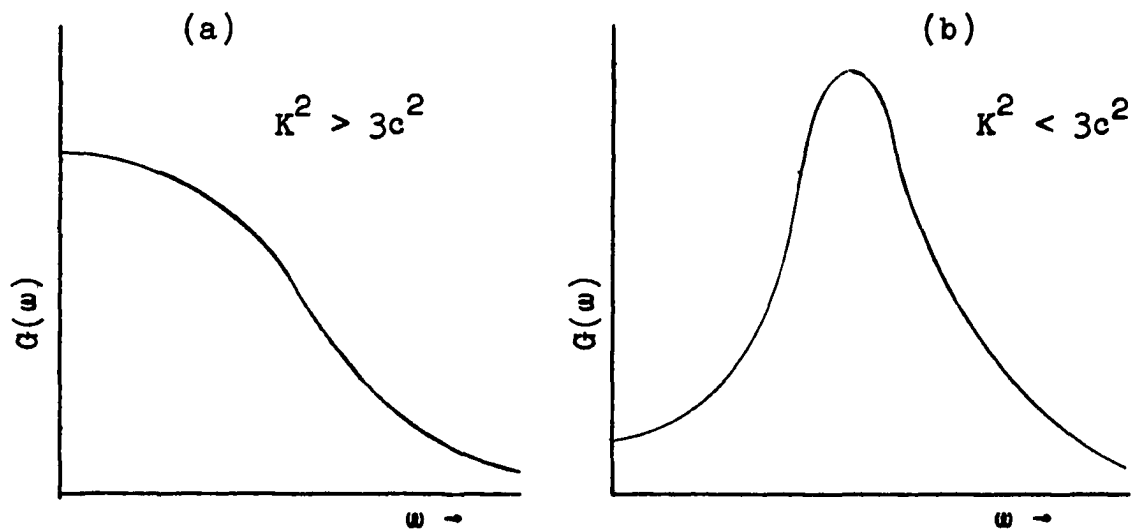


Figure 8. Typical Spectral Density Corresponding to a Cosine-Exponential Correlation Function

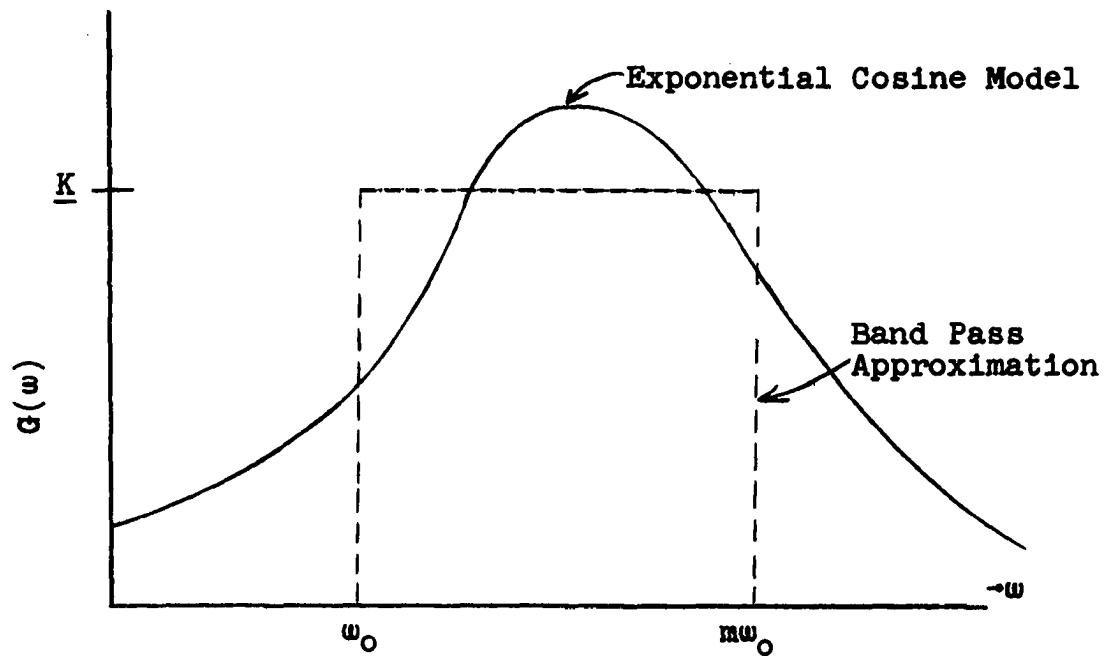


Figure 9. Comparison Between Spectral Densities of Cosine Exponential and Band-Pass-Approximation Noise Models.

$$\begin{aligned}
R_{XX}(\tau) &= \int_0^{\infty} G_{XX}(\omega) \cos \omega \tau d\omega = \underline{K} \int_{\omega_0}^{m\omega_0} \cos \omega \tau d\omega \\
&= \frac{\underline{K}}{\tau} [\sin m\omega_0 \tau - \sin \omega_0 \tau],
\end{aligned} \tag{4-39}$$

and the mean square value is

$$R_{XX}(0) = \lim_{\tau \rightarrow 0} \frac{\underline{K}}{\tau} [\sin m\omega_0 \tau - \sin \omega_0 \tau] = \underline{K} \omega_0 (m - 1) . \tag{4-40}$$

Thus the normalized correlation function $\rho(\tau)$ is

$$\rho(\tau) = \frac{R_{XX}(\tau)}{R_{XX}(0)} = \frac{\sin m\omega_0 \tau - \sin \omega_0 \tau}{(m - 1)\omega_0 \tau} . \tag{4-41}$$

The correlation function $\rho(\tau)$ is given in terms of two parameters m and ω_0 which completely characterize it. If the mean value for zero-crossing intervals and the variance of zero crossing intervals are expressed in terms of $\rho(\tau)$, then m and ω_0 can be determined.

The mean zero-crossing interval can be evaluated by substitution of Equation (4-41) into (4-27). The derivation is presented in Appendix C and the final result is

$$\gamma = \frac{\pi}{\omega_0} \sqrt{\frac{3}{m^2 + m + 1}} . \tag{4-42}$$

The variance of zero-crossing distribution is given by

$$\sigma^2 = \int_0^{\infty} (\tau - \gamma)^2 P(\tau) d\tau . \tag{4-43}$$

Evaluation of Equation (4-43) presents a certain difficulty: In the previous section an approximation for $P(\tau)$ in terms

of $\rho(\tau)$, Equation (4-35), was obtained which is, however, inaccurate for large values of τ . Better approximations to this problem have been obtained by other investigators [30], but these methods are too complicated and cannot be solved explicitly for $\rho(\tau)$ (except numerically). It turns out that for a band pass spectrum the approximation of Equation (4-35) adequately describes $P(\tau)$ for values of τ up to about 1.5 - 2 times the mean interval γ . Figure 10 represents a comparison between theoretical and experimental results based on McFadden's and Longuet-Higgins' publications. It can be seen that $P(\tau)$ coincides with the experimental results and with Longuet-Higgins' approximation $P_0^{(2)}(\tau)$ for small values of τ . It is important to notice that the actual distribution is almost symmetric with respect to γ . Thus if symmetry of $P(\tau)$ with respect to γ is assumed, Equation (4-43) can be rewritten as

$$\sigma^2 = 2 \int_0^{\gamma} (\tau - \gamma)^2 P(\tau) d\tau \quad (4-44)$$

and Equation (4-35) can be utilized to express $P(\tau)$ in terms of $\rho(\tau)$ since it is a satisfactory approximation over the range $0 - \gamma$. The mathematics of the derivation is presented in Appendix D and the final result which is expressed in terms of (σ/γ) is

$$\left(\frac{\sigma}{\gamma} \right)^2 = \frac{2}{\pi\eta} \int_0^{\eta} \sin^{-1} \left[\frac{\sin m\tau - \sin \tau}{(m-1)\tau} \right] d\tau \quad (4-45)$$

where

$$\eta = \pi \sqrt{\frac{3}{m^2 + m + 1}} .$$

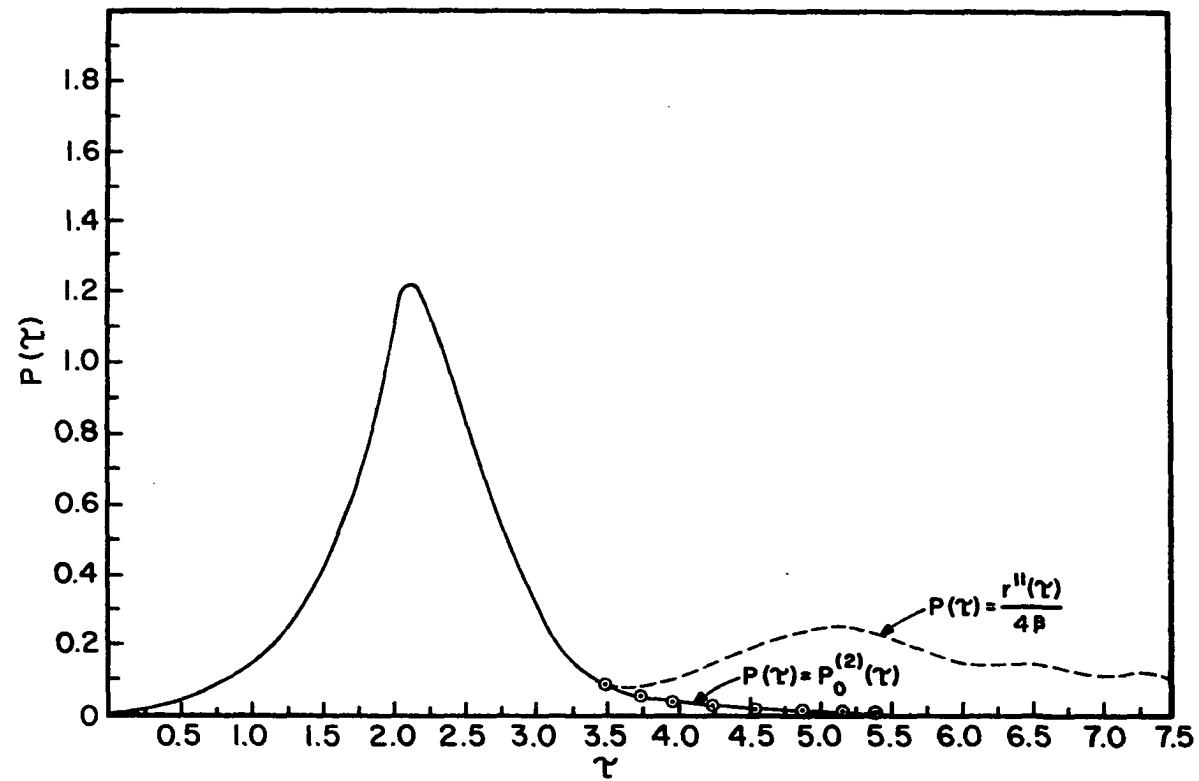


FIGURE 10. TYPICAL ZERO CROSSING DISTRIBUTION FOR BAND-PASS NOISE MODEL COMPUTED BY MEANS OF $P(\tau) = \frac{r''(\tau)}{4\beta}$ AS COMPARED TO $P_0^{(2)}(\tau)$ PROPOSED BY LONGUENT-HIGGINS. EXPERIMENTAL RESULTS OF RANIEL PLOTTED FOR COMPARISON.

Equation (4-45) cannot be solved analytically. A plot of the numerical solution of σ/γ as a function of m is given in Figures 11 and 12. It is interesting to notice that the ratio σ/γ is not a function of the parameter ω_0 and depends only on the band-width of the spectrum m . From Figure 11 it can also be noted that for a given value γ , the standard deviation σ of zero-crossing intervals increases with an increase in the bandwidth of the spectrum. When $m=1$ only the single frequency ω_0 is present and the variance of zero-crossings becomes 0. In this case the random function becomes completely periodic. The mean zero crossing interval given by Equation (4-42) is a function of both ω_0 and m . Notice in particular that the lower the frequencies the higher is the mean zero crossing interval.

The parameters m and ω_0 can now be evaluated from Equation (4-45) and Figure 11. The mean zero crossing interval γ and the standard deviation σ can be evaluated experimentally by standard methods.

Evaluation of the Statistical Properties of the Integrals

Once the parameters m and ω_0 are known and the spectral density is given by

$$G_{xx}(\omega) = \begin{cases} \frac{K}{\omega} & \omega_0 \leq \omega \leq m\omega_0 \\ 0 & \text{elsewhere} \end{cases} \quad (4-38)$$

the mean square error of the integral can be evaluated directly. From the fact that integration between the limits 0, t corresponds to passing the noise through a process with

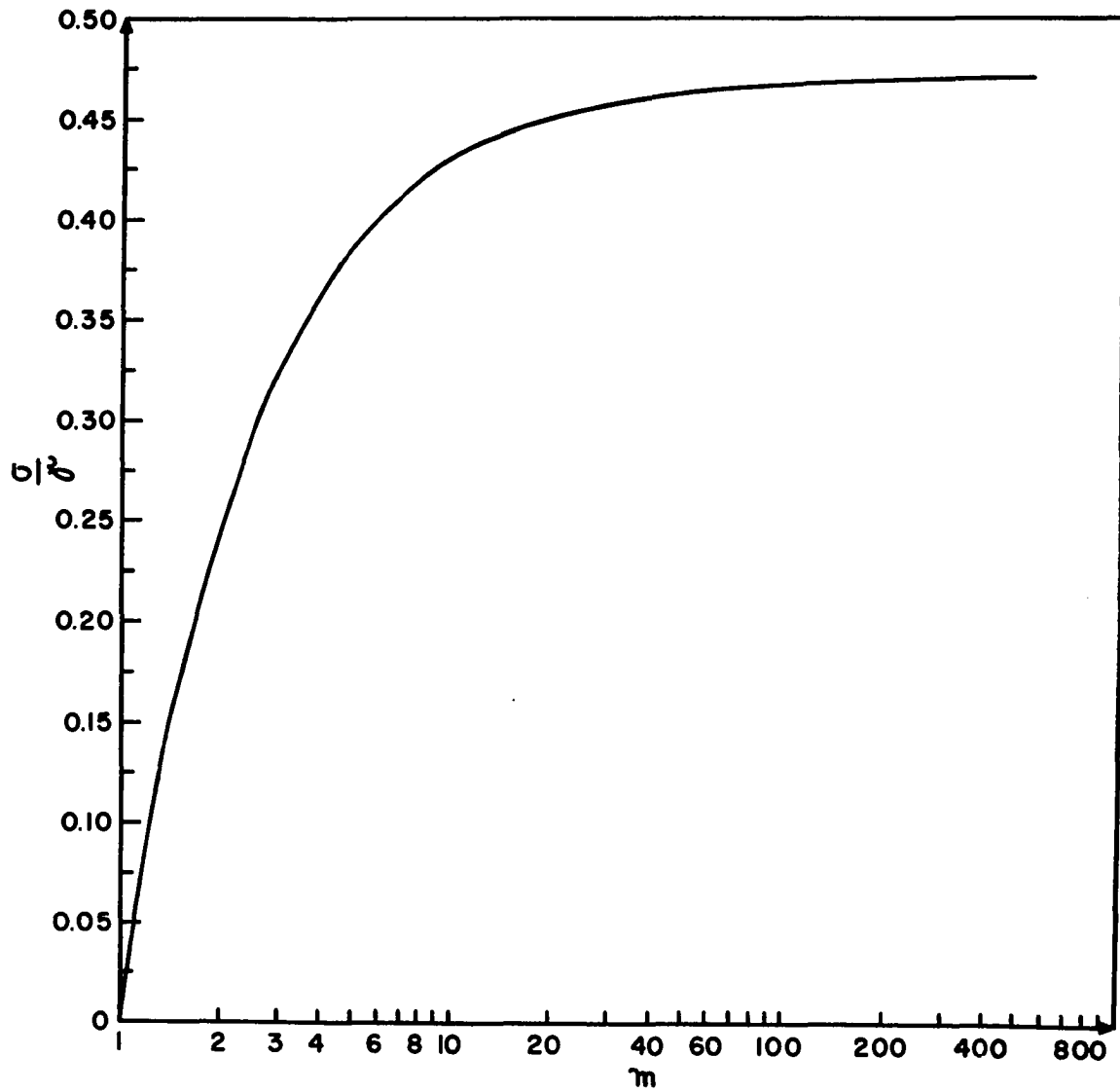


FIGURE II. NUMERICAL SOLUTION OF EQUATION (4-45)

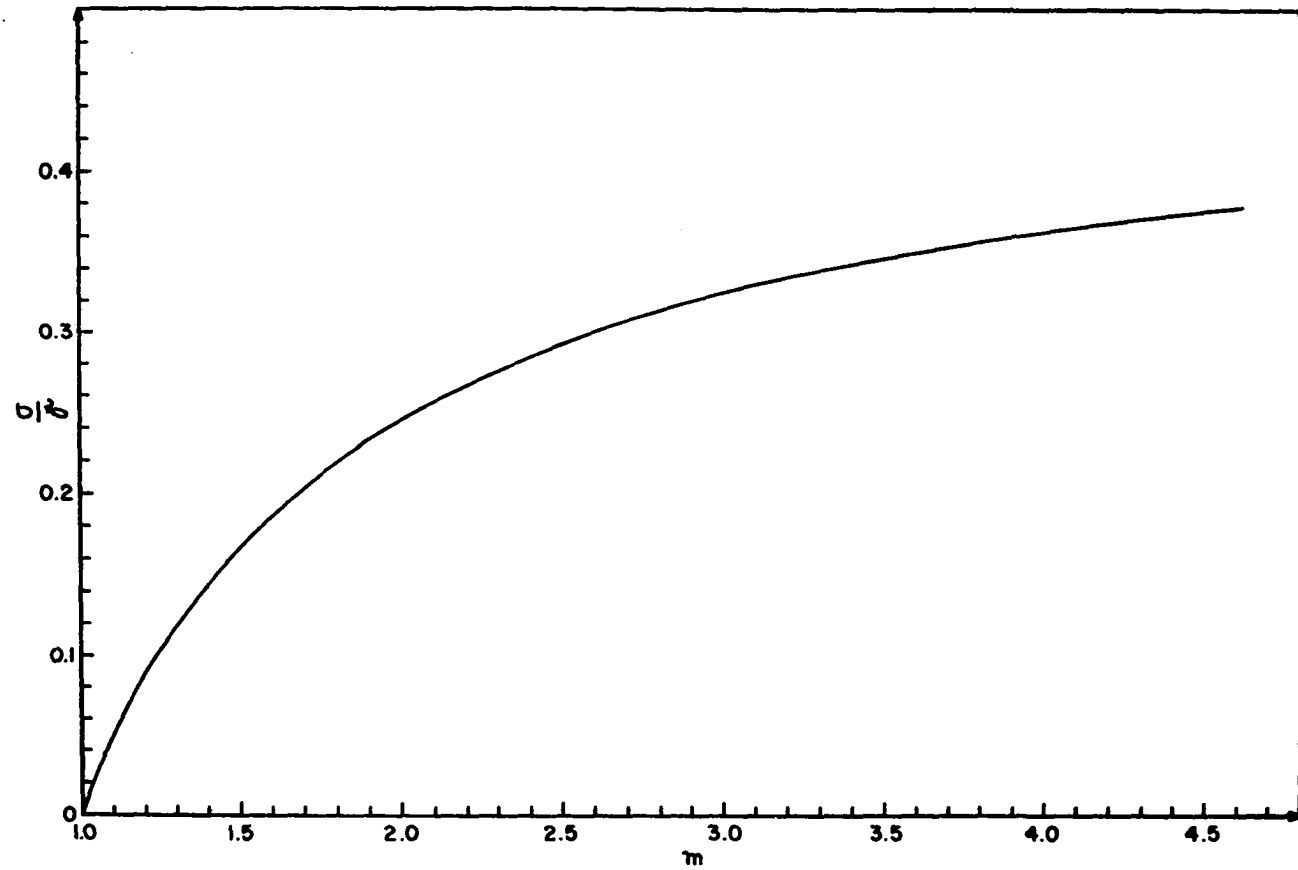


FIGURE 12. NUMERICAL SOLUTION OF EQUATION (4-45)-(LOW RANGE)

a transfer function $1/s$ it follows [27] that:

$$G_{yy}(\omega) = \frac{G_{xx}(\omega)}{|j\omega| |-j\omega|} = \frac{G_{xx}(\omega)}{\omega^2} \quad (4-39)$$

The mean square error of the integral is then given by

$$R_{yy}(0) = \int_0^{\infty} G_{yy}(\omega) d\omega \quad (4-40)$$

Substitution of Equations (4-38) and (4-39) into (4-40) results in

$$R_{yy}(0) = \underline{K} \int_{\omega_0}^{m\omega_0} \frac{1}{\omega^2} d\omega = \frac{\underline{K}}{\omega_0} \left[1 - \frac{1}{m} \right] \quad (4-41)$$

For second integration

$$G_{zz}(\omega) = \frac{G_{yy}(\omega)}{\omega^2} = \frac{G_{xx}(\omega)}{\omega^4} \quad (4-42)$$

Thus the mean square error for the second integral is

$$R_{zz}(0) = \underline{K} \int_{\omega_0}^{m\omega_0} \frac{1}{\omega^4} d\omega = \frac{\underline{K}}{3\omega_0^3} \left[1 - \frac{1}{m^3} \right] \quad (4-43)$$

For higher orders of integrations the above process is repeated.

The ratio between the mean square error of the integrated response and that of the original function is given by

$$\frac{R_{yy}(0)}{R_{xx}(0)} = \frac{\frac{\underline{K}}{\omega_0} \left(\frac{m-1}{m} \right)}{\underline{K}(m-1)\omega_0} = \frac{1}{m\omega_0^2} \quad (4-44)$$

$$\text{or:} \quad \frac{\sigma_y}{\sigma_x} = \frac{1}{\omega_0 \sqrt{m}} \quad (4-45)$$

For the second integration the respective results are

$$\frac{R_{zz}(0)}{R_{xx}(0)} = \frac{m^2 + m + 1}{3^4 \omega_0^4} \quad (4-46)$$

and

$$\frac{\sigma_z}{\sigma_x} = \frac{1}{\omega_0^2} \sqrt{\frac{m^2 + m + 1}{3m^3}} \quad (4-47)$$

Analog Computer Studies

Analog computer simulation studies were performed to evaluate the advantages and limitations of the proposed method for estimation of error propagation through integration of noisy response curves. The primary purpose of this study was to determine the validity of the approximative, band-pass, noise model as a tool for prediction of error propagation.

The analog computer circuit diagram is shown in Figure 13. White noise (the spectral density function of which is constant over all frequencies) was generated by a General Radio Co. Type 1390-A random noise generator. The noise was amplified by a gain of 100 and fed through a potentiometer into a second order linear system. The system was programmed on a Donner Model 3100 D analog computer as modified by Bishop and Sims [5]. The system response was integrated, and both the response and its integral were recorded on a Sanborn, six-channel recorder, Model 156-1100C. Since the random noise generator exhibited a D.C. bias, an adjustable constant voltage was fed into the system to compensate for the drift. A sample record is shown in Figure 14.

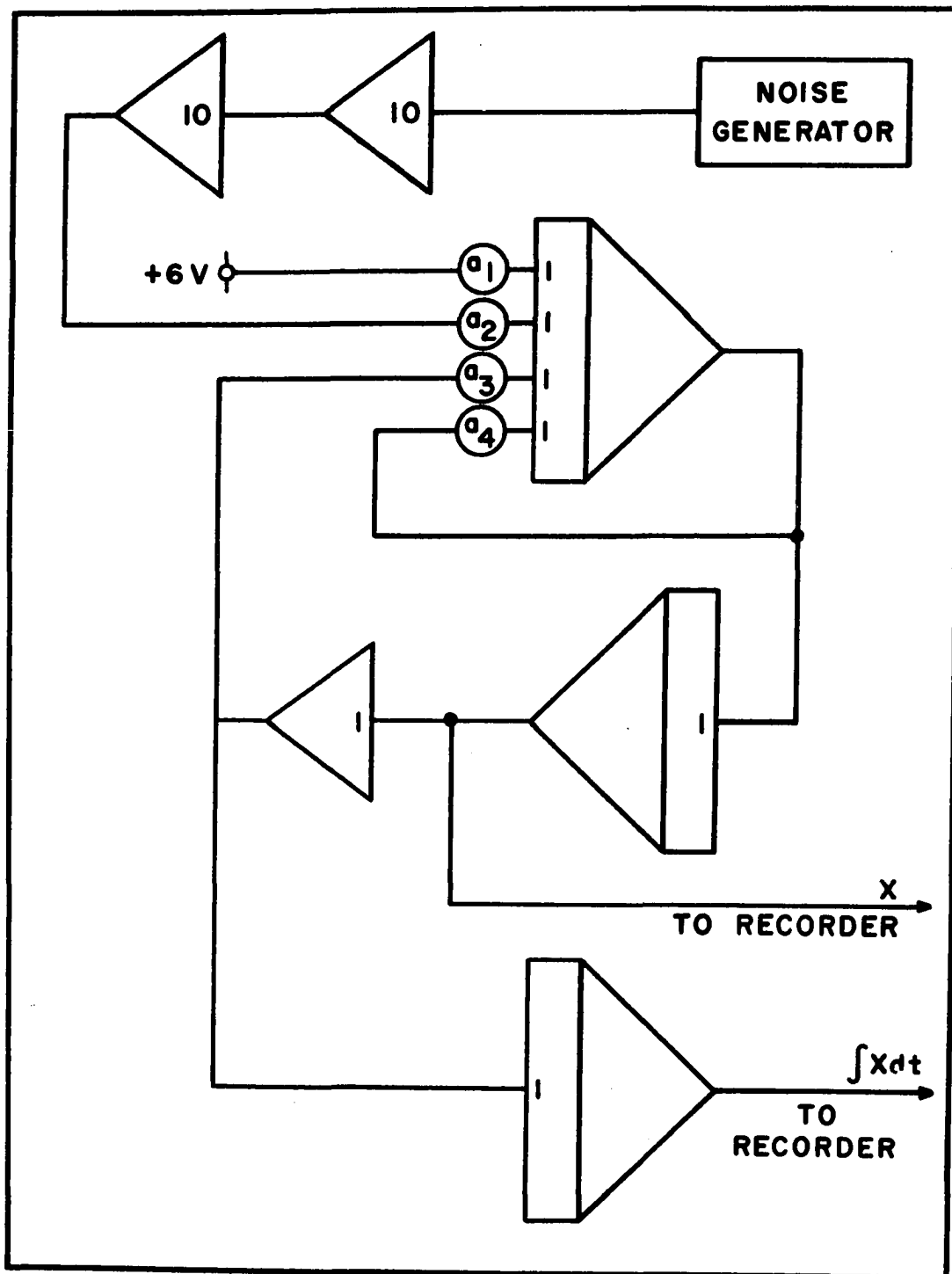


FIGURE 13. ANALOG COMPUTER CIRCUIT
DIAGRAM FOR STUDIES OF ERROR
PROPAGATION IN NOISE INTEGRATION

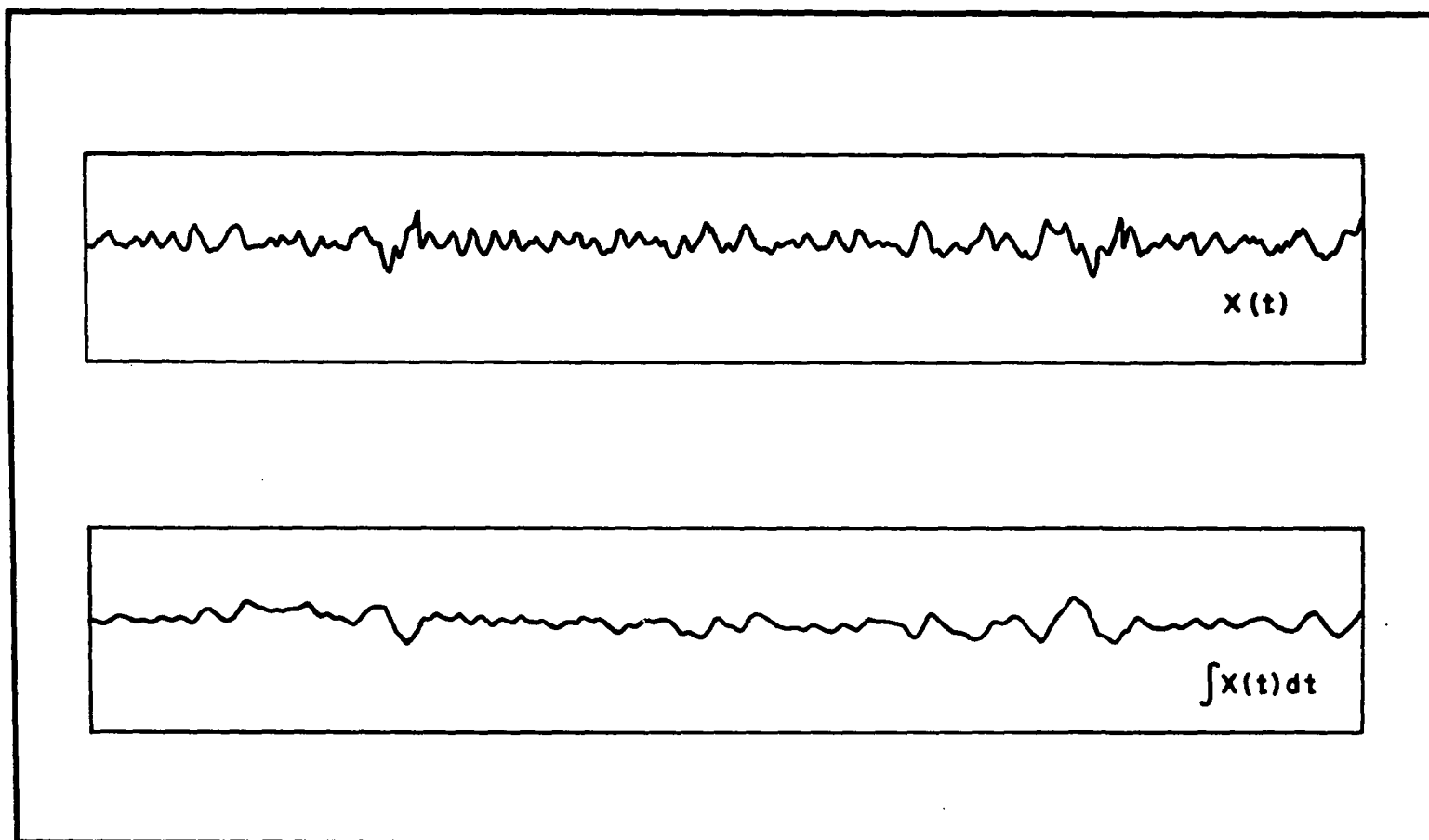


FIGURE 14. SAMPLE RECORD OF RANDOM FUNCTION AND ITS INTEGRAL

In order to estimate the error in the integral of the response curve from the random characteristics of the response itself, it was necessary to evaluate the approximate band pass spectral density of the response. The number and lengths of zero crossing intervals were evaluated from a given record. The mean zero crossing interval γ was determined from knowledge of the total record length and the number of intervals. The standard deviation σ was determined from the interval lengths and the mean interval. The first order approximation of zero crossing intervals in terms of the correlation function of the noise, Equation (4-15), has been previously shown to be valid only for small zero-crossing intervals. It was therefore decided to evaluate the standard deviation σ from only the zero-crossing intervals that are smaller than the mean. This approach proved to be satisfactory.

Thus σ was given by:

$$\sigma = \left[\frac{1}{n} \sum_{\substack{i=1 \\ l_i \leq \gamma}}^n (l_i - \gamma)^2 \right]^{1/2},$$

where n is the number of intervals, l_i is the length of the i^{th} interval and γ is the mean interval length.

Given the ratio σ/γ , the value for m in Equation (4-38) is then found from (4-42). The ratio between the error in the integral and the error in the original response function is then given by Equation (4-45).

The results of the analog computer studies are

presented in Table 1. The various intermediate steps of the calculations are also shown. The values ξ in Table 1 are the damping factor of the second order system; the transfer function of which can be expressed as

$$F(s) = \frac{1}{s^2 + 2\xi\bar{\omega}s + \bar{\omega}^2}$$

where s is the Laplace operator and $\bar{\omega}$ is the undamped natural frequency.

In Figure 15 a comparison between the theoretical spectral density function of the output noise and the approximative, experimentally-evaluated, band pass model is shown for representative cases. The theoretical spectral density of the response noise is given by

$$G(\omega) = F(j\omega) \cdot F(-j\omega) \cdot \bar{K} = \frac{\bar{K}}{(\omega^2 - \bar{\omega}^2)^2 + 4\xi^2 \omega^2 \bar{\omega}^2},$$

where $\bar{K}$ is the spectral density of the input white noise. The band pass model is expressed by Equation (4-38) in terms of m , ω_0 and $\underline{K}$, where $\underline{K}$ is yet to be determined. Since both noise models must have the same mean square value it follows that

$$R_{XX}(0) = \lim_{\tau \rightarrow 0} \int_0^\infty G_{XX}(\omega) \cos \omega\tau d\omega = \int_0^\infty G_{XX}(\omega) d\omega = \underline{K}(m-1)\omega_0.$$

The term on the right expresses the mean square value of the approximative band pass model, Equation (4-40). Thus,

$$\int_0^\infty \frac{\bar{K}d\omega}{(\omega^2 - \bar{\omega}^2)^2 + 4\xi^2 \omega^2 \bar{\omega}^2} = \underline{K}(m-1)\omega_0.$$

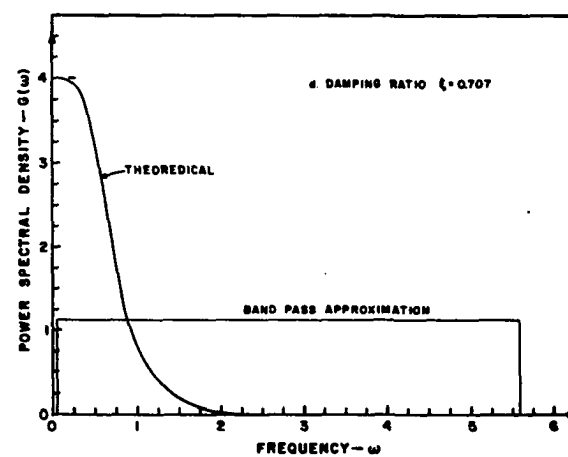
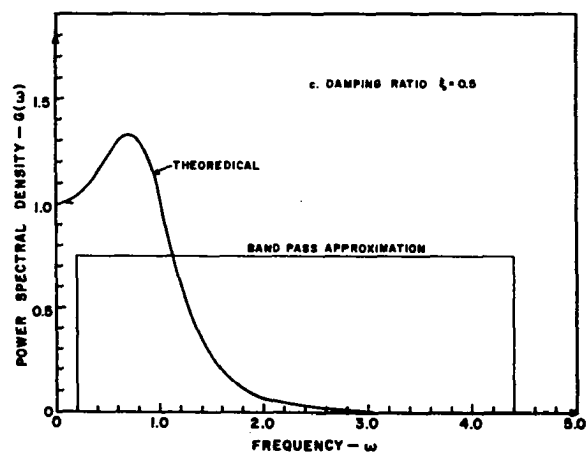
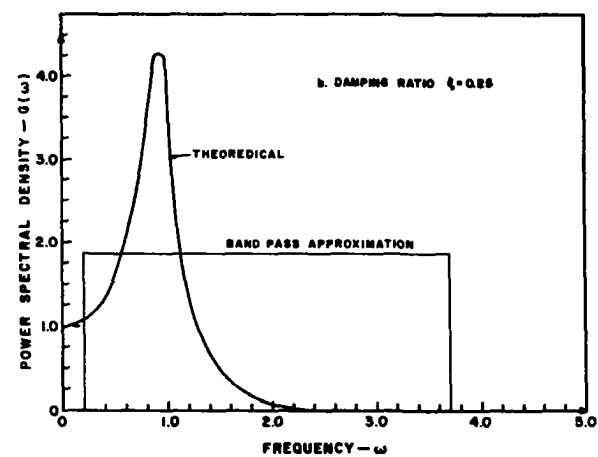
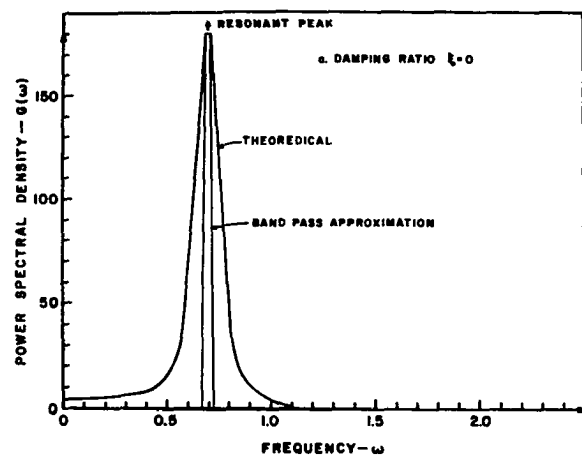


FIGURE 15. COMPARISON BETWEEN EXPERIMENTAL BAND PASS APPROXIMATION AND THEORETICAL SPECTRAL DENSITY FUNCTIONS.

The above integral was evaluated by contour integration to yield:

$$\frac{\pi \bar{K}}{2\xi\omega^3} = \underline{K}(m-1)\omega_0 .$$

If $\bar{K}$ is arbitrarily scaled to be 1 then

$$\underline{K} = \frac{\pi}{2\xi\omega^3(m-1)\omega_0}$$

All the parameters in the above equation are known, thus $\underline{K}$ can be evaluated.

Discussion

The results of the analog computer studies of error propagation into the integrals of the response curves are summarized in Table 1. In the ninth column the predicted ratios σ_y/σ_x of standard deviations of the integrals to the standard deviations of the responses is listed. The observed value for this ratio (determined from actual measurement of amplitude ratios) is given in the next column. The deviation between the two columns in percent is given in the last column. The error of prediction varies between 0 percent up to about 50 percent. This prediction is a satisfactory one, considering the crude noise model that has been assumed. For the present purposes the accuracy of these results is certainly sufficient as it is desired for an order of magnitude estimate, rather than the generation of accurate numbers. The prediction is in general conservative; i.e., there is a general tendency to obtain an overestimation of the error, which

TABLE 1

RESULTS OF ANALOG COMPUTER STUDIES OF ERROR PROPAGATION
INTO INTEGRALS OF RESPONSE CURVES

Test No.	ξ	Total Number of Zero Crossings	γ (sec)	σ (sec)	σ/γ	m	ω_0 (1/sec)	σ_y/σ_x Calculated	σ_y/σ_x Obser.	% Error
1	0	20	4.46	0.146	0.0328	1.03	0.697	1.413	1.271	10.0
2	0	36	3.18	0.153	0.0487	1.04	0.978	1.0	1.0	0.0
3	0.25	63	1.42	0.64	0.45	20.0	0.1858	1.201	0.79	34.2
4	0.3535	52	0.9959	0.2931	0.294	2.5	1.747	0.361	0.362	0.0
5	0.459	44	1.0829	0.3914	0.361	3.9	1.146	0.441	0.461	-4.5
6	0.5	79	1.177	0.528	0.449	19.5	0.23	0.967	0.50	48.2
7	0.5	76	1.133	0.496	0.434	11.0	0.413	0.73	0.581	20.4
8	0.5	55	1.172	0.569	0.460	37.5	0.1225	1.089	0.88	19.2
9	0.5	39	1.107	0.315	0.285	2.4	1.62	0.483	0.312	35.4
10	0.5	249	1.153	0.485	0.42	8.0	0.552	0.835	0.72	13.8
11	0.707	46	1.158	0.54	0.466	90.0	0.05	0.463	0.538	-16.2

is desirable since in the error propagation analysis one is interested in obtaining upper bounds of error.

Figure 15 shows a comparison between the theoretical spectral density function of the response noise and the "approximative" band pass model. It is quite clear that as an approximation to the actual spectral density function the idealized band pass model is totally meaningless. However, this comparison was not the purpose of the present work, and one should not interpret it as such. It should be noticed that for the case where $\xi = 0$ and the system exhibits an infinite resonant peak, the band pass model approximates the spectral density quite well. As ξ becomes larger the approximation becomes worse and shifts toward the higher frequency range. This shift may be a compensating effect for the loss of the D.C. and the low frequency components. However, too much physical meaning should not be associated with the band pass model since the model itself is physically incorrect and unrealizable.

Better estimation of the spectral density function (or correlation function) from zero-crossing information would be possible if a plausible and physically realizable noise model were assumed. The exponential cosine correlation function, for example, would be such a model. Of course, a higher price in terms of additional computation would have to be paid for the evaluation.

For the present work the particular noise model selected was not critical. However, it is the approach rather

than the particular noise model that is of significance. Through further refinement the method may eventually serve as a valid and relatively simple approach for evaluating auto-correlation functions and power spectral densities.

Discrete Analysis

In many cases it may be necessary or desirable to sample the response records and perform the integration numerically. One then faces the following situation:

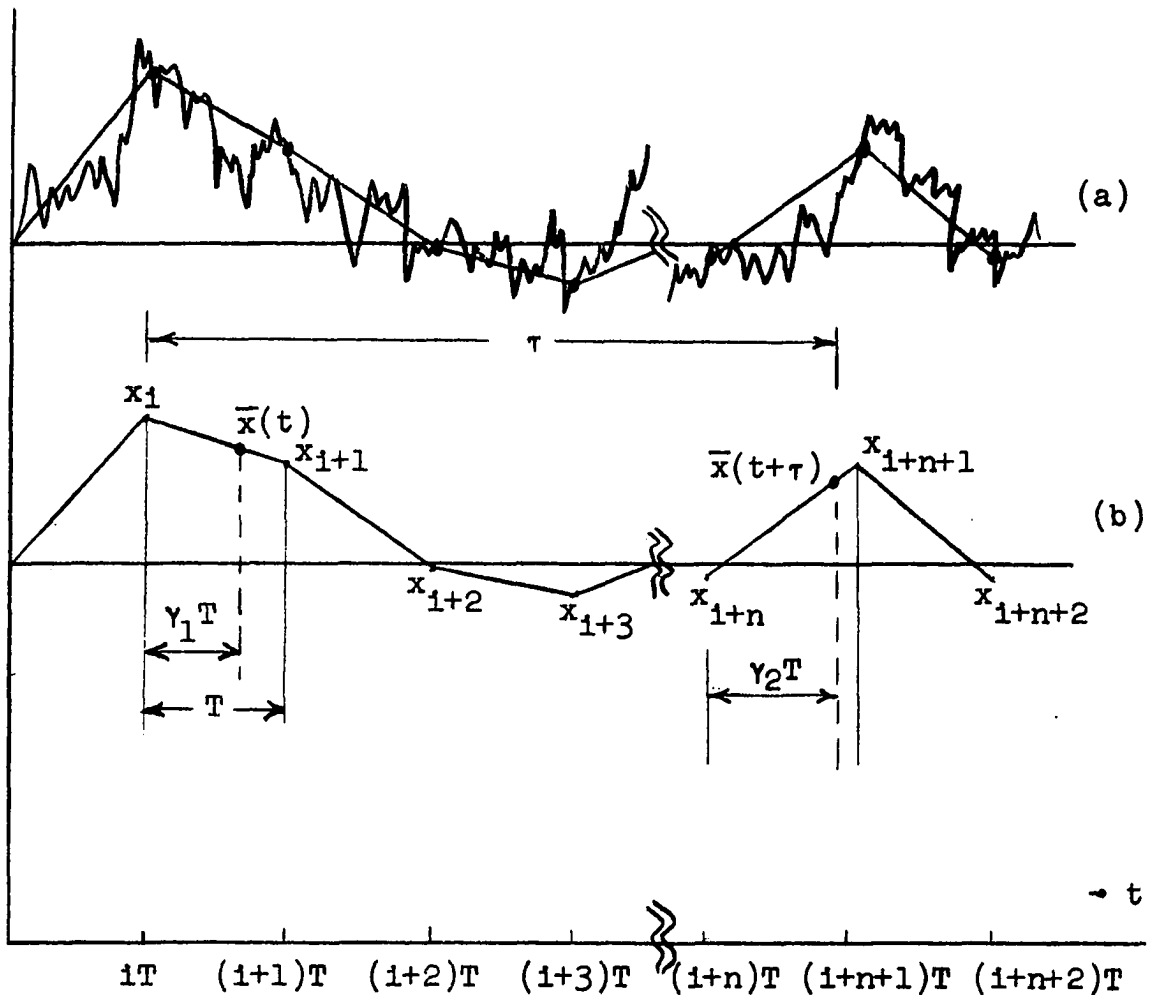


Figure 16. Discrete Sampling of Random Noise

Figure (16-a) describes the continuous random process which is sampled at intervals T . The resultant function in Figure (16-b) is a discrete process in which the sampled points are connected by line segments. The propagation of error through integration will of course also be a function of the particular integration scheme used.

In the present analysis the trapezoidal rule for integration is applied. The statistical distribution of the sampled points is, as can be expected, a function of the statistical properties of the continuous random process. It will be assumed here that the correlation function for the continuous random process is known. This information can be obtained by the methods described in the previous section.

Let the sampling interval be denoted by T , and let x_j be the corresponding sampled points on the continuous random function. The points $\bar{x}(t)$ are points on the discrete record. It is desired to evaluate the expectation $R_{\bar{x}\bar{x}}(\tau) = E[\bar{x}(t)\bar{x}(t + \tau)]$. The interval τ can be expressed as $\tau = (n + \theta)T$ where $\theta \leq 1$ and n is an integer. If $\bar{x}(t)$ is assumed to lie between x_1 and x_{i+1} then $\bar{x}(t+\tau)$ may lie either between x_{i+n} and x_{i+n+1} or between x_{i+n+1} and x_{i+n+2} . Define γ as the fractional distance of $\bar{x}(t)$ between adjacent x_j 's ($0 \leq \gamma \leq 1$). Then for $\bar{x}(t)$ where

$$iT < t < (i + 1)T$$

$$\text{if} \quad 0 \leq \gamma_1 \leq 1 - \theta$$

$$\text{then} \quad x_{i+n} < \bar{x}(t + \tau) < x_{i+n+1}$$

$$\text{such that} \quad 0 \leq \gamma_2 \leq 1.$$

However if $1 - \theta \leq \gamma_1 \leq 1$

then $x_{1+n+1} < \bar{x}(t + \tau) < x_{1+n+2}$

such that $0 \leq \gamma_2 \leq \theta$.

Thus if $\bar{x}(t)$ is denoted by $\bar{x}_1$ and $\bar{x}(t+\tau)$ by $\bar{x}_2$, then

$$\bar{x}_1 = x_1(1 - \gamma_1) + x_{1+1}\gamma_1$$

$$\bar{x}_2 = \begin{cases} x_{1+n}(1 - \gamma_2) + x_{1+n+1}\gamma_2 & \text{if } (1+n)T < t+\tau < (1+n+1)T \\ x_{1+n+1}(1 - \gamma_2) + x_{1+n+2}\gamma_2 & \text{if } (1+n+1)T < t+\tau < (1+n+2)T \end{cases} \quad (4-48)$$

It can be seen from Figure (16-b) that $\gamma_2 = \gamma_1 + \theta$ when $0 \leq \gamma_1 \leq 1 - \theta$ and that $\gamma_2 = \gamma_1 + \theta - 1$ when $1 - \theta \leq \gamma_1 \leq 1$.

Substitution of the above values for γ_2 in Equation (4-48) yields

$$(a) \bar{x}_2 = x_{1+n}(1 - \gamma_1 - \theta) + x_{1+n+1}(\gamma_1 + \theta)$$

$$\text{when} \quad (1+n)T < t+\tau < (1+n+1)T \quad (4-49)$$

$$(b) \bar{x}_2 = x_{1+n+1}(2 - \gamma_1 - \theta) + x_{1+n+2}(\gamma_1 + \theta - 1)$$

$$\text{when} \quad (1+n+1)T < t+\tau < (1+n+2)T.$$

From which it immediately follows that

$$(a) \bar{x}_1 \bar{x}_2 = [x_1(1 - \gamma_1) + x_{1+1}\gamma_1][x_{1+n}(1 - \gamma_1 - \theta) + x_{1+n+1}(\gamma_1 + \theta)]$$

$$\text{for } 0 \leq \gamma_1 \leq 1 - \theta$$

(4-50)

$$\begin{aligned}
 (b) \quad \bar{x}_1 \bar{x}_2 &= [x_1(1 - \gamma_1) + x_{i+1}\gamma_1][x_{i+n-1}(2 - \gamma_1 - \theta) \\
 &\quad + x_{i+n+2}(\gamma_1 + \theta - 1)] \\
 &\text{for } 1 - \theta \leq \gamma_1 \leq 1
 \end{aligned} \tag{4-50}$$

continued

The autocorrelation function $R_{xx}(\tau)$ is then given by

$$\begin{aligned}
 R_{xx}(\tau) &= \lim_{m \rightarrow \infty} \frac{1}{2m+1} \sum_{i=-m}^m \left\{ \int_0^{1-\theta} [x_1(1 - \gamma_1) + x_{i+1}\gamma_1] \cdot \right. \\
 &\quad [x_{i+n}(1 - \gamma_1 - \theta) + x_{i+n+1}(\gamma_1 + \theta)] d\gamma_1 \\
 &\quad + \int_{1-\theta}^1 [x_1(1 - \gamma_1) + x_{i+1}\gamma_1][x_{i+n+1}(2 - \gamma_1 - \theta) \\
 &\quad \left. + x_{i+n+2}(\gamma_1 + \theta - 1)] d\gamma_1 \right\} \tag{4-51}
 \end{aligned}$$

Evaluation of the integral terms in Equation (4-51) leads to

$$\begin{aligned}
 R_{xx}(\tau) &= \lim_{m \rightarrow \infty} \frac{1}{2m+1} \sum_{i=-m}^m \left\{ x_{i+1}x_{i+n} \frac{1}{6} (1 - \theta)^3 \right. \\
 &\quad + x_1x_{i+n+1} [\theta(1 - \theta) + \frac{1}{6} (1 - \theta)^3] \\
 &\quad + x_{i+1}x_{i+n+1} [\frac{1}{3} (1 - \theta)^3 + \frac{1}{2} \theta(1 - \theta)^2] \\
 &\quad + x_1x_{i+n} [\frac{1}{2} (1 - \theta)^2 - \frac{1}{6} (1 - \theta)^3] + x_1x_{i+n+1} [\frac{1}{2} \theta^2 \\
 &\quad - \frac{1}{6} \theta^3] + x_1x_{i+n+2} \frac{1}{6} \theta^3 + x_{i+1}x_{i+n+1} [(1 - \theta)\theta + \frac{1}{6} \theta^3] \\
 &\quad \left. + x_{i+1}x_{i+n+2} [\frac{1}{2} \theta^2 - \frac{1}{6} \theta^3] \right\} \tag{4-52}
 \end{aligned}$$

Since θ is constant, Equation (4-52) can be further simplified to yield, after collecting terms,

$$\begin{aligned}
R_{\overline{XX}}(\tau) = & R_{XX}[(n-1)T] \left[\frac{1}{6} (1-\theta)^3 \right] + R_{XX}[nT] \left[\frac{1}{3} (1-\theta)^3 \right. \\
& + \frac{1}{2} \theta (1-\theta)^2 + \frac{1}{2} (1-\theta)^2 - \frac{1}{6} (1-\theta)^3 + \theta(1-\theta) \\
& + \left. \frac{1}{6} \theta^3 \right] + R_{XX}[(n+1)T] \left[\theta(1-\theta) + \frac{1}{6} (1-\theta)^3 \right. \\
& + \left. \frac{1}{2} \theta^2 - \frac{1}{6} \theta^3 \right] + R_{XX}[(n+2)T] \frac{1}{6} \theta^3 \quad (4-53)
\end{aligned}$$

which after further simplification yields the final result

$$\begin{aligned}
R_{\overline{XX}}(\tau) = & R_{XX}[(n-1)T] \left[\frac{1}{6} - \frac{1}{2} \theta + \frac{1}{2} \theta^2 - \frac{1}{6} \theta^3 \right] \\
& + R_{XX}(nT) \left[\frac{2}{3} - \theta^2 + \frac{1}{2} \theta^3 \right] + R_{XX}[(n+1)T] \left[\frac{1}{6} + \frac{1}{2} \theta \right. \\
& + \left. \frac{1}{2} \theta^2 - \frac{1}{2} \theta^3 \right] + \frac{1}{6} R_{XX}[(n+2)T] \theta^3, \quad (4-54)
\end{aligned}$$

where $\tau = (\theta + n)T$.

The mean square value of $\overline{x}$ is given by

$$R_{\overline{XX}}(0) = \frac{2}{3} R_{XX}(0) + \frac{1}{3} R_{XX}(T). \quad (4-55)$$

Obviously the integral of the sampled random process $\overline{x}$ is generally nonstationary, and a drift from the true value will occur. As the sampling interval T becomes smaller, the correlation function $R_{\overline{XX}}(\tau)$ approaches $R_{XX}(\tau)$ and the mean square value $R_{\overline{XX}}(0)$ approaches $R_{XX}(0)$. It is therefore generally desirable to reduce the sampling interval if the noise magnitude is significant. Evaluation of the correlation function of the integral of $\overline{x}(t)$ is possible by methods similar to those used for the continuous analysis. The results, however, are quite complex and not tractable. When the sampling interval is small enough (i.e., the correlation

function of the discrete case $R_{xx}(\tau)$ is close enough to the correlation function of the continuous case $R_{xx}(\tau)$ the continuous analysis should be used as an approximation.

CHAPTER V

FORMALIZATION OF THE TECHNIQUE

In the previous chapters, various individual aspects of the identification technique have been discussed. Here an integrated and formalized description of the actual procedure will be presented.

In order to implement the identification technique the number of linearly independent response functions corresponding to the suspected order of the system have to be experimentally generated. The techniques for generating these responses as well as the method for determining the correct order of the system have been or will be discussed.

The error propagation analysis associated with the identification serves not only the purpose of estimation of the uncertainty of the results, but must also provide a tool for the minimization of this uncertainty. It is quite evident at this point that the proposed identification method hinges strongly on this latter aspect of the error analysis. When accurate estimates of the errors associated with the identification are desired, good estimates of the uncertainty of the data are required. However, when the error propagation serves primarily as a tool for the execution of the identification

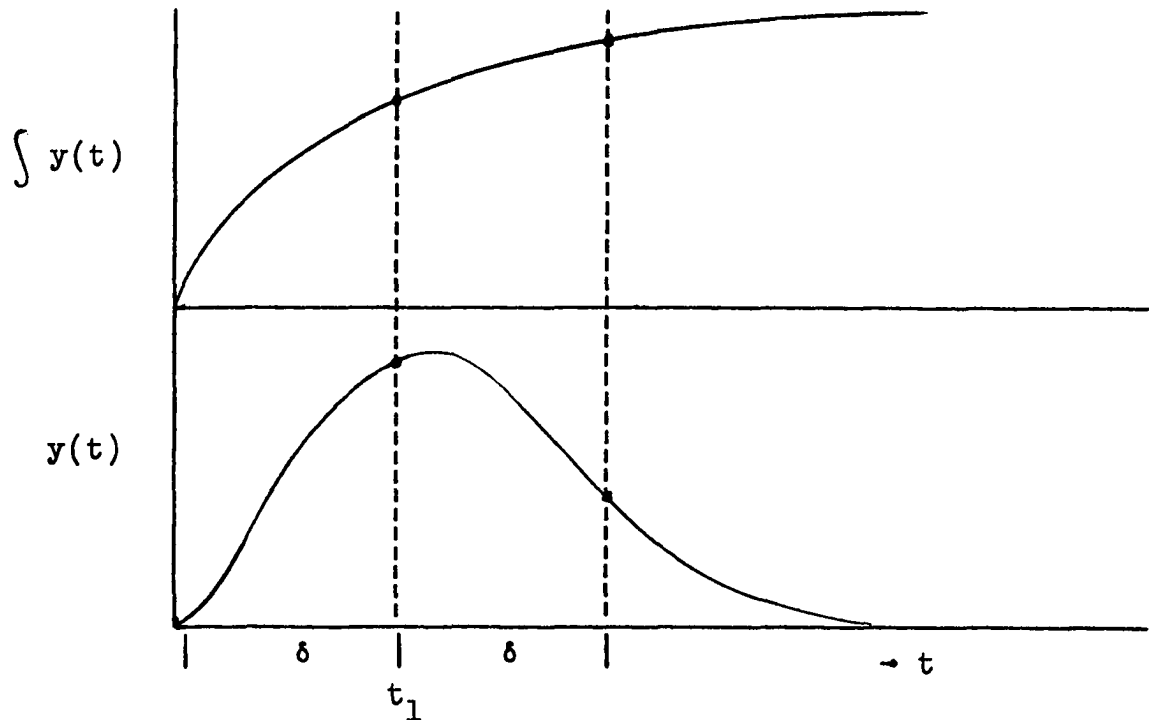
process, no stringent requirements are necessary when estimating the error of the data, since the error minimization characteristics depend only indirectly on the magnitude of this error.

In the last section of Chapter III it has been shown how the uncertainty of the identified poles depends on the selection of the integration parameter δ . In order to obtain the best identification, δ must be selected so as to minimize the relative error associated with the poles. However, since the values for the poles are unknown, δ cannot be obtained a priori. Furthermore, since all the relative errors are not minimized simultaneously (at the same δ) there does not, in general, exist a single unique δ at which the identification for all poles should be carried out. If the responses were noise free and the error were a smooth function of δ with only one relative minimum (except for the case of the imaginary parts of complex poles) as in Figure 5, a simple iteration procedure could be established for minimization of the errors associated with the various poles. Each pole would then be treated individually and would be evaluated at a particular value of δ . Consequently the best and most rapid identification of the system would result.

In practice, since the response functions are generally noisy, the errors (as a function of δ) fluctuate and an iterative procedure cannot easily be employed. In the present work no attempt was made to establish the most efficient process for the error minimization. Rather, an extensive

study of the error propagation characteristics was performed.

In order to describe the actual identification and error minimization procedure utilized, consider the illustration presented below where one of the two response functions and its integral of a second order system are schematically shown.



The fundamental matrices $\Phi_{n-1}(t_{n-1})$ and $\Phi_{n-1}(t_{n-1} - \delta)$ are constructed from the response curves and their integrals for the selected δ . The limit of integration δ is initially selected as a small value such that, in the investigator's judgement, it is not larger than a fraction of the lowest time constant of the system to be identified. Thus the initially generated fundamental matrices are evaluated over a relatively narrow segment of the response function. The

identification routine is then initialized to generate the natural frequencies and their errors. At this point δ is increased by a constant interval and the process repeated until the records are completely exhausted. The value of δ at which the smallest relative error is obtained for each individual pole is considered as the value of δ at which this pole is identified. Figure 17 shows schematically the formalized structure of the identification process. Digital computer programs for the identification procedure corresponding to the schematic structure of Figure 17, written in the Osage Algol compiler language,* are listed in Appendix E. It should be noticed that programs for preparation of the data, i.e., construction of the fundamental matrices $\Phi_{n-1}(t_{n-1} - \delta)$ and $\Phi_{n-1}(t_{n-1})$, have not been included in Appendix E. They are of no general use and depend entirely on the measuring and recording facilities available. Figure 17 follows step by step through the computational stages required for the identification and the error analysis. The error evaluation and the main identification computation are performed simultaneously, i.e. the error evaluation for each computational stage is performed immediately following the calculations for this stage. The bold faced boxes of Figure 17 represent the main computational stages while the bottom part of the diagram represents the error minimization section of the

*The Osage Algol language is a slightly modified version of the Algol 60 language [39,56], used on the Osage high speed computer at the University of Oklahoma Computer Center.

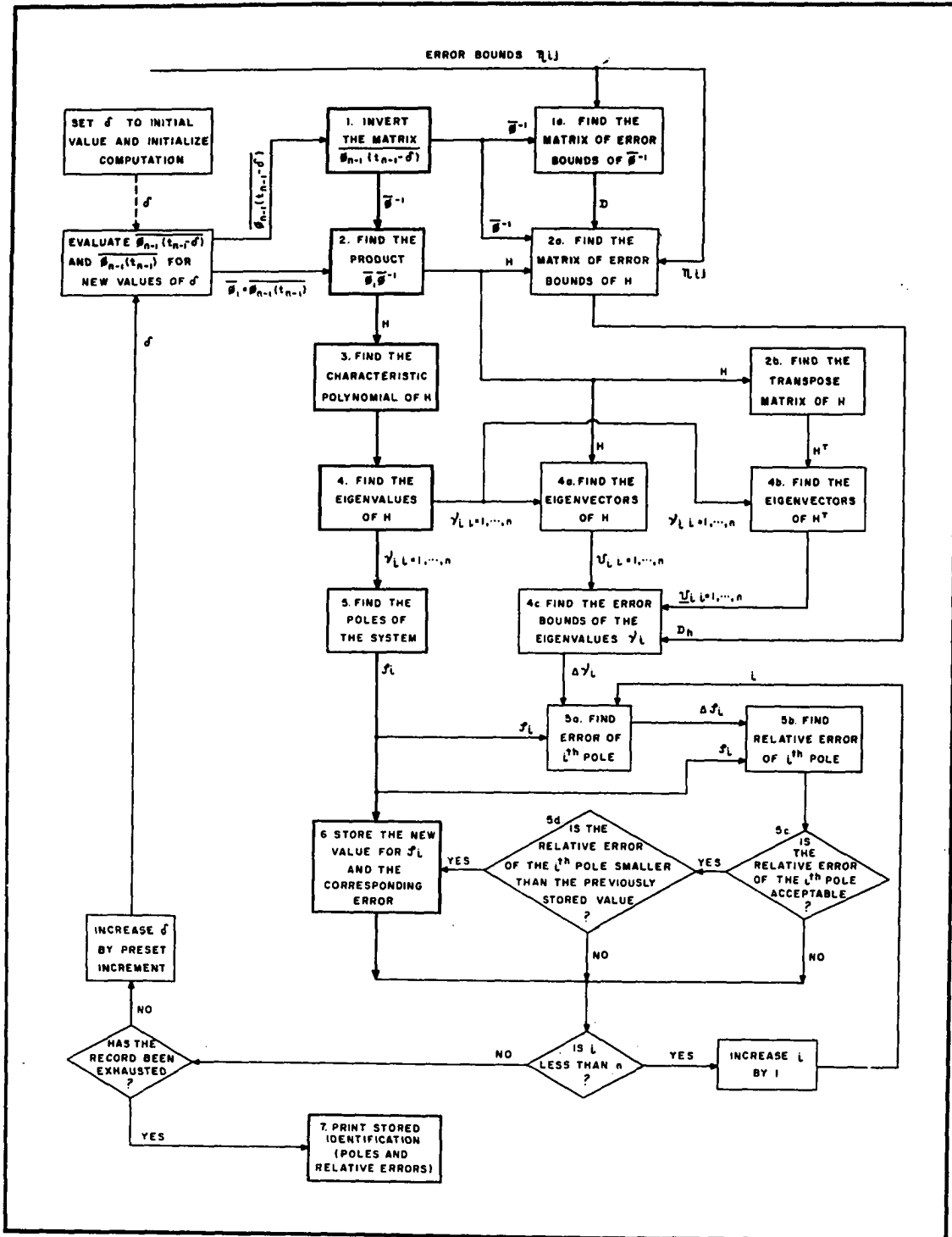


FIGURE 17. FORMALIZED BLOCK DIAGRAM OF THE IDENTIFICATION TECHNIQUE.

calculation.

In the following paragraphs several comments concerning the actual utilization of the technique are made.

1) Selection of the System Order. In the present technique the order of the system to be identified has to be assumed beforehand. Frequently one may have a preconceived idea of the order of the equations describing the tested plant or of the desired order of the model. In these cases a model of the desired order is to be found. When no knowledge of the system order is available, the order of the model can be arbitrarily assumed, and the identification can be performed with respect to this order. If the order assumed is too high the identification will nevertheless be carried out properly and the actual order will be indicated by the results through the error estimates. If the assumed order is too low the error criteria will indicate such. This problem is further discussed in Chapter VI in association with the results of the experimental investigation.

2) Type of Required Plant Test Data. The only requirement of the test data is that the n response functions be linearly independent. For identification of the General Model the responses have to be zero-input responses, i.e. the responses generated during the relaxation of the system following a forced input.

3) Construction of the Fundamental Matrices $\Phi_{n-1}(t_{n-1})$ and $\Phi_{n-1}(t_{n-1} - \delta)$. The construction of the fundamental matrices is performed by formally expanding Equations (2-77) -

(2-80). For illustration purposes the expanded expressions for second and third order systems are given below. Thus for second order systems

$$\overline{\varphi_{1j}(t_1 - \delta)} = \varphi_{j,2-1}(t_1) - \varphi_{j,2-1}(t_1 - \delta) \quad (5-1)$$

and

$$\overline{\varphi_{1j}(t_1)} = \varphi_{j,2-1}(t_1 + \delta) - \varphi_{j,2-1}(t_1), \quad (5-2)$$

where in the above equations $\overline{\varphi_{1j}(\)}$ represents the $1j^{\text{th}}$ term of the fundamental matrix. The subscript j on the right hand side of the equations represents the j^{th} response record, and the second index represents the $(2-1)^{\text{th}}$ integral of the response curve. Thus for example, $\varphi_{2,1}(t_1 + \delta)$ is the value of the first integral of the second response function evaluated at $t_1 + \delta$. Similarly for third order systems

$$\overline{\varphi_{1j}(t_2 - \delta)} = \varphi_{j,3-1}(t_2 + \delta) - 2\varphi_{j,3-1}(t_2) + \varphi_{j,3-1}(t_2 - \delta) \quad (5-3)$$

and

$$\overline{\varphi_{1j}(t_2)} = \varphi_{j,3-1}(t_2 + 2\delta) - 2\varphi_{j,3-1}(t_2 + \delta) + \varphi_{j,3-1}(t_2). \quad (5-4)$$

It should be remembered that in the above equations the response functions and their respective integrals are perturbation variables, i.e. evaluated with respect to steady state as the zero operation level.

4) Evaluation of the Imaginary Parts of Complex Poles.

It follows from Equation (3-48) that

$$\operatorname{Re}(v_1) = e^{\delta \operatorname{Re}(\rho_1)} \cos [\delta \operatorname{Im}(\rho_1)] , \quad (5-5)$$

from which

$$\delta \operatorname{Im}(\rho_1) = \cos^{-1} \left[\frac{\operatorname{Re}(v_1)}{e^{\delta \operatorname{Re}(\rho_1)}} \right] . \quad (5-6)$$

Since the inverse trigonometric functions are multi-valued functions of their arguments, it follows that $\delta \operatorname{Im}(\rho_1)$ is given only within a multiplicity of $2\pi k$, where k is unknown. Thus without knowledge of k , $\operatorname{Im}(\rho_1)$ cannot be determined. This difficulty can be eliminated by evaluation of the system at two close values of δ and by proper manipulation of the results. This procedure is, however, not necessary and much simpler solutions of the problem are possible. The easiest and usually a satisfactory approach for evaluation of k is through inspection of the response records from which $\operatorname{Im}(\rho_1)$ can be determined approximately. Another approach to determine an approximate value for $\operatorname{Im}(\rho_1)$ is through the error estimation. It has been shown in the last section of Chapter III that the relative error of the imaginary parts of complex poles possesses a periodic oscillation as a function of δ and approaches infinity at intervals of $\delta \operatorname{Im}(\rho_1) = \pi$. Thus, from the error plot as a function of δ , $\operatorname{Im}(\rho_1)$ can be obtained sufficiently accurately to determine k .

5) Selection of the parameter t_{n-1} . Although the computer programs in Appendix E have provision for variation of t_{n-1} , this parameter should always be taken equal to δ as was discussed in Chapter III.

CHAPTER VI

ANALOG COMPUTER SIMULATION STUDIES

The experimental investigation of the identification technique was carried out for the following purposes:

(1) To investigate and demonstrate the capability of the technique to identify systems over a wide range of natural frequency ratios for both real and complex poles, in particular in the regions where the identification is usually difficult.

(2) To study the error propagation and error minimization characteristics of the technique and to compare the results to conclusions from the theoretical studies.

(3) To investigate the sensitivity of the technique to noise and to study the error propagation characteristics under these conditions.

(4) To investigate the effect of estimation of error bounds of the response data on the prediction of the error in the results.

(5) To investigate methods for securing linearly independent responses and to study the sensitivity of the identification method when only weak linear independence between the responses is present.

(6) To investigate the problem of misidentification, i.e. the identification when the order of the system is not assumed correctly.

(7) To study the effect of variation of the integration parameter t_{n-1} .

(8) To investigate the possibility of smoothing out the noise and improving the results through higher order integration of the response records, i.e. through the utilization of fundamental matrices $\bar{\Phi}$ of order higher than $n-1$.

(9) To compare the technique to the pulse testing and Fourier transformation identification method.

Construction of an experimental apparatus that permits an investigation over such a wide range of conditions is not feasible. Furthermore, accurate control over the operating conditions and precise knowledge of the parameters of such an apparatus could not be easily obtained. An effective way to perform the proposed study was to simulate the systems on an analog computer where the equations could be accurately programmed and their performance controlled to obtain the responses. Identification calculations were performed on the Osage computer.

The investigation was divided into three general phases:

(1) Analog computer simulation and identification studies of second order systems.

(2) Analog computer simulation and identification studies of third order systems.

(3) Identification of actual chemical plant test data by means of a digital computer.

The third phase of the investigation was carried out to demonstrate the identification of an actual chemical plant of high degree of complexity which exhibits strong noise and nonlinear effects.

Analog Computer Simulations

The simulation studies were performed on the modified Donner Model 3100D Analog Computer [5]. The computer is a medium size analog computer with an operating range of ± 100 volts. The circuit diagram is shown in Figure 18. The analog computer was divided into four major functional sections: (1) the simulated system, (2) the noise generator circuitry, (3) the forcing (or pulsing) circuitry and (4) the output devices. Each of the above functional sections is described and discussed below.

(1) The Simulated System: is a third order, constant parameter, linear differential system consisting of the dashed box in Figure 18 containing integrators I1, I2, and I3. The mathematical equations describing the system are

$$\begin{aligned}\dot{y}_1(t) &= a_{11}y_1(t) + a_{12}y_2(t) + a_{13}y_3(t) + x_1(t) + x(t) \\ \dot{y}_2(t) &= a_{21}y_1(t) + a_{22}y_2(t) + a_{23}y_3(t) + x_2(t) \\ \dot{y}_3(t) &= a_{31}y_1(t) + a_{32}y_2(t) + a_{33}y_3(t) + x_3(t),\end{aligned}\tag{6-1}$$

where $y_1(t)$, $y_2(t)$ and $y_3(t)$ are the outputs of integrators

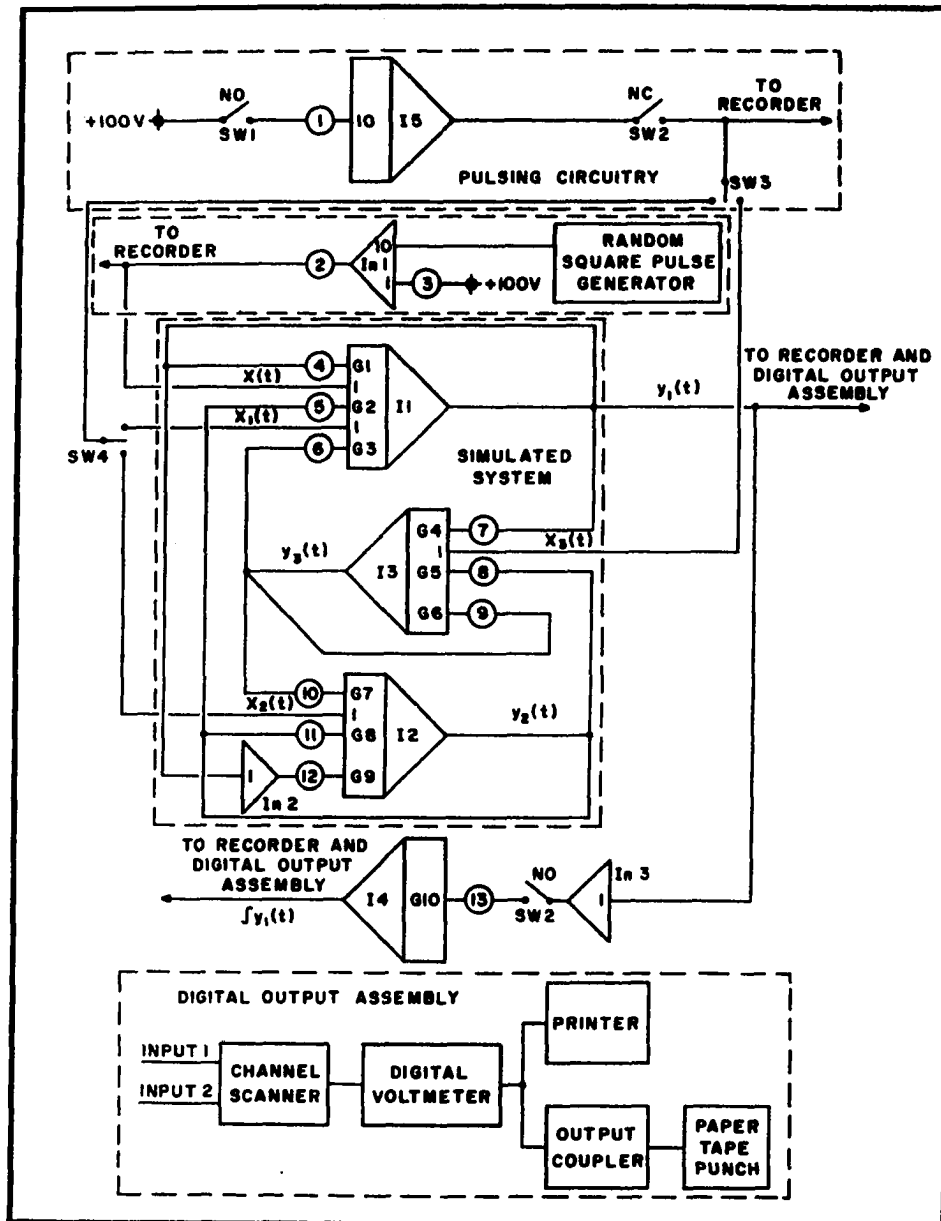


FIGURE 18. ANALOG COMPUTER CIRCUIT DIAGRAM AND DIGITAL OUTPUT ASSEMBLY FOR SIMULATION STUDIES.

I1, I2 and I3 respectively, $x_1(t)$, $x_2(t)$ and $x_3(t)$ are the respective forcing inputs, and $x(t)$ is the noise input into integrator I1. The parameters a_{ij} are adjusted by proper setting of the coefficient potentiometers 4-12 and the input gains G1-G9. The inverter In2 is tentatively programmed for the purpose of sign adjustments of the coefficients. The system also simulates a second order system by setting potentiometers 6-10 to zero, thereby reducing Equation (6-1) to

$$\begin{aligned}\dot{y}_1(t) &= a_{11}y_1(t) + a_{12}y_2(t) + x_1(t) + x(t) \\ \dot{y}_2(t) &= a_{21}y_1(t) + a_{22}y_2(t) + x_2(t) .\end{aligned}\tag{6-2}$$

The system was tested and identified with respect to $y_1(t)$, and by elimination of $y_2(t)$ and $y_3(t)$ from Equation (6-1) the following scalar third order equation is obtained:

$$\ddot{\ddot{y}}_1(t) + a_1\ddot{y}_1(t) + a_2\dot{y}_1(t) + a_3y_1(t) = f(t) + f_x(t) , \tag{6-3}$$

where

$$\begin{aligned}a_1 &= a_{11} + a_{22} + a_{33} \\ a_2 &= a_{11}a_{22} + a_{11}a_{33} + a_{22}a_{33} - a_{12}a_{21} - a_{13}a_{31} - a_{23}a_{32} \\ a_3 &= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{32}a_{21} - a_{11}a_{23}a_{32} - a_{22}a_{31}a_{13} \\ &\quad - a_{33}a_{12}a_{21} .\end{aligned}$$

The forcing function $f(t)$ in (6-3) is the combined term obtained as a result of the forcing functions $x_1(t)$, $x_2(t)$ and $x_3(t)$. The forcing function $f_x(t)$ is the result of $x(t)$ and

is itself a random function.

The second order system of Equation (6-2) is reduced by elimination of $y_2(t)$ to

$$\ddot{y}_1(t) + a_1 \dot{y}_1(t) + a_2 y_1(t) = f(t) + f_x(t)$$

where

$$a_1 = a_{11} + a_{22}$$

$$a_2 = a_{11}a_{22} - a_{12}a_{21}$$

(2) The Noise Generator Circuitry: The noise generator was a random square pulse generator. It consisted of a radioactive source and a counting tube that generated high frequency random impulses, a time base oscillator and an electronic coincidence circuitry. Whenever a time base pulse and a random impulse occurred simultaneously, switching of the square pulse occurred. The mean frequency of the square pulse was adjustable, within the range of about 7 to 17 zero-crossings per second, by variation of the time base frequency. The probability distribution of the random pulse is Poissonian, but since its frequency was high compared to the system's natural frequencies, it could be considered Gaussian (by the central limit theorem [27]). A sample of a high speed recording of the random noise is given in Figure 27a. Since the noise generator provided a pulse with amplitude range 0-(-)20 volts, it was necessary to adjust the DC level in order to obtain a symmetric pulse (± 10 volts). This requirement was accomplished by superposition of a DC voltage that was adjustable through potentiometer 8.

(3) The pulsing circuitry: This circuit consisted of integrator I5 which was fed an adjustable voltage through potentiometer 1 and a series of switches. When switch SW1 is closed and SW2 is in normal position, a ramp function is generated which fed to one of the integrators I1, I2 or I3 by proper setting of switches SW3 and SW4. When SW1 is switched to normal position, the ramp generation is interrupted and a constant voltage is fed to the integrator. When SW2 is switched the input is removed from the system and integration of output $y_1(t)$ begins in I4. By variation of the switching procedure it is possible to generate a wide variety of pulse shapes and durations. Slight modification of the circuitry of switches SW3 and SW4 permitted the pulsation of two or all the integrators simultaneously.

(4) Output devices: These components consisted of the digital output assembly and of a Sanborn six channel recorder model 156-1100C. The digital output assembly that was used to record the system's responses consisted of: (1) Dymec model 2900B channel scanner, (2) Hewlett-Packard model 3440A digital voltmeter, (3) Hewlett-Packard model H75-562A digital recorder (printer), (4) Dymec model 2540 coupler and (5) Friden model SP-2 paper-tape punch. The sampling frequency of the digital output assembly was one reading per second, and the recorded values were simultaneously printed and punched on paper-tape. The punched paper-tape was directly usable as input to the Osage computer, and thereby the need for manual data reduction was eliminated. The printed data list

served for checking the punched output. The channel scanner was sometimes used for the purpose of alternately recording the response function and its integral. The Sanborn six channel recorder was used to record the input noise, the pulsing function and the various system variables for visual inspection and for estimation of errors (uncertainty) in the responses.

Procedure for Simulation Studies and Identification

The simulated systems were programmed on the analog computer (Figure 18) according to Equation (6-1) or (6-2). The setting of the parameters a_{1j} was done by adjustment of the coefficient potentiometers 4-12 and the respective gains G1-G9. The gains were set by proper programming of the input resistors and the feedback capacitors of integrators I1, I2 and I3. Thus, for example, Table 2 shows the proper settings for a third order system that possesses the following natural frequencies:

$$\rho_1 = -0.10607 \text{ [1/sec]}$$

$$\rho_2 = -0.004904 + j \ 0.0136 \text{ [1/sec]}$$

$$\rho_3 = -0.004904 + j \ 0.0136 \text{ [1/sec]}$$

In programming the various systems it was necessary to take into consideration some physical limitations of the analog computer and of the output devices. Since the sampling frequency of the digital output was one reading per second, it was not possible to program systems with time constants smaller than 5 to 10 seconds if a reasonable number of

TABLE 2

EXAMPLE OF POTENTIOMETER AND GAIN SETTINGS
FOR THIRD ORDER SYSTEM*

1j	a_{1j}	Pot. No.	Pot. Setting	Input Resistor	Gain
11	0.014780	4	0.1478	1M	0.1
12	0.085850	6	0.8585	1M	0.1
13	0.006912	5	0.6912	10M	0.01
21	0.000985	12	0.0985	10M	0.01
22	0.008905	11	0.8905	10M	0.01
23	0.003690	10	0.3690	10M	0.01
31	0.088650	7	0.8865	1M	0.1
32	0.099800	8	0.9980	1M	0.1
33	0.092200	9	0.9220	1M	0.1

*All three feedback capacitors were 10 μ f.

samples over the time constant period was desired. On the other hand, it was difficult and undesirable to program very long time constants due to drifts in the analog computer performance and due to the sensitivity of the computer in these ranges of operation. These limitations become quite important under certain testing conditions as will be noted later.

The tests were performed by forcing the different inputs individually (or simultaneously) in order to obtain the necessary independent responses. Usually the forcing

functions were square, trapezoidal or triangular pulses, the duration of which was determined such as to introduce the largest possible amount of energy into the system without overloading the computer. The strength of the responses of the tested outputs ($y_1(t)$) depended largely on the inter-coupling between the variables. Thus for poorly coupled systems, the responses obtained in $y_1(t)$, when $y_2(t)$ or $y_3(t)$ were pulsed, were rather weak. This limitation affected the quality of identification. To overcome such limitations more than one variable was pulsed simultaneously, thereby improving the response characteristics.

For second order systems an attempt was made to generate the integrals of the response functions on the analog computer and to record both the response functions and the integrals simultaneously (on the digital output it was of course necessary to sample both functions alternately). This procedure was successful only for short tests due to the fact that it was necessary to attenuate the integral functions in order not to exceed the operating range of the analog computer and to avoid overloading the amplifiers. For relatively long tests the attenuation factors became too high, thereby reducing the accuracy of the recorded integral functions. Thus for most tests only the response functions were recorded and the integration was performed numerically on the digital computer.

The estimation of error in the response data was done as follows: for tests in which no noise was superimposed,

the uncertainty of the response functions and their integrals was estimated as 1/2 percent of full scale deflection, i.e., 1/2 volt. This value corresponds roughly to the reliability of the analog computer operation. For tests where noise was introduced, the error in the response functions was estimated by measuring only the maximum amplitude of fluctuation of the response to noise, i.e., steady state operation. The estimation of error in the integrals was conducted according to the methods of Chapter IV.

The fundamental matrices were generated on the Osage computer according to the instructions in Chapter V. The final data were fed as input to the computer program (Appendix E) which carried out the identification.

CHAPTER VII

DISCUSSION OF RESULTS

Identification of Second Order Systems with Real Poles

A typical segment of a Sanborn recording of a set of linearly independent response functions and the corresponding forcing pulses of a simulated second order system are shown in Figure 19. The two forcing functions applied were trapezoidal pulses and were used to force $y_1(t)$ and $y_2(t)$ individually, thereby securing linear independence of the responses. The system tested had two real natural frequencies of -0.002 [1/sec] and -0.1 [1/sec], the ratio of which is 50. This ratio of natural frequencies is moderately high, but it is in a range where identification by conventional methods is possible.

The results of identification of this system are shown in Figure 20. In Figure 20a the error prediction as a function of δ is presented, and in Figure 20b the corresponding identification results are shown.

The error prediction corresponds to estimated uncertainty of the data of 0.5 volts (0.5 percent of full scale deviation). This estimate is based on a general knowledge of the accuracy of the analog computer, but is quite arbitrary

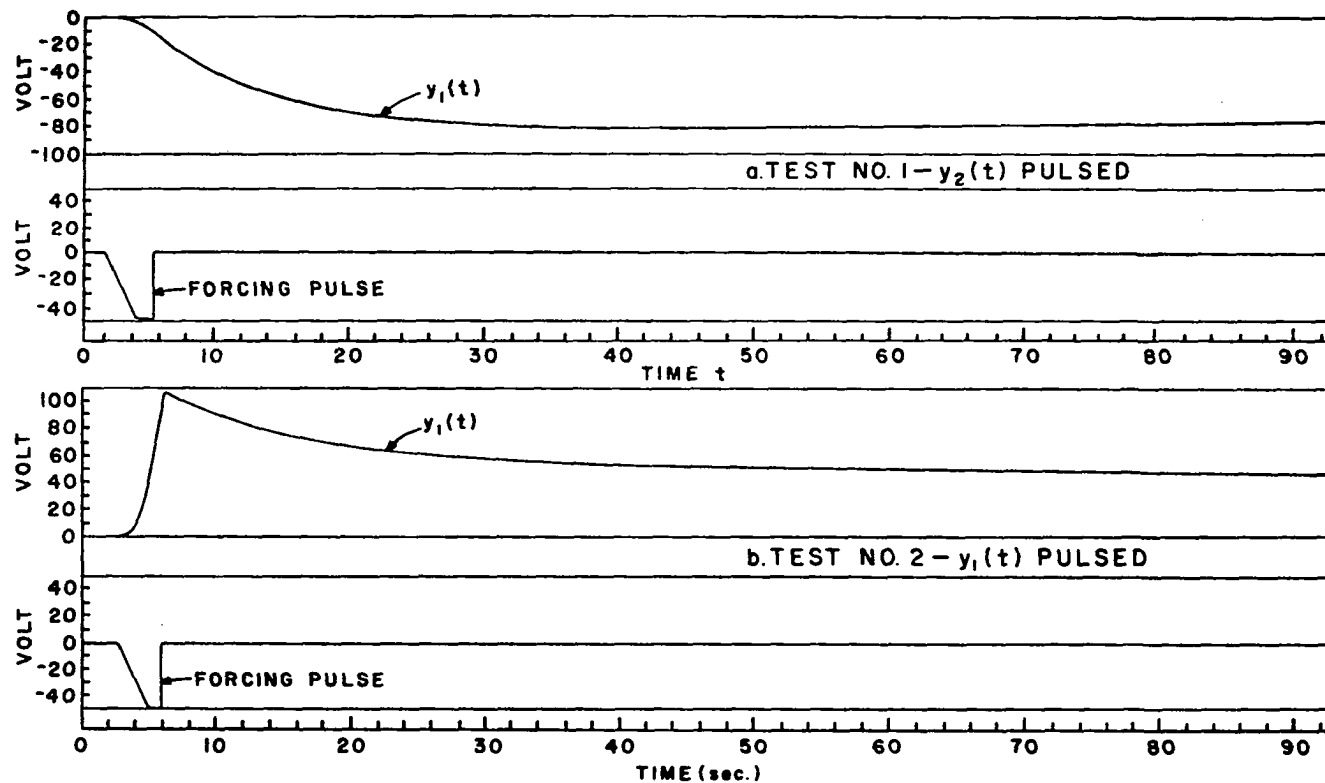


FIGURE 19. TYPICAL SET OF LINEARLY INDEPENDENT RESPONSE RECORDS AND FORCING PULSES FOR IDENTIFICATION OF A SECOND ORDER SYSTEM WITH REAL POLES: $\mathcal{J}_1 = 0.002$ (1/sec.), $\mathcal{J}_2 = 0.1$ (1/sec.). NO SUPERIMPOSED NOISE.

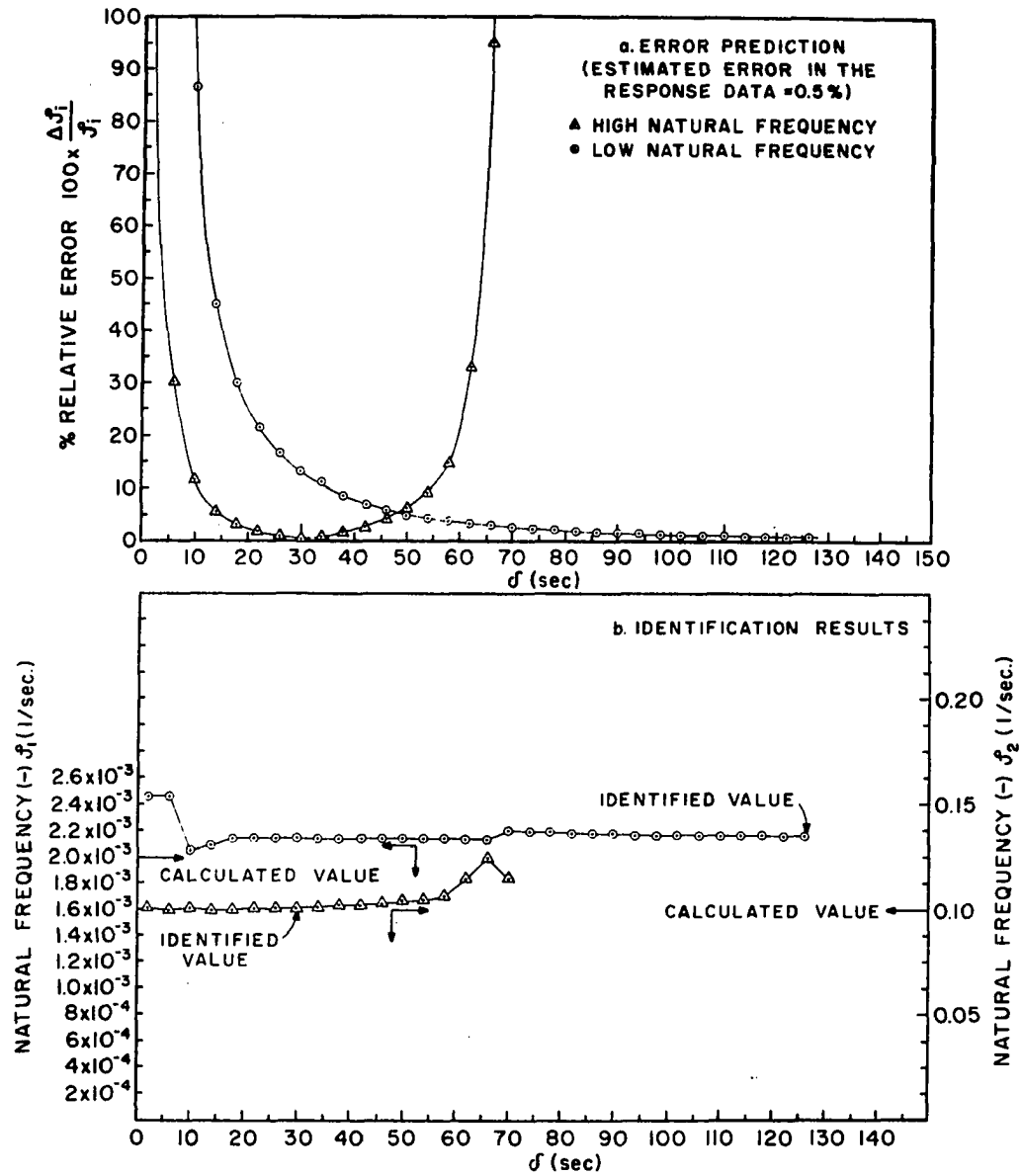


FIGURE 20. IDENTIFICATION OF SECOND ORDER SYSTEM WITH REAL POLES:
 $\dot{f}_1 = -0.002$ (1/sec.), $\dot{f}_2 = -0.1$ (1/sec.). NO SUPERIMPOSED NOISE.

since it does not take into account the uncertainty introduced by errors in the potentiometer settings. This error is difficult to estimate since it depends on factors such as the range of the potentiometer scale used and other external factors which could not be evaluated properly. However, these uncertainties do not affect the response functions, but rather they result in deviations between the system that was actually programmed and the system that was intended to be programmed; i.e., the calculated system. In reality, the estimation of 0.5 percent error in the response functions was not so much intended to predict the actual errors in the identified poles but rather to serve as a tool in the minimization of the predicted errors.

Comparison of the actual error prediction in Figure 20a with the simplified theoretical curves in Figure 5 shows the effect of the term Δv_1 which in Figure 5 was assumed to be constant. According to the simplified Equation (3-54) the error should be minimized when δ is equal to the time constant, which for the high natural frequency corresponds to error minimization at $\delta = 10$ seconds. In actuality the error was minimized at $\delta = 30$ seconds. This difference is a result of the effect of Δv_1 which is in reality a function of δ . As will be seen later this drift of the error minimization to higher values of δ was consistent throughout all the tests.

In the plots of the relative predicted error only the range up to 100 percent was included since beyond this value the identification is totally meaningless. However,

as is obvious from the steep slopes of the error curves, the predicted error rises quickly to values as high as 10^5 percent. This rapid increase is an advantage because it very clearly indicates the error minimization region even when noise is present (as will be seen later).

The predicted relative error curve of the low natural frequency flattens out at relatively low values of δ as compared to the value of the corresponding time constant, which is in agreement with theoretical considerations. This behavior is desirable since it allows proper identification over relatively short records. It also eliminates the need to perform tests that are so long as the highest time constant, particularly when the natural frequencies are widely separated. In effect the identification procedure can be considered as a process of decoupling the various natural frequencies, each of which is identified in the region of δ where its effect is most significant.

The results of the identification, Figure 20b, also illustrate some interesting points. The identification of the high frequency is very accurate and stable at low values of δ , which is the range where the corresponding relative error is minimized. For larger values of δ , where the predicted error increases again, the identification becomes unstable and begins to drift. This difficulty is of course expected since as δ increases the high frequency term becomes relatively insignificant and its reliability doubtful. The identification curve for the high frequency does not extend

beyond $\delta = 70$ seconds since the computation became unreliable.* This phenomenon was present in most of the identification tests when relatively widely separated natural frequencies were present.

The identification of the low natural frequency was essentially stable over the whole region except at very low values of δ . The fact that the identified value of the low natural frequency was not in exact agreement with the calculated value (a deviation of about 9 percent) probably is a result of error in the potentiometer settings on the analog computer and not the identification itself.

In general, the smoothness of the error curves and the stability of the identification (as a function of δ) is seen to be a good criterion for the quality and reliability of the identification. The actual final identification is taken to be the values of the natural frequencies at which the calculated relative errors are minimized. The predicted relative errors of the identified natural frequencies for the system of Figure 20 are less than 1 percent, which in this case is a realistic estimate.

For the identification that was discussed above, t_1 was taken to be equal to δ as was recommended in Chapters III and V. The effect of performing the identification at values of t_{n-1} larger** than δ was also investigated experimentally.

*The argument of a logarithmic term in the computation became negative.

**Smaller values are impossible.

Figure 21 shows the effect of variation of t_1 on the identification of a second order system with real poles $\rho_1 = -0.0646$ [1/sec] and $\rho_2 = -1.9354$ [1/sec]. The predicted relative error is shown as a function of t_1 for various values of δ . As t_1 increases the error rises rapidly to very high values. The exception in the case where $\delta = 1.5$ seconds (where the error decreases when t_1 is increased from 1.5 to 3 seconds) is due to the fact that this particular change causes a relative improvement of the identification. On the other hand the overall results are not improved as can be seen from the dashed line that connects the points for which $\delta = t_1$. The general conclusion confirms the theoretical considerations that t_{n-1} must be invariably selected to be equal to δ . Therefore, throughout the remaining studies, $\delta = t_1$.

Figure 22 shows the identification of a second order system with extremely wide separation between the natural frequencies (ratio of natural frequencies is 1000). The high frequency of such systems can usually not be identified by standard techniques. As before, the identification of the high frequency term is very stable at low values of δ and begins to drift as δ increases beyond the value where the error is minimized. The relative bandwidth of stable identification of this frequency is lower than in Figure 20, which is to be expected. It should be noticed that the predicted error for the high natural frequency is much larger than the actual instability of the identification. The increase is a result of the relatively high sensitivity of

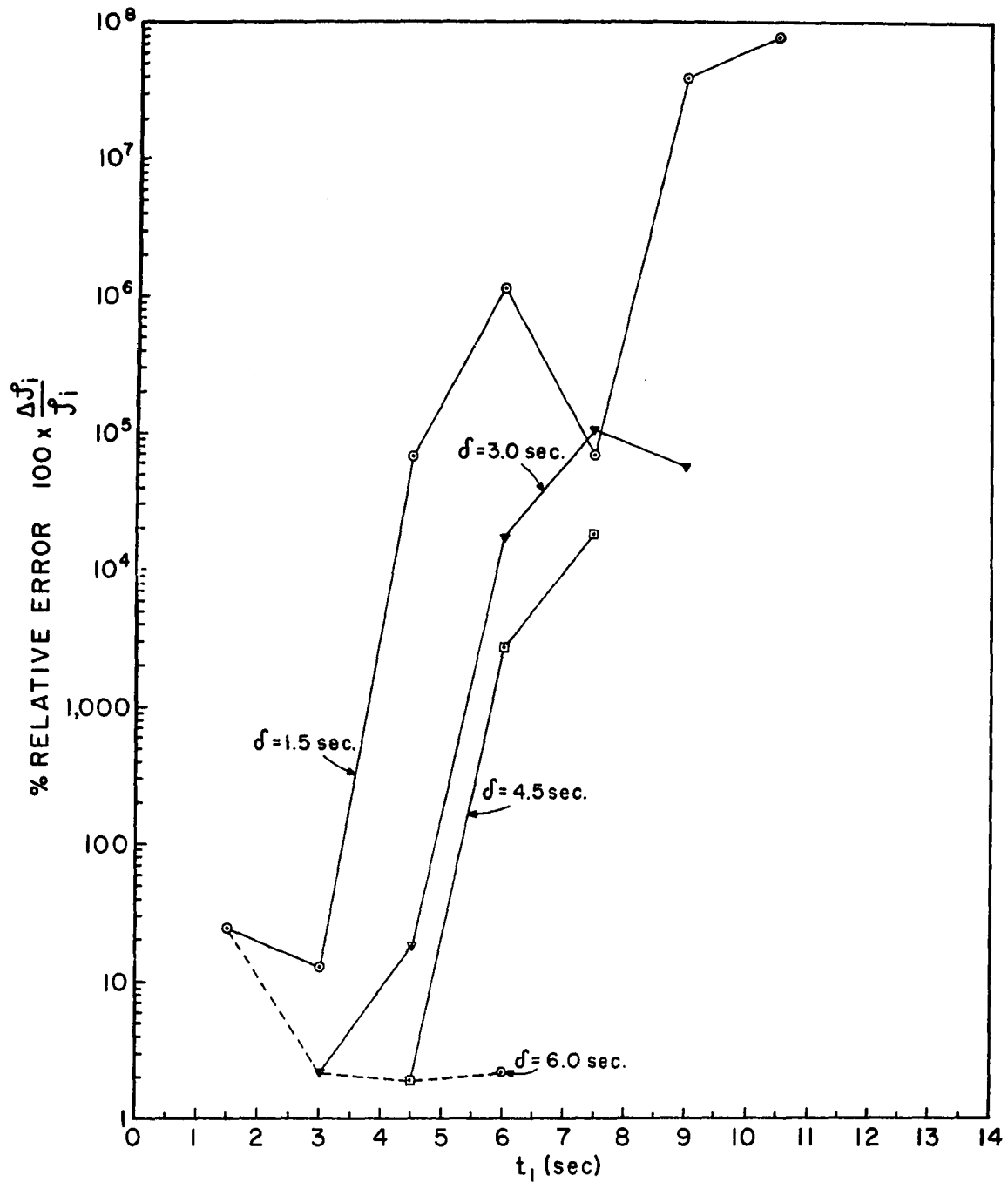


FIGURE 21. EFFECT OF VARIATION OF t_{n-1}
ON QUALITY OF IDENTIFICATION

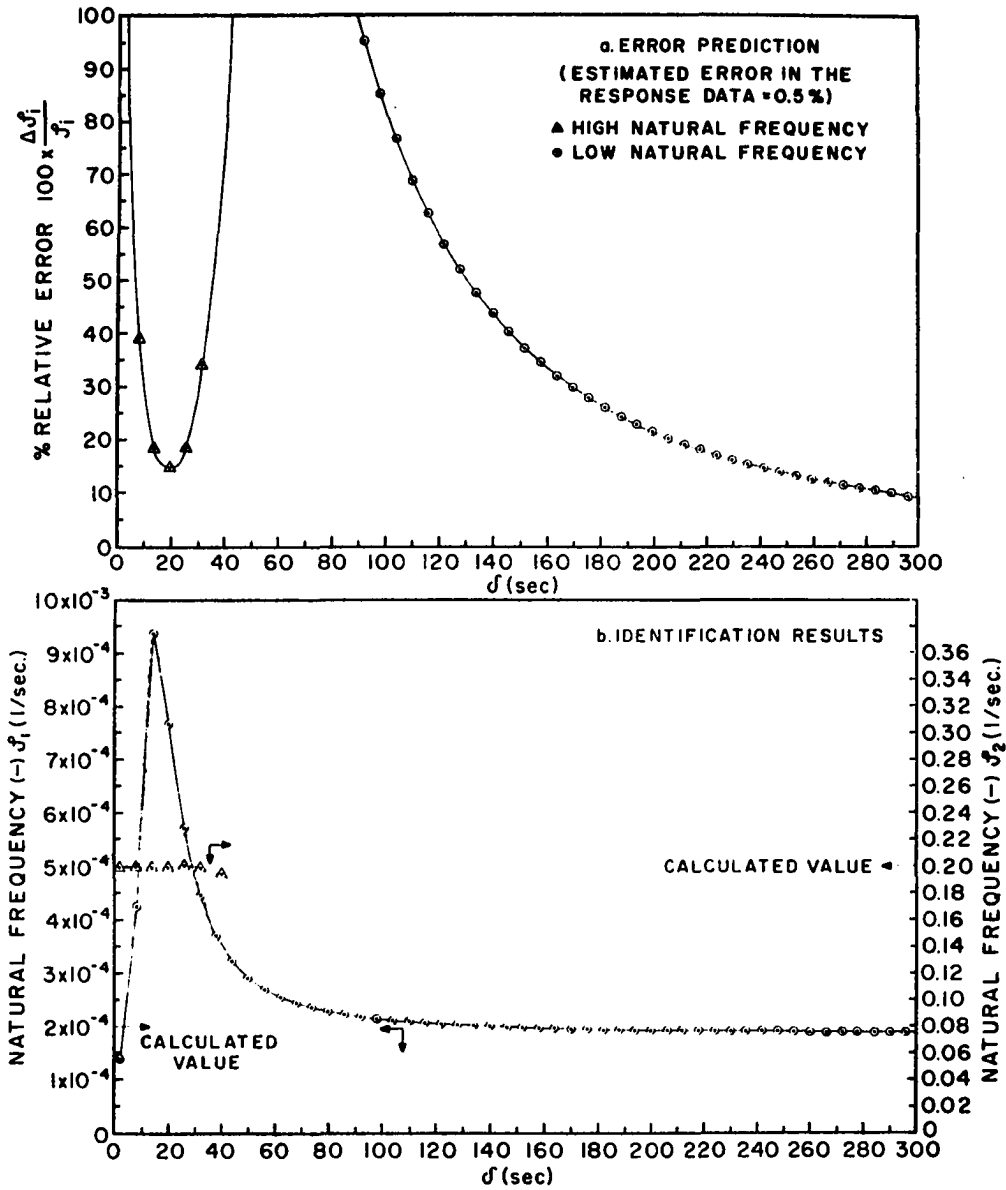


FIGURE 22. IDENTIFICATION OF SECOND ORDER SYSTEM WITH REAL POLES:
 $\zeta_1 = -0.0002$ (1/sec.), $\zeta_2 = -0.2$ (1/sec.). NO SUPERIMPOSED NOISE.

this natural frequency.

The identification of the low natural frequency is unstable at low values of δ . It becomes stable as δ increases and is in agreement with the corresponding error curve. The slight drift in the identified value of the low frequency is more than likely due to drifts in the analog computer operation which became significant when prolonged tests were performed.

Of particular note is the fact that the low natural frequency is properly identified (and the predicted error is relatively low) at values of δ which are considerably lower than the corresponding time constants. As was stated before, this characteristic is of great importance in identification of chemical processes which usually possess very long time constants and respond very slowly.

The identification of a system with very close poles is presented in Figure 23. The sampling frequency was relatively low due to physical limitations of the output equipment (Cf. Chapter VI). Therefore the points in the graph were connected by straight line segments rather than by curves. The identification is stable throughout the tested region (except for very low values of δ at which the identification is always extremely sensitive). The identification of both natural frequencies occurs simultaneously as can be seen from the error minimization characteristics in Figure 23b.

In general the identification of second order systems with real poles was very satisfactory throughout the tested

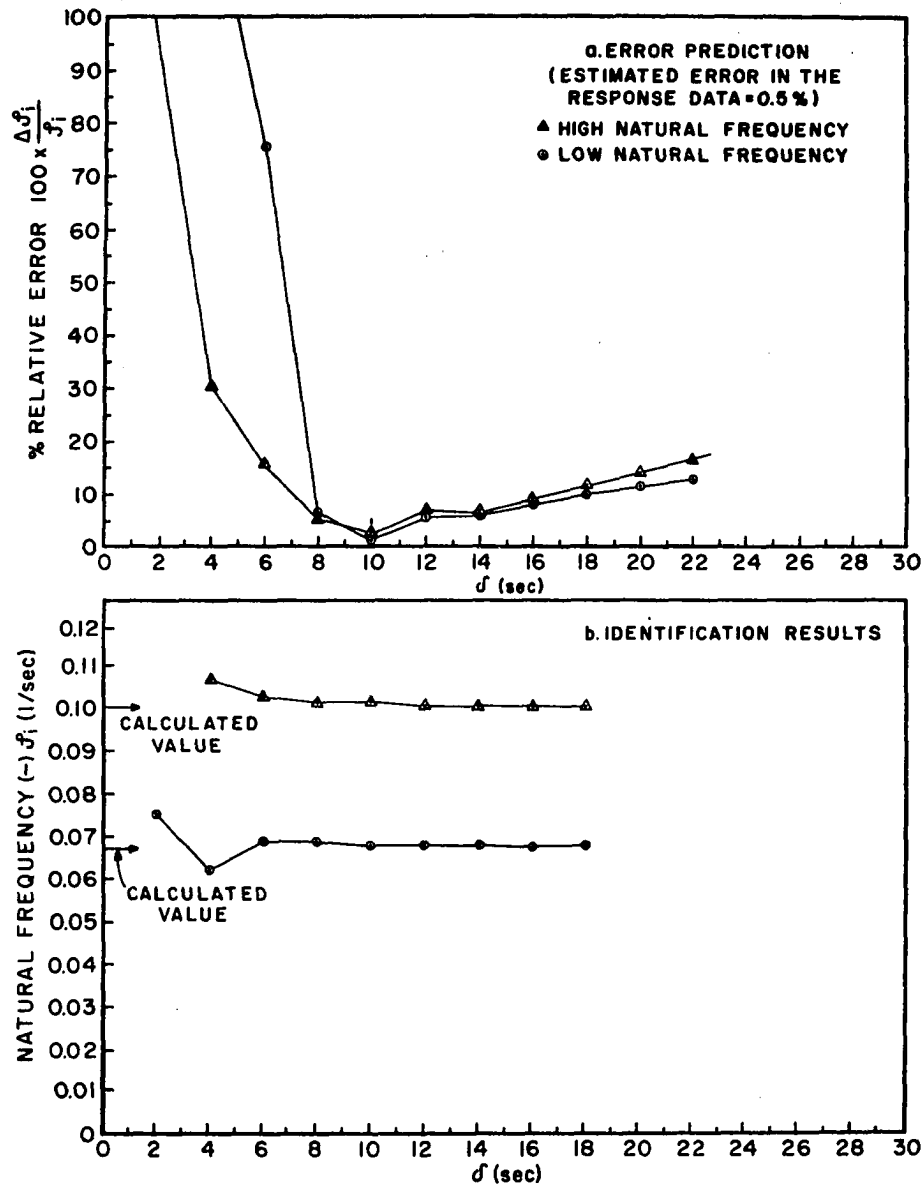


FIGURE 23. IDENTIFICATION OF SECOND ORDER SYSTEM WITH REAL POLES:
 $\hat{f}_1 = -0.0667$ (1/sec.), $\hat{f}_2 = -0.1$ (1/sec.). NO SUPERIMPOSED NOISE.

region and no particular limitations were discovered. A summary of final identification results of a variety of tests is given in Table 3. The fact that there is no consistent pattern in the error prediction as a function of the natural frequency ratios is due to variations in the degree of linear independence and the quality of the responses. This problem of linear independence of the responses will be discussed later in further detail.

Identification of Second Order Systems with Complex Poles

Figure 24 presents a typical set of linearly independent response records and the integrals which were obtained from a simulated second order system with complex poles. The forcing functions which are also shown were narrow triangular pulses. The results of the identification of this system are shown in Figure 25. This system which has a pair of complex poles $p_{1,2} = -0.06 \mp j 0.19078$ possesses a damping factor $\xi = 0.3$ and has very stable characteristics. As can be seen from Figure 25b the identification of both the real and the imaginary parts of the natural frequency is very stable throughout the tested region of δ . The error propagation characteristics of the real part of the natural frequency are quite similar to the case of real poles although its flat region is relatively wide.

Of particular interest is the behavior of the relative error of the imaginary part of the pole. The relative error

TABLE 3
FINAL RESULTS OF IDENTIFICATION OF SECOND ORDER
SYSTEMS WITH REAL POLES*

Test Series No.	Calculated System			Identified System					
	$-p_1$ [1/sec]	$-p_2$ [1/sec]	p_2/p_1	p_1 [1/sec]			p_2 [1/sec]		
				Identified Value	% Relative Predicted Error	δ^{**} [sec]	Identified Value	% Relative Predicted Error	δ^{**} [sec]
1200	0.09429	0.1077	1.14	0.09602	1.729	12	0.1087	2.314	12
800	0.06667	0.1	1.5	0.06742	1.204	10	0.10764	2.545	10
1600	0.05	0.2	4.0	0.0493	0.147	12	0.1997	0.909	4
1700	0.02	0.2	10.0	0.01955	2.246	18	0.199	0.936	20
1800	0.002	0.1	50.0	0.002149	0.71	126	0.1011	0.324	30
1900	0.002	0.2	100.0	0.001878	6.80	60	0.1974	0.248	14
6000	0.0004	0.2	500.0	0.000412	7.37	158	0.1879	7.49	16
6100	0.0002	0.2	1000.0	0.0001885	9.01	288	0.1997	14.6	20

*All identifications were based on estimation of 0.5% error in the response data.

**Value at which the predicted error was minimized and pole evaluated.

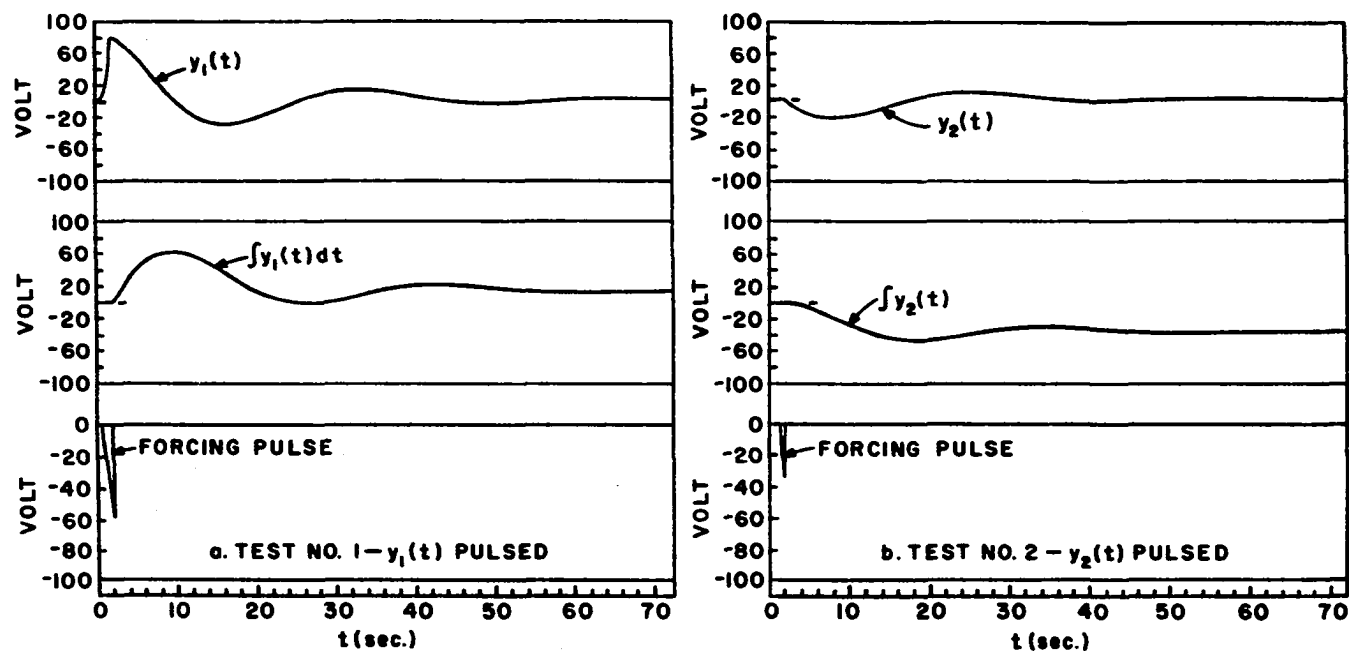


FIGURE 24. TYPICAL SET OF LINEARLY INDEPENDENT RESPONSES AND THEIR INTEGRALS FOR IDENTIFICATION OF SECOND ORDER SYSTEM WITH COMPLEX POLES: $\mathcal{P}_{1,2} = -0.06 \pm j0.19078$ (1/sec.). NO SUPERIMPOSED NOISE.

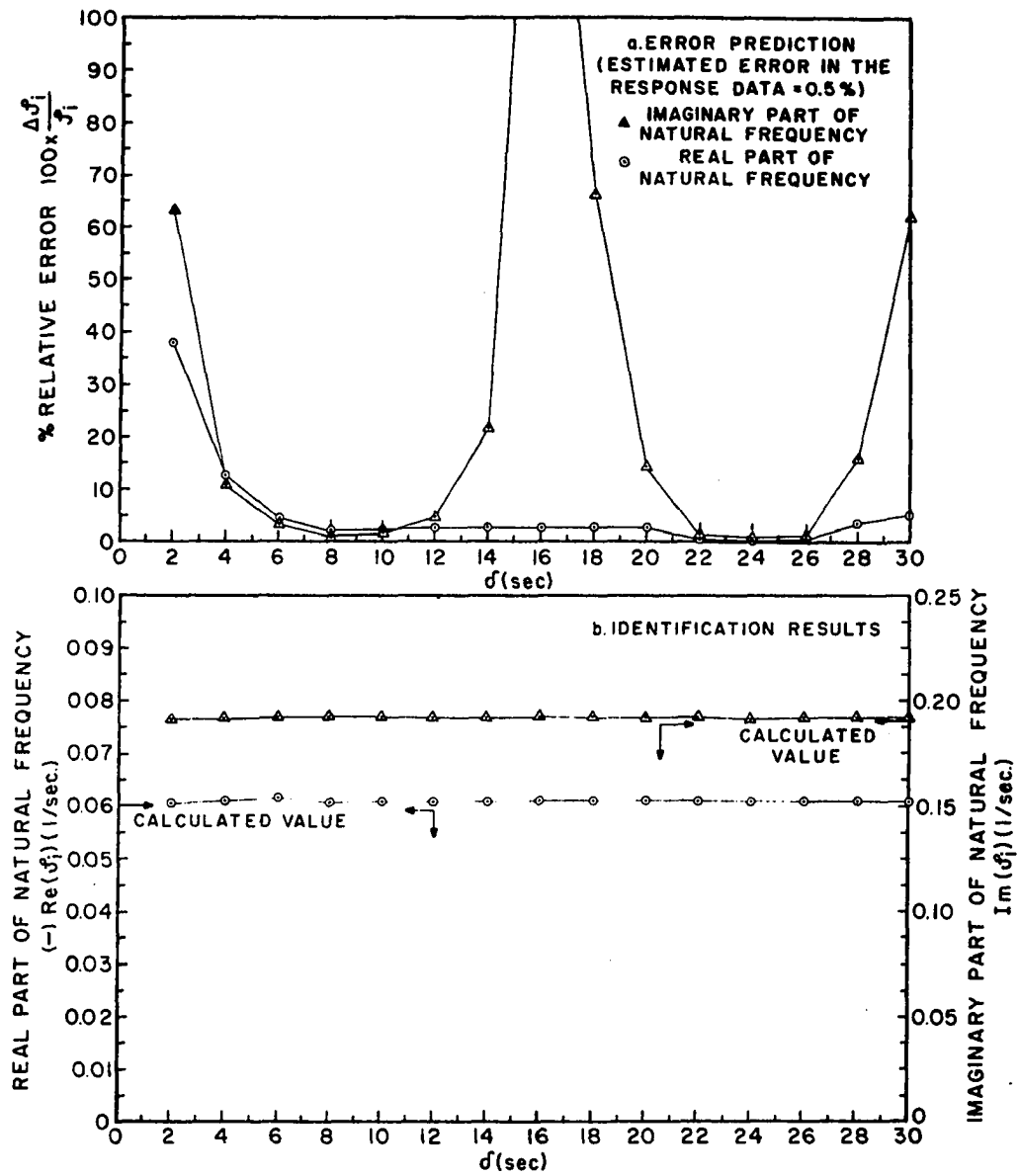


FIGURE 25. IDENTIFICATION OF SECOND ORDER SYSTEM WITH COMPLEX POLES: $j_{1,2} = -0.06 \pm j0.19078$. NO SUPERIMPOSED NOISE.

possesses a periodic oscillation and approaches infinity at intervals when $\delta \text{Im}[\rho_1] = \pi$, a result which is in agreement with the theoretical considerations of the last section of Chapter III. The branches of the error function change in height since the relative error is a ratio between a term that is similar to the error function of real poles and a sinusoidal term (Cf. Equation (3-55)). Therefore the relative error of the imaginary part is minimized at the minimum of the branch closest to the value of δ at which the relative error of the real part is minimized. Although this is difficult to see in Figure 25 due to the flatness of the error curve of the real part of the natural frequency, it is clearly apparent in Figure 26 which shows the error propagation characteristics of a complex system with $\xi = 0.7$. The second branch of the error in Figure 26 is already higher than 100 percent.

As the damping factor increases, the width of the branches of the error curve of $\text{Im}(\rho_1)$ increases. This can be seen by comparison between Figure 25a and Figure 26. When ξ approaches unity the imaginary part of the pole approaches zero and its identification becomes very sensitive. Table 4 shows the identification of a system with two identical poles ($\xi = 1$). It can be noticed that even slight errors in the response records affect the ability of the identification results to indicate either a complex system with a very small and uncertain imaginary part or a real system with two distinct but close natural frequencies. When ξ approaches zero,

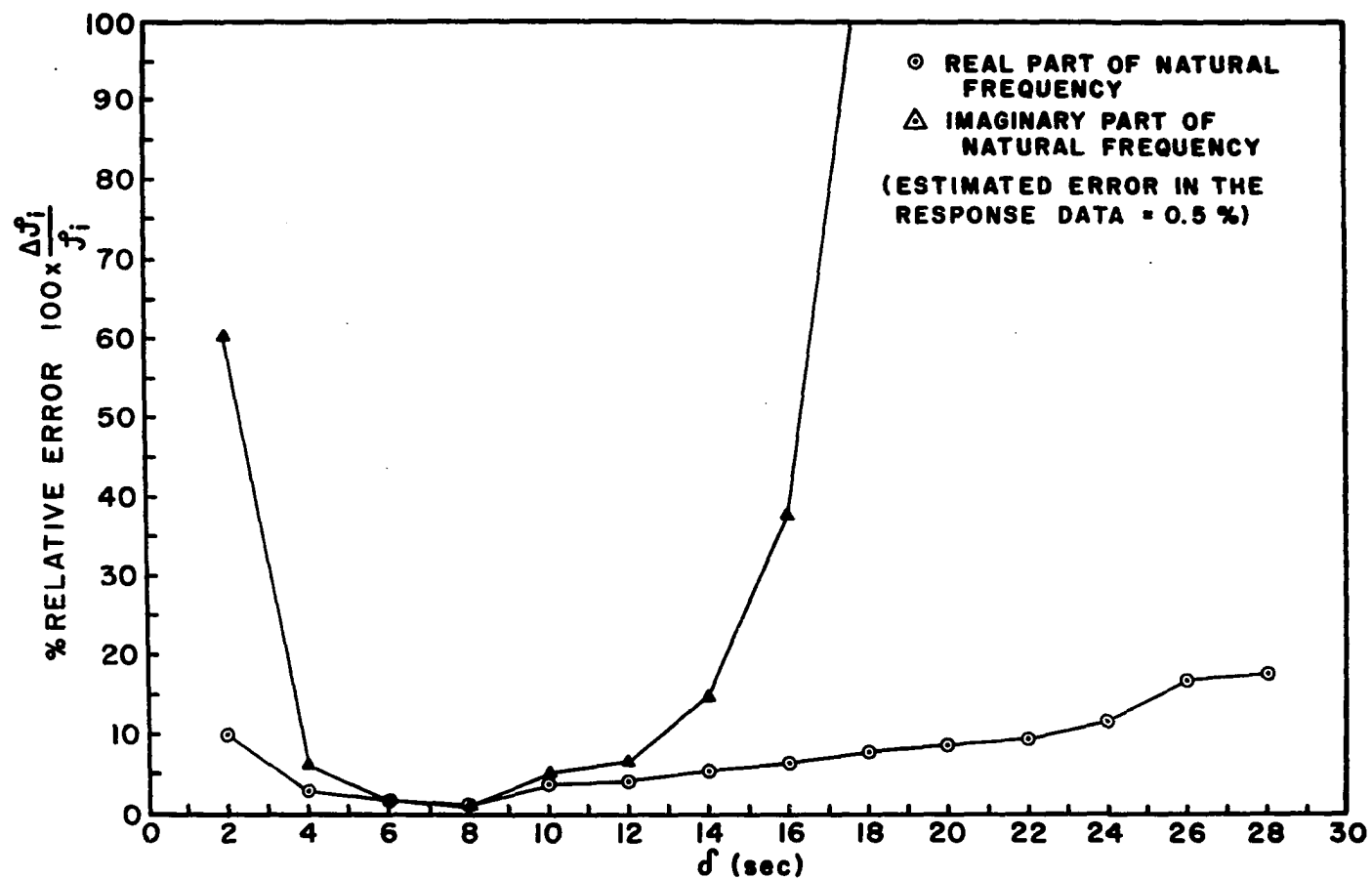


FIGURE 26. ERROR PROPAGATION CHARACTERISTICS OF SECOND ORDER SYSTEM WITH COMPLEX POLES: $\mathcal{J}_{1,2} = -0.14 \pm j0.142828$. NO SUPERIMPOSED NOISE

TABLE 4
IDENTIFICATION OF A SECOND ORDER SYSTEM WITH IDENTICAL
POLES ($\xi = 1$) : $\rho_1 = \rho_2 = 0.1[1/\text{sec}]^*$

$\delta[\text{sec}]$	$-\text{Re}(\rho)$		$\text{Im}(\rho)$		$-\rho_1$		$-\rho_2$	
	Identified Value [1/sec]	% Predicted Error	Identified Value [1/sec]	% Predicted Error	Identified Value [1/sec]	% Predicted Error	Identified Value [1/sec]	% Predicted Error
2	0.1002	48.042	0.01663	1.5×10^4	0.09608 0.09827	154.8 144.7	0.1058 0.1043	130.4 128.7
4	0.1016	19.514	0.0109	7.2×10^3				
6								
8								
10	0.1019	8.737	0.0102	1.4×10^3				
12	0.1019	3.886	0.00974	1.1×10^3				
14	0.1018	0.893	0.00993	97.62				
16	0.1018	1.364	0.0101	104.09				
18	0.1019	1.891	0.0090	151.23				
20	0.1020	2.491	0.00884	169.71				
22	0.1020	3.007	0.00908	179.45				

*Identification was based on estimation of 0.5% error in the response data.

the system response approaches a sustained oscillation. For such systems the identification of the real part of the natural frequency becomes sensitive. It must be emphasized, however, that in these border regions (when ξ approaches unity or zero) the parts that become difficult to identify are insignificant, and their effect can be practically neglected.

In general the identification is stable for both the real and the imaginary parts of the natural frequencies for values of ξ between 0.1-0.95. Table 5 gives a summary of the final results of identification of systems over this range.

Effect of Noise and Error in Response Data
on Quality of Identification

A sample of a high speed recording of the noise function generated by the random square pulse generator is shown in Figure 27a. The mean frequency of this noise function is 13.7 zero-crossings per second. This random noise function was attenuated and introduced into the tested systems to generate the noise in the response functions, Figure 18. Figure 27b and c show the noisy response functions of a second order system with widely separated natural frequencies. The noise level in these response records was estimated as 3.4 volts (i.e. 3.4 percent of full scale deviation). The corresponding identification results of this system are shown in Figure 28. Due to the noise the relative error curve is no longer smooth

TABLE 5
FINAL RESULTS OF IDENTIFICATION OF SECOND ORDER
SYSTEMS WITH COMPLEX POLES*

Test Series No.	Calculated System			Identified System					
	-Re(p1) [1/sec]	Im(p1) [1/sec]	ζ	Re(p1) [1/sec]			Im(p1) [1/sec]		
				Identified Value	% Relative Predicted Error	δ^{**} [sec]	Identified Value	% Relative Predicted Error	δ^{**} [sec]
2100	0.09	0.04358	0.9	0.0908	3.34	20	0.0445	10.8	20
2200	0.16	0.12	0.8	0.1613	4.42	8	0.1208	6.16	8
2300	0.14	0.142828	0.7	0.1408	1.155	8	0.1437	1.004	8
2400	0.12	0.16	0.6	0.1208	2.749	10	0.16126	5.53	10
2500	0.10	0.1732	0.5	0.1008	0.742	6	0.1743	0.454	6
2600	0.08	0.1833	0.4	0.08107	0.493	24	0.1847	0.385	24
2700	0.06	0.19078	0.3	0.0609	0.501	24	0.1913	1.123	24
2800	0.04	0.19596	0.2	0.0405	1.566	14	0.1970	0.573	8
2900	0.02	0.199	0.1	0.0202	0.635	18	0.1994	0.208	6

*All identifications were based on estimation of 0.5% error in the response data.

**Value at which the predicted error was minimized and pole evaluated.

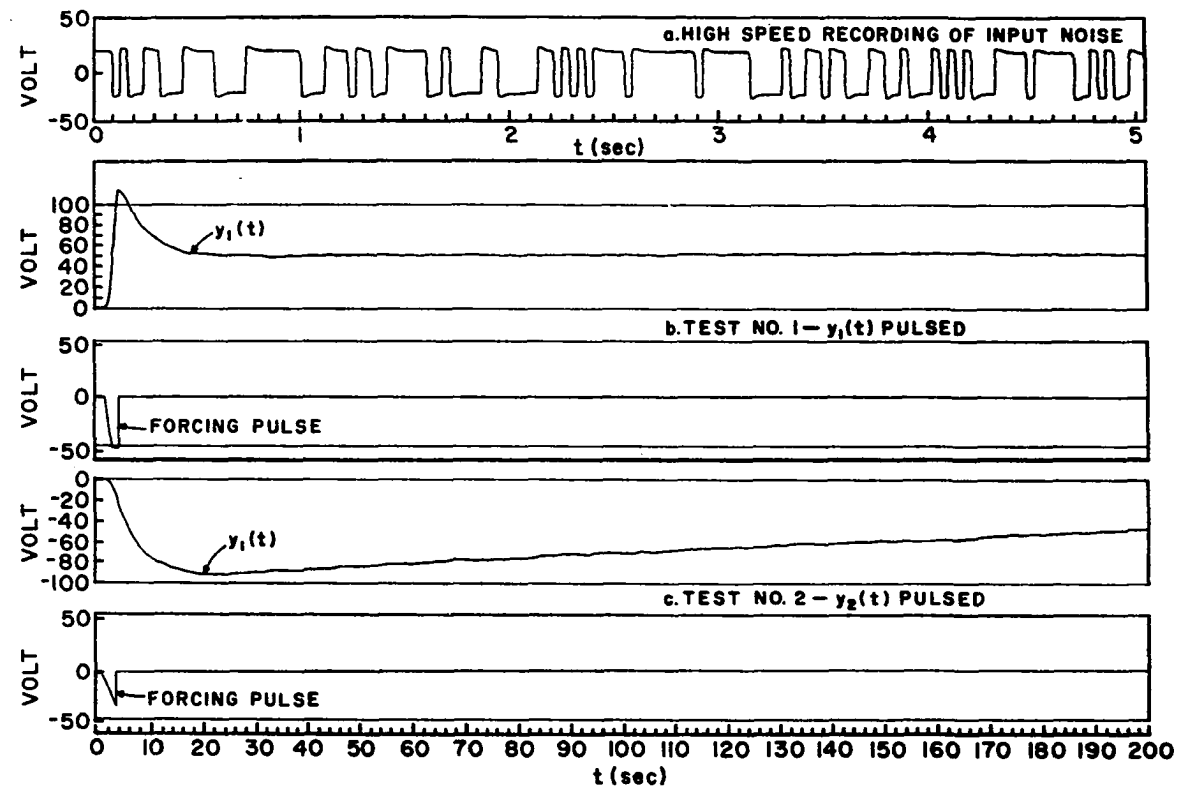


FIGURE 27. a. SAMPLE OF HIGH SPEED NOISE RECORDING, b & c. NOISY RESPONSES OF SECOND ORDER SYSTEM WITH REAL POLES: $\sigma_1 = -0.002$ (1/sec.), $\sigma_2 = -0.2$ (1/sec.) ESTIMATED NOISE LEVEL: 3.4 VOLT (3.4 %).

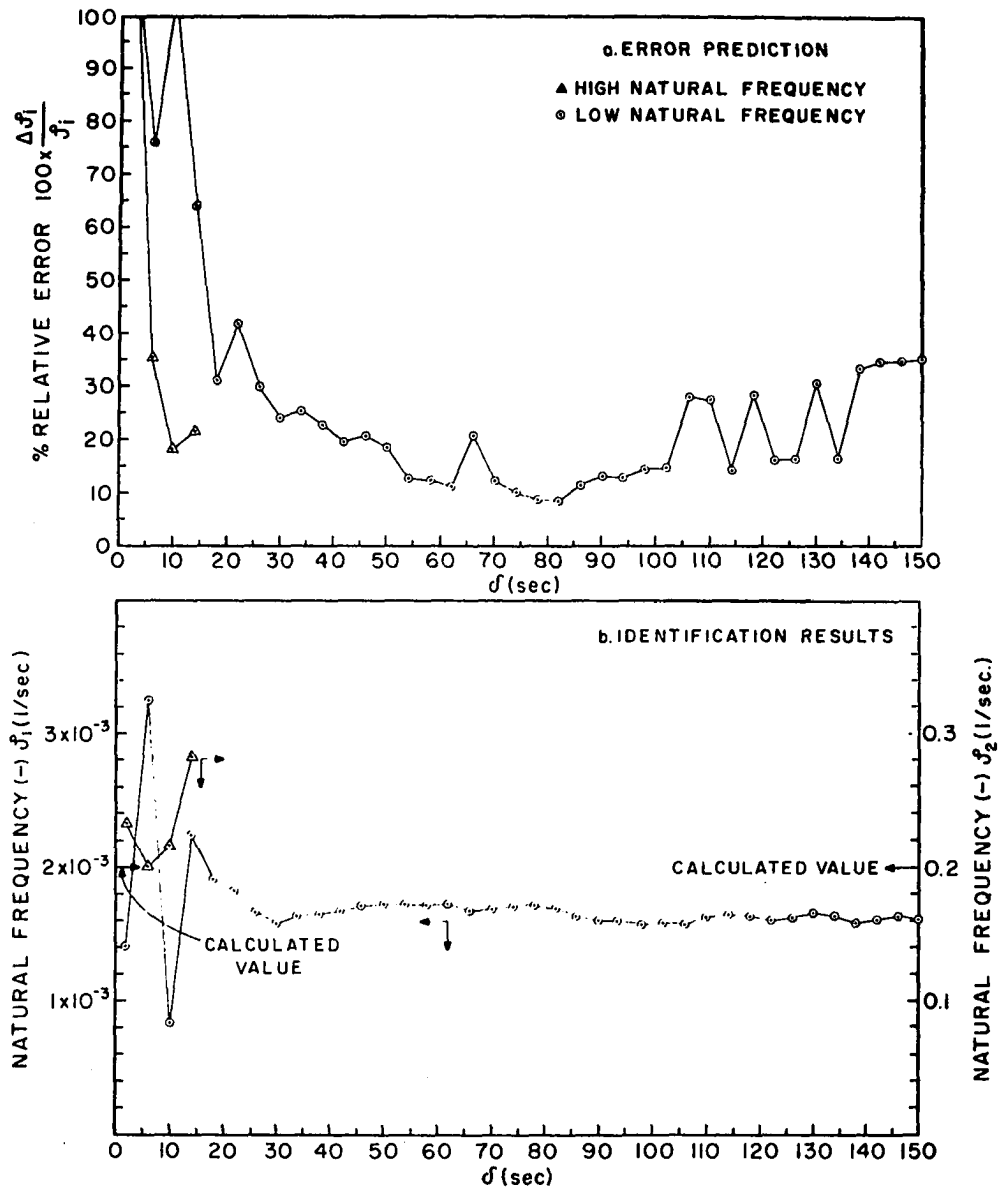


FIGURE 28. IDENTIFICATION OF SECOND ORDER SYSTEM WITH REAL POLES:
 $f_1 = -0.002$ (1/sec.), $f_2 = -0.2$ (1/sec.). ESTIMATED NOISE LEVEL:
 3.4% IN RESPONSES, 2.7% IN INTEGRALS.

but rather possesses random characteristics. However, the error minimization region is still clearly identifiable even for the high frequency. It is important to notice that the bandwidth of stability of the high natural frequency is very narrow because of the noise, indicating the high sensitivity of its identification. This effect is expected since the high natural frequency interacts with the noise and its response is thereby masked. At a high enough noise level there would be no distinguishable difference between the effect of noise and that of the high natural frequency present. Thus for all practical purposes, under such circumstances, the system could be identified at best as first order. Figure 28b shows the results of the identification. It can be seen that the high frequency term possesses a relatively stable value only over the narrow region where the corresponding error is minimized. The stability of identification of the low natural frequency is only slightly affected by the presence of the noise. The fact that the identified value of the low natural frequency deviates by about 9 percent from the calculated value is probably due to errors in potentiometer settings on the analog computer.

The identification of a system with complex poles in the presence of noise is shown in Figure 29. The effect of noise magnitude on the quality of the identification is studied. For comparison the identification results of this system ($\xi = 0.5$) without superimposed noise is also presented. Typical noisy response records are shown in Figure 30.

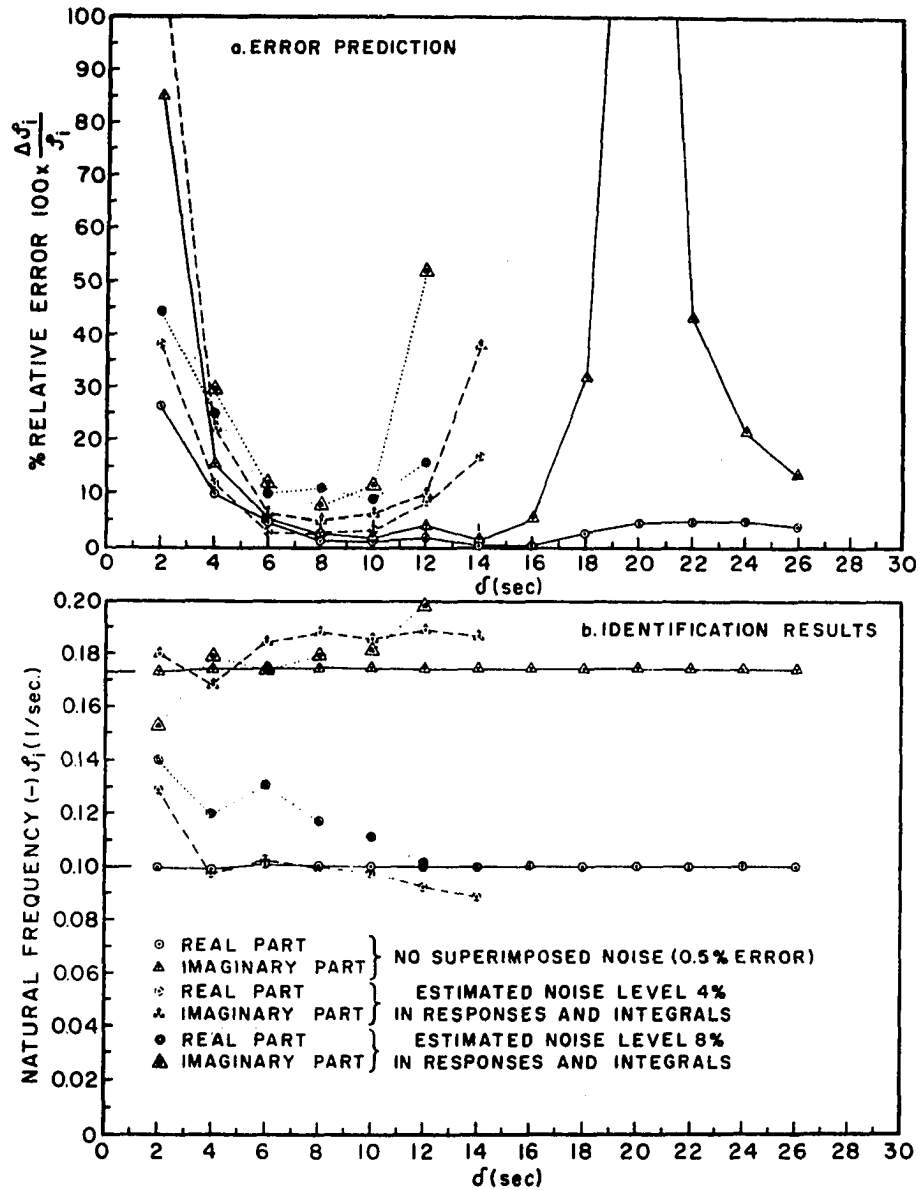


FIGURE 29. IDENTIFICATION OF SECOND ORDER SYSTEM WITH COMPLEX POLES: $\hat{f}_{1,2} = -0.1 \pm j0.1732$ (1/sec.). EFFECT OF NOISE MAGNITUDE ON QUALITY OF IDENTIFICATION.

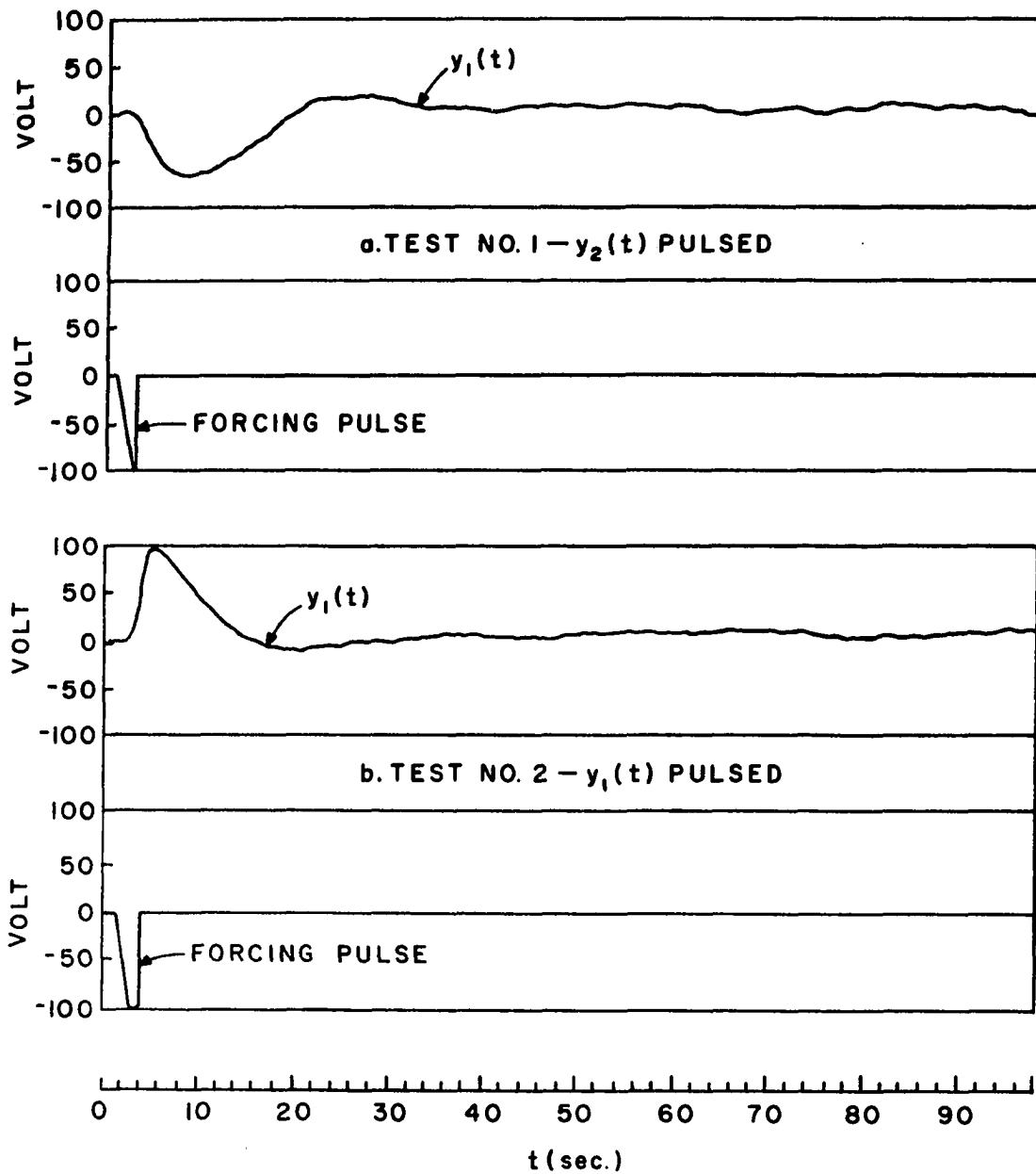


FIGURE 30. NOISY RESPONSES OF SECOND ORDER SYSTEM
WITH COMPLEX POLES: $p_{1,2} = -0.1 \pm j0.1732$ (1/sec.)
ESTIMATED NOISE LEVEL: 8 VOLTS (8%)

The most important point that is seen from Figure 29 is the significant narrowing of the bandwidth of δ over which the system can be identified, due to increased noise: beyond $\delta = 12$ seconds for the system with 8 percent noise and beyond $\delta = 14$ seconds for the system with 4 percent noise, negative logarithmic arguments were obtained in the identification calculations, indicating total unreliability of the response data in these regions. This breakdown can be appreciated if the response records of Figure 30 are consulted: It can be seen that beyond about $t = 15$ seconds little, if any, of the significant response is still present in the recorded signal. Within the identifiable region the error curves are quite stable and well behaved even for the system with 8 percent noise level.

The results of the identification in Figure 29b illustrate an interesting and important phenomenon. For the system with no superimposed noise the identification is extremely stable and the results agree exactly with the calculated values. As noise is introduced into the system, a drift in the identified values (as a function of δ) becomes evident which increases as the noise level is raised. This effect is not a result of the presence of random noise, which accounts only for the fluctuations in the identified values but not for the drift. It is a consequence of the fact that the input noise function was not properly balanced and therefore introduced a bias into the system's response functions. This bias, which, in effect, is a shift in the steady state

was not taken into account and was not subtracted from the responses prior to the integration and the construction of the fundamental matrices. This difficulty illustrates the need for correct knowledge of steady state operating conditions when the present identification method is utilized. However, this phenomenon also suggests the means for elimination of the difficulty: when the steady state conditions are not accurately known, or when biases which are not accounted for are present in the response data, a simple correction procedure can be established for elimination of the error, based on the drift as a criterion.

Effect of Error Estimation or Quality of Identification

Figure 31 shows the effect of variation of the estimates of error bounds for the response data on the minimization characteristics of the predicted relative errors of the natural frequencies. The response functions that were used for this identification were free of superimposed noise. The irregularity in the results of identification of the low natural frequency is apparently due to some disturbance in the analog computer operation during the testing. When the estimated uncertainty of the response data is low, the predicted band of stability of the high natural frequency is very wide. As the assumed error bounds of the data are increased, the predicted bandwidth of stability of the high frequency becomes narrower, and the error minimization shifts toward lower values of δ . This trend is in agreement with

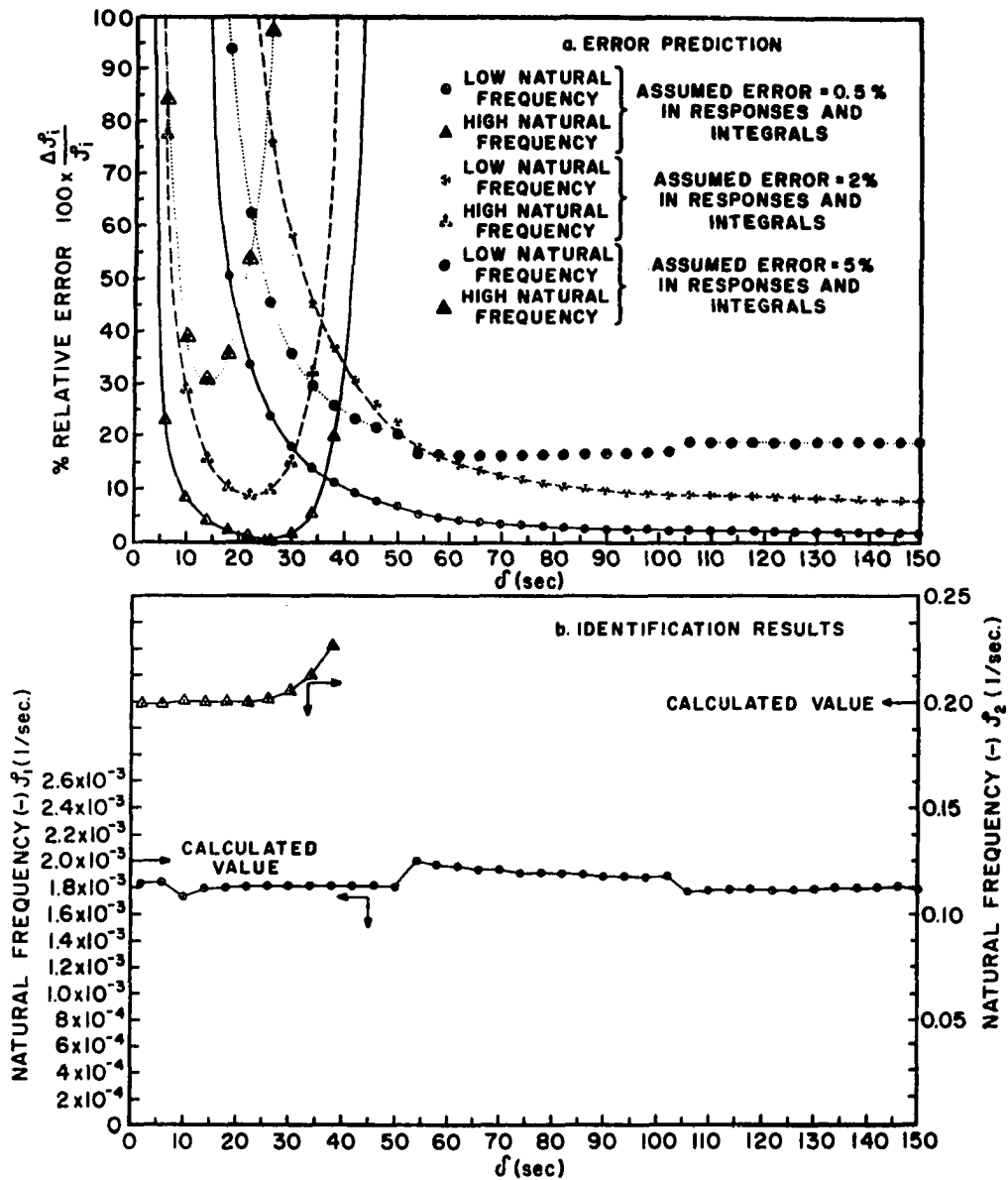


FIGURE 31. IDENTIFICATION OF SECOND ORDER SYSTEM WITH REAL POLES: $\hat{f}_1 = -0.002$ (1/sec.), $\hat{f}_2 = -0.2$ (1/sec.). EFFECT OF ERROR ESTIMATION ON QUALITY OF IDENTIFICATION. NO SUPERIMPOSED NOISE.

observations made in the previous section.

Figure 32 shows the effect of error estimation in the case where noise was present in the response data. When the error was underestimated, the predicted relative error of both natural frequencies was minimized in unstable regions. For the high frequency, the minimization occurred at too high a value of δ and for the low frequency at too low a value of δ , which can lead to erroneous identification. Furthermore, the predicted error level is unrealistically low. When the error was correctly estimated reliable results were obtained. Overestimation of the error results in safe evaluation of the location of error minimization but drastically exaggerates prediction of the uncertainties in the results. Thus the relative error of the high frequency in Figure 32a for the high error estimation is never below 100 percent. This result is another form of erroneous identification since it implies that the high natural frequency does not have any significant effect in the presence of the noise, which is not correct.

It can be concluded that as the actual noise level increases, it becomes more and more important to have good error estimates if reliable identification is desired. Overestimation of the error is usually safer than, and superior to, underestimation.

Methods for Securing Linear Independent Responses and Effects of Weak Linear Independence

The effect of weak linear independence of the

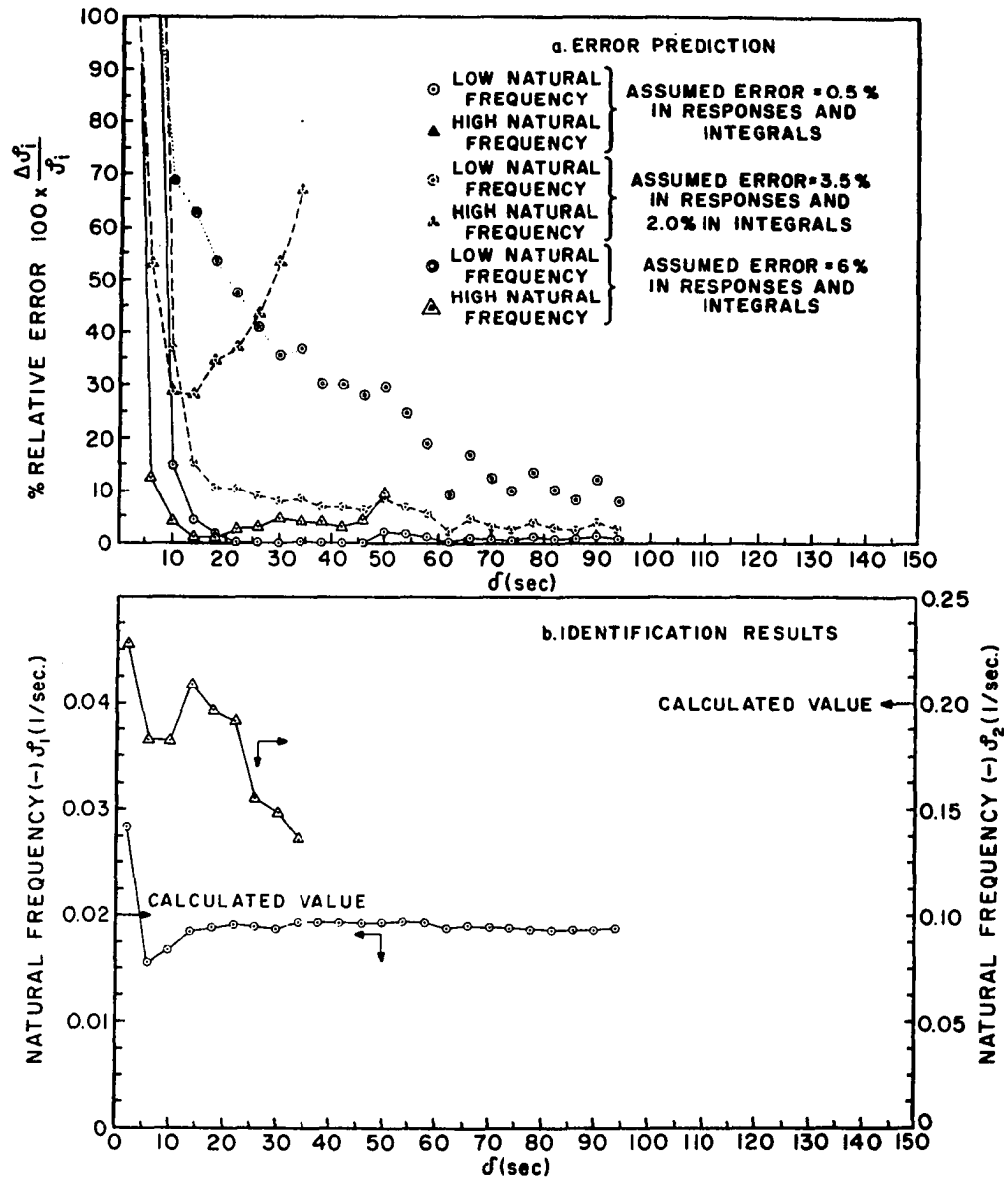


FIGURE 32. IDENTIFICATION OF SECOND ORDER SYSTEM WITH REAL POLES' $\hat{J}_1 = -0.02$ (1/sec.), $\hat{J}_2 = -0.2$ (1/sec.). EFFECT OF ERROR ESTIMATION ON QUALITY OF IDENTIFICATION. ESTIMATED NOISE LEVEL: 3.5% IN RESPONSES, 2% IN INTEGRALS.

response functions is very similar to the effect of error. Although weak linear independence does not introduce random characteristics into the results, it magnifies the effect of existing error. From a theoretical viewpoint linear independence can be secured by forcing a certain input variable with pulses that contain different spectral density distributions. Practical experimentation with this method failed to provide the required responses and no significant linear independence was obtained, even when virtually no noise was present in the system. As noted previously, the method that was successfully applied for obtaining linear independent responses was forcing a different variable, or combination of variables, each time.

Effect of Utilization of Higher Order $\bar{\Phi}$ Matrices

The results that were obtained by utilization of fundamental matrices $\bar{\Phi}$ of order higher than $n-1$ in order to improve the results of identification were usually unfavorable.

Although the noise effects were considerably reduced through the process of integration, it appeared that unaccountable errors such as bias and steady state deviations were considerably amplified through the higher order integration process. This technique is therefore not recommended.

Identification of Third Order Systems

The results of identification of some representative third order systems are illustrated in Figures 33, 34 and 35.

Figure 33 shows the results of identification of a

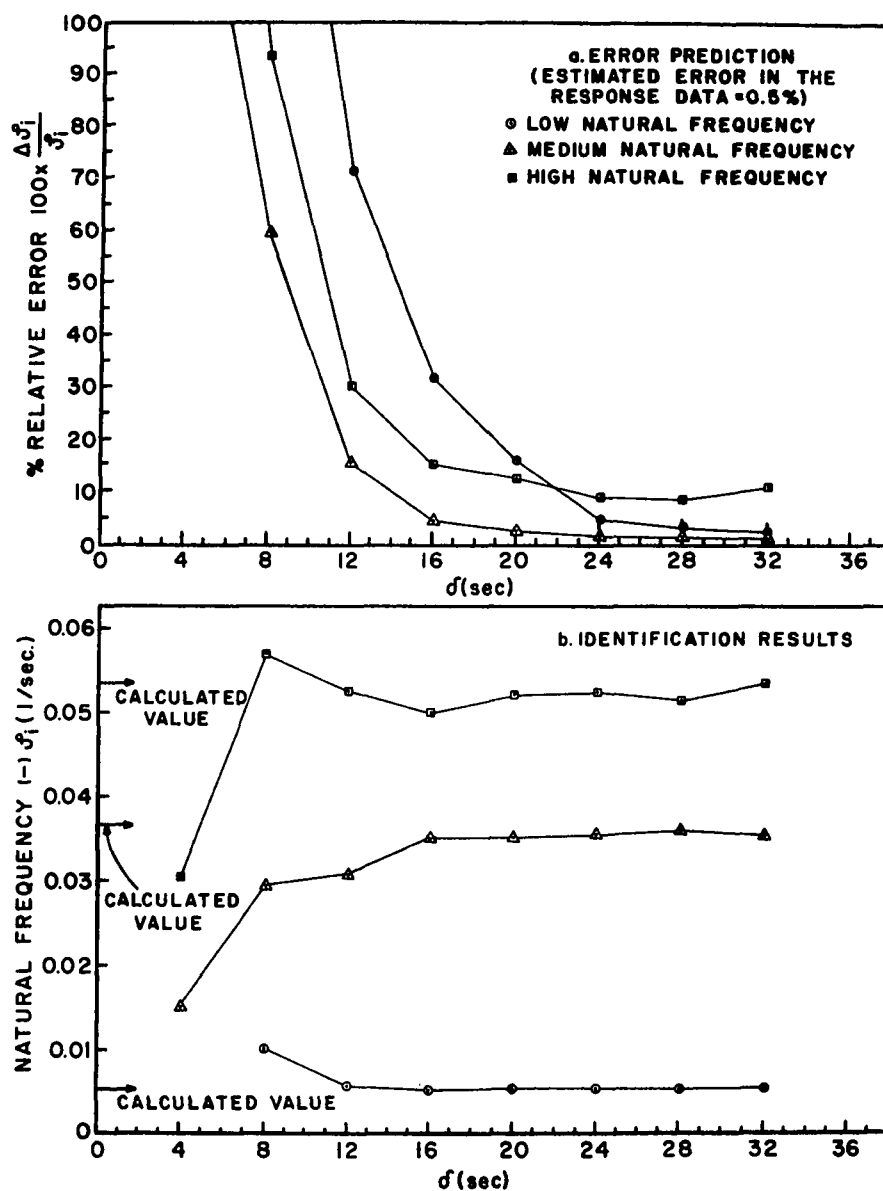


FIGURE 33. IDENTIFICATION OF THIRD ORDER SYSTEM WITH REAL POLES: $\hat{\sigma}_1 = -0.0054027$ (1/sec.), $\hat{\sigma}_2 = -0.036722$ (1/sec.), $\hat{\sigma}_3 = -0.053567$ (1/sec.). NO SUPERIMPOSED NOISE.

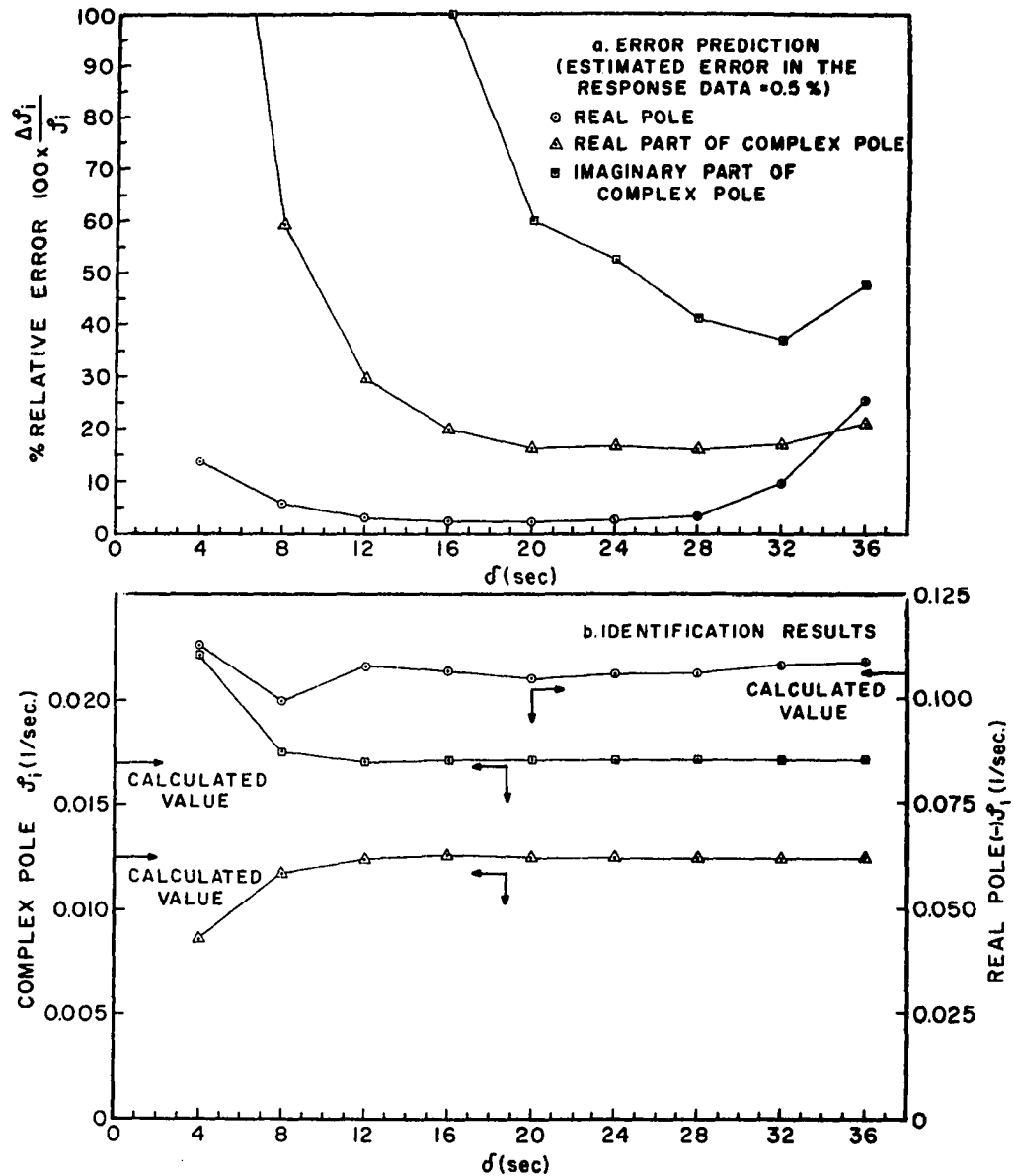


FIGURE 34. IDENTIFICATION OF THIRD ORDER SYSTEM WITH ONE
REAL POLE: $\hat{s}_1 = -0.10636$ (1/sec.) AND A PAIR OF COMPLEX POLES:
 $\hat{s}_{2,3} = -0.012517 \pm j0.01705$. NO SUPERIMPOSED NOISE.

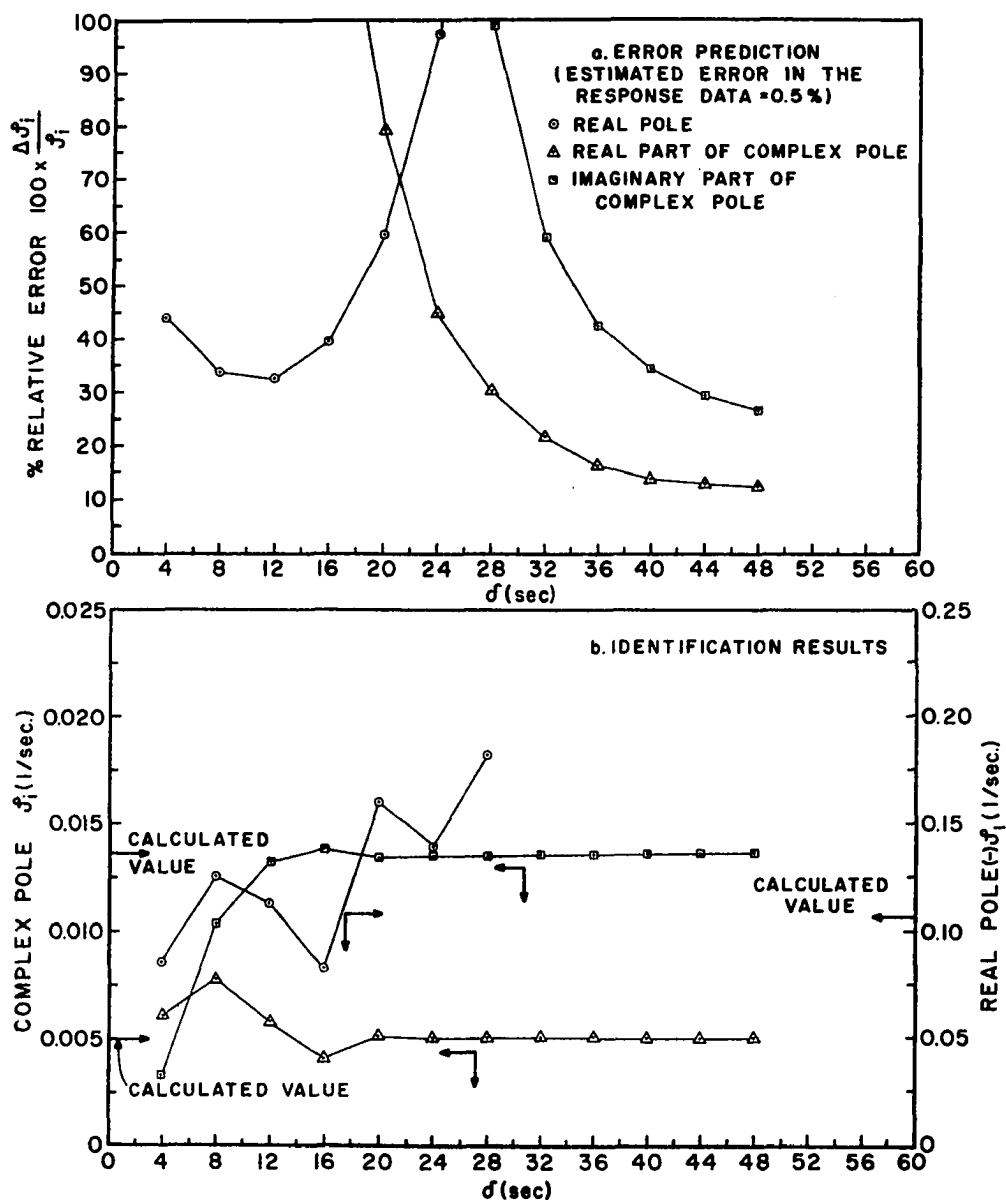


FIGURE 35. IDENTIFICATION OF THIRD ORDER SYSTEM WITH ONE REAL POLE: $\hat{p}_1 = +0.10607$ (1/sec.) AND A PAIR OF COMPLEX POLES: $\hat{p}_{2,3} = -0.00490459 \pm j0.01361$ (1/sec.). NO SUPERIMPOSED NOISE.

system with three moderately separated real poles. The identification is generally stable except for the range of low δ where some deviation is evident. The error minimization characteristics are quite stable and properly behaved.

Figure 34 illustrates the results of identification of a system with one real pole and a pair of complex poles. The complex pole has a damping factor $\xi = 0.531$ and the ratio between the real pole and the real part of the complex pole is 8.5. Figure 34b indicates very stable identification results as expected for such a system.

The results of identification of a rather sensitive third order system are illustrated in Figure 35. This system possesses one real pole and a pair of complex poles. The damping factor ξ of the complex pole is 0.341 and the ratio between the real pole and the real part of the complex pole is 23. It is evident from Figure 35b that the identification of the real pole becomes quite difficult and unstable due to its relatively high frequency. The error prediction in Figure 35a is in accord with the theoretical considerations.

Several general observations can be made concerning the identification of third (and higher) order systems:

- (1) The deviation between the predicted and real errors becomes exceedingly large as the system's order increases. This behavior is evident in Figures 34a and 35a where the predicted errors are noticeably higher than the calculated deviations. The reason for this phenomenon is the fact that the error propagation analysis evaluates the extreme

bounds of errors, i.e. the largest possible errors, rather than the expected (most likely) deviations. As the order of the system increases, the probability of attaining the extreme error bounds becomes exceedingly low because of the increased number of permutations of matrix configurations. In fact, the proportionality is $1/n^2$.

It might be suggested that it would be desirable to perform an error analysis based on the expected deviations. This approach is questionable if the increased complexity and computational requirements are taken into consideration.

(2) The generation of responses that are sufficiently independent (with regard to the noise level) becomes more difficult as the order of the system is increased. This statement is quite obvious and the equivalent of its implications is true for any identification method.

(3) Computational sensitivity may become a factor to be considered when higher order systems are to be identified.

(4) The largest interval δ over which the identification can be performed is $(1/n)^{\text{th}}$ of the response record length. Thus for third order systems δ can extend to one third of the record length (as compared to one half for second order systems). In other words longer response records are required for identification of higher order systems. Thus the relative bandwidth of stability for the high frequencies becomes narrower and their identification becomes more sensitive.

In spite of the above remarks the results of the

identification of third order systems are very satisfactory and indicate favorable possibilities for identification of even higher order models.

The Problem of Misidentification

Figure 36 illustrates the results of an attempt to identify a third order model for a second order system. Thus, three response functions were used for the identification and three natural frequencies were evaluated.

The two truly existing natural frequencies ($\rho_1 = -0.02$ [1/sec], $\rho_2 = -0.2$ [1/sec]) are correctly identified and the results are very stable. The predicted relative error (not illustrated) behaves properly and in accordance with theory.

The third erroneous natural frequency is obtained as a result of errors present in the responses (which is the only reason for the responses not being absolutely linearly dependent). This apparent natural frequency is very unstable and oscillates strongly, as a function of δ , even toward positive values which would imply the system's instability! The predicted relative error for this natural frequency was in the order of 10^5 percent throughout the tested range. Consequently this natural frequency would be rejected if any reasonable acceptance criterion were established. Thus, for all practical purposes, this system is in effect identified as being of second order.

The error analysis, being a very sensitive gauge of

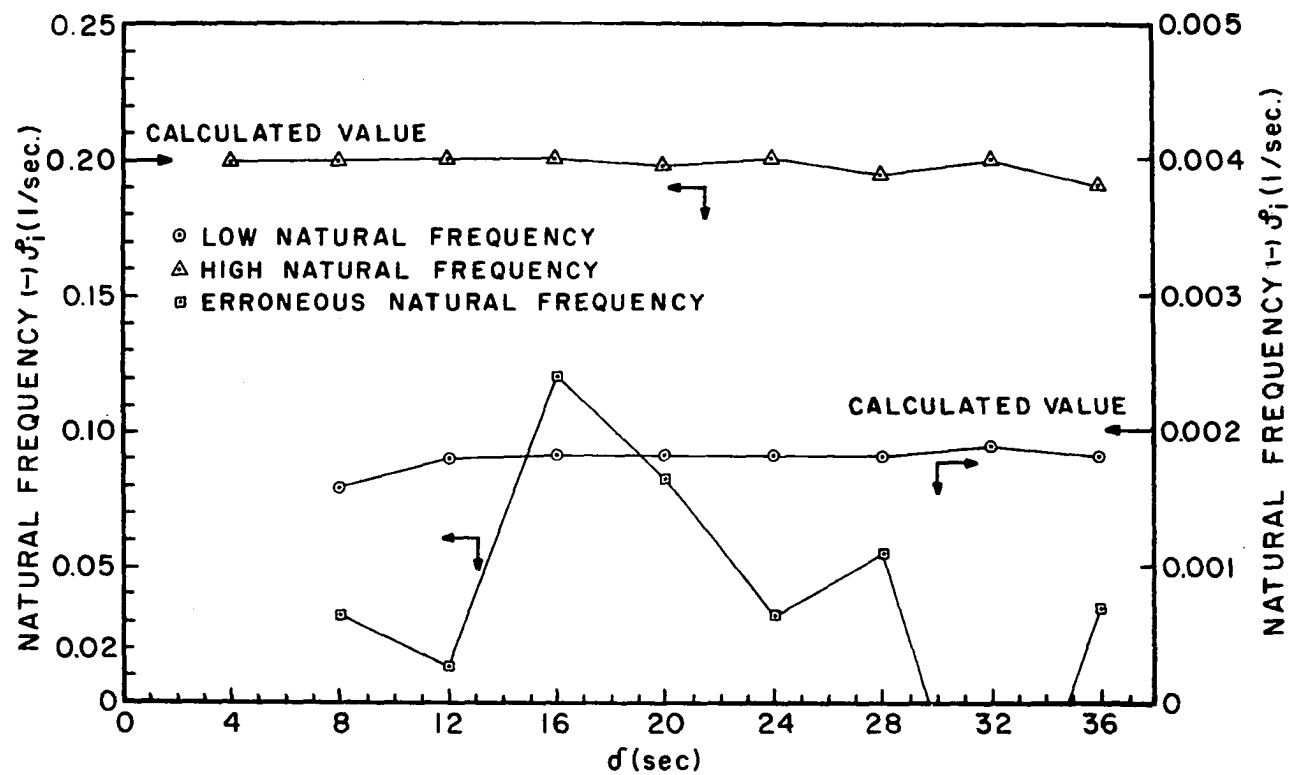


FIGURE 36. IDENTIFICATION OF A THIRD ORDER MODEL FOR A SECOND ORDER SYSTEM.

the reliability of the identification, indicates whenever difficulties in the identification are present. The cause for difficulty of identification is invariably weak, linear independence between the responses. The error analysis indicates which of the natural frequencies is poorly presented in these responses and thereby, implicitly, determines the effective order of the system (as evidenced through these particular responses). Since the reason for this weak linear independence cannot be determined directly through the identification, engineering judgment has to be applied.

Identification of Actual Chemical Plant Test Data

Actual chemical plant test data were obtained for identification by the time domain technique. These test data of an operating plant were originally generated by a petroleum company for the purpose of analysis and identification by the pulse testing and Fourier transformation method. The response data and the respective forcing pulses are presented in Figure 37.

Before presenting the results of the identification, several remarks pertaining to the test data are in order:

(1) No particular effort was made to generate linearly independent response functions since for the Fourier transformation identification method only one test is required. Although the responses are satisfactory from this aspect they were obtained by simply forcing the different inputs without any particular emphasis on linear independence.

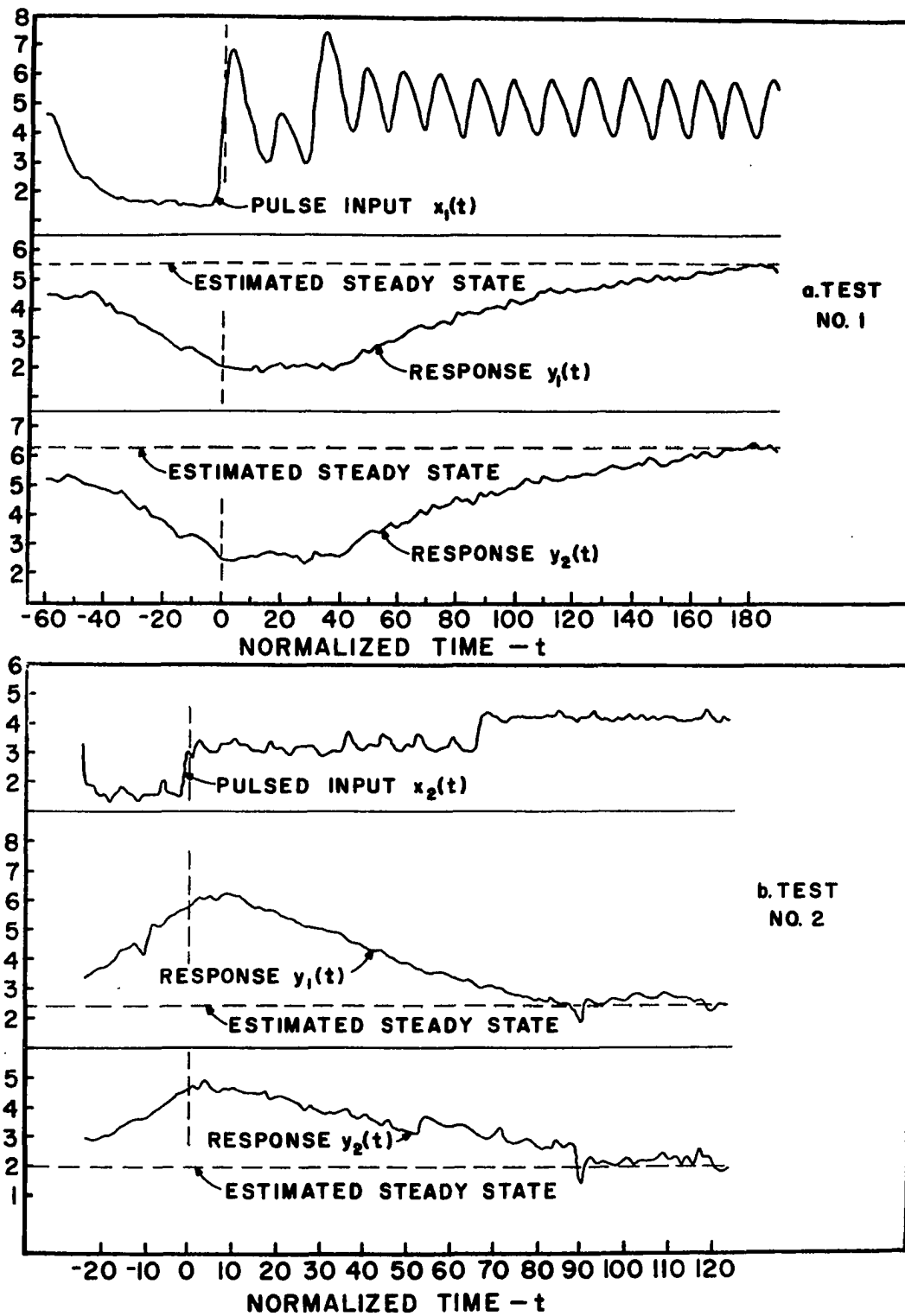


FIGURE 37. PLANT PULSE TEST DATA

(2) The system possesses pronounced nonlinearities as can be observed by inspection of the response functions. Of particular note is the dead-time (or time delay) which is clearly recognized in the responses of test no. 2 in Figure 37. A definite change of slope is evidenced in both response functions $y_1(t)$ and $y_2(t)$, which occurs at $t = 90$ and is the result of a sudden change in the input level which occurs at $t = 65$. Thus, there is a dead time of about 25 time units.

(3) The strong oscillations of the input function in test no. 1 do not (visually) appear to affect the response characteristics. The apparent reason is the relatively high oscillation frequency as compared to the plant's natural frequencies. The change in the operation level of the input of test no. 2 has a profound effect on the response functions, and as a consequence (as will be seen later) also a harmful effect on the identification results.

(4) No steady state data were available. The steady state operating levels (indicated by the dashed lines in Figure 37) were estimated by a trial and error identification procedure where an effort was made to eliminate drifts in the identification results (cf. discussion on page 178).

(5) The noise level in the response functions is very high and was estimated as 20 percent of the maximum amplitude deviation.

(6) No attempts were made to smooth the response records. The identification was made based on readings off the actual response functions.

(7) The order of the system was not known a priori. However, since only two sets of test data were available, a second order model was the highest possible order or identification. As the identification results prove, second order identification was satisfactory.

Two sets of identification tests were made: one was based on the responses of $y_1(t)$ and the other on the responses of $y_2(t)$. Since only the general model was sought, both identification test sets were expected to yield the same results, i.e., the system's poles.

The results of the identification are presented in Tables 6 and 7.

In Table 6 the results of the identification based on the response records of $y_1(t)$ are listed. At low values of δ two real poles are identified, which are quite unreliable as indicated both by the wide fluctuations in the identified values (as a function of δ) and by the respective error predictions. The results become more stable, however, as δ increases. As δ reaches 18 TU (time units) a pair of complex poles begins to appear, which has a real part of about 0.0105 [1/TU] and a damping factor of about 0.71. This identification is steadily obtained, (within $\pm$ 30 percent), as δ increases up to 42 TU.

It should be noted that the identification of the imaginary part of the pole is less reliable than the real part, which is evident both from the identified values and from the predicted relative errors. The error for the real

TABLE 6
IDENTIFICATION OF CHEMICAL PLANT TEST DATA
(BASED ON RESPONSE RECORDS OF $y_1(t)$ IN FIGURE 37)*

$\delta[TU]**$	$-\text{Re}(\rho)$		$\text{Im}(\rho)$		$-\rho_1$		$-\rho_2$	
	Identified Value [1/TU]	Per Cent Predicted Error	Identified Value [1/TU]	Per Cent Predicted Error	Identified Value [1/TU]	Per Cent Predicted Error	Identified Value [1/TU]	Per Cent Predicted Error
2					0.02154	387.0	0.84186	4739.5
6					0.01079	70.8	0.1289	267.4
10					0.01459	89.3	0.1019	189.7
14					0.01882	73.0	0.07799	87.3
18	0.01047	36.0	0.0129	216.7				
22	0.007015	39.2	0.01646	77.6				
26	0.01335	32.6	0.0167	48.3				
30	0.01052	20.4	0.00954	72.9				
34	0.01078	11.6	0.0170	43.2				
38	0.01052	7.75	0.011	39.2				
42	0.01319	13.9	0.0107	31.6				
46	0.02032	8.74	0.0124	36.5				
50	0.02172	9.35	0.0105	52.5				
54	0.0185	11.73	0.01455	93.2				

*Error prediction is based on estimation of error bounds of 20 per cent (of full scale deviation) in the response functions and 25 per cent in the integrals.

**TU are normalized time units used in the response records in Figure 37.

TABLE 7
IDENTIFICATION OF CHEMICAL PLANT TEST DATA
(BASED ON RESPONSE RECORDS OF $y_2(t)$ IN FIGURE 37)*

δ [TU]**	$-\text{Re}(\rho)$		$\text{Im}(\rho)$		$-\rho_1$		$-\rho_2$	
	Identified Value [1/TU]	Per Cent Predicted Error	Identified Value [1/TU]	Per Cent Predicted Error	Identified Value [1/TU]	Per Cent Predicted Error	Identified Value [1/TU]	Per Cent Predicted Error
2					0.01171	308.3	---	---
6					0.1667	143.5	0.01833	2428.3
10					0.05397	102.1	0.01053	32.4
14					0.01712	124.9	0.0401	88.2
18					0.02338	51.0	0.00947	18.34
22	0.00492	61.37	0.01445	116.32				
26	0.01302	12.00	0.01845	23.39				
30	0.01343	11.07	0.01935	17.13				
34	0.00778	20.15	0.0185	18.22				
38	0.01296	13.34	0.0173	21.33				
42	0.01383	13.08	0.01905	15.4				
46	0.02683	18.11	0.0097	126.0				
50					0.02258	23.70	0.07004	119.53
54					0.02097	17.36	---	---

*Error prediction is based on estimation of error bounds of 20 per cent (of full scale deviation) in the response functions and 25 per cent in the integrals.

**TU are normalized time units used in the response records in Figure 37.

part is minimized at $\delta = 38$ TU and the imaginary part at $\delta = 42$ TU. The final identification, (corresponding to the value of δ at which the error minimization occurs), is $\rho_{1,2} = -0.01052 \mp j 0.0107$ [1/TU].

As δ reaches 46 TU there occurs a sudden jump in the identification of the real part of the pole to a value of about -0.02 [1/TU], without any significant simultaneous change in the predicted relative error. This upset is obviously a result of the change in the input level in test no. 2 which propagates into the responses at $t = 90$ TU. The identification (of a second order system) reaches this value of t when $\delta = t/2$, i.e., $\delta = 45$ TU. This value is in agreement with the results of the identification.

From Table 7, which presents the identification results based on the response records of $y_2(t)$, observations similar to those above, concerning the identification characteristics, can be made. The identified values for the poles are $\rho_{1,2} = -0.01343 \mp j 0.01905$ [1/TU]. The deviation in the identification of the real parts between Table 6 and Table 7 are rather insignificant, considering the uncertainty level of the results. However the deviation in the identification of the imaginary parts, (about 40 percent based on the value of Table 7), is of significance and is probably a result of slight errors in evaluation of the steady state conditions.

No other information about the dynamics of the tested plant was available and therefore no comparison of results

could be made. Nevertheless, the identification appears to be very satisfactory, and in spite of the limitations and uncertainties associated with the test data, quite reliable results are obtained.

The effect of nonlinearities has not been investigated in this research and therefore no conclusions concerning their effect on the identification can be made.

Noise can be tolerated even when present at relatively high levels so long as its frequency is high compared to the system's natural frequencies. Low frequency disturbances during the testing period may have a disastrous effect on the identification as is evidenced from the results due to the change in the input level of test no. 2.

CHAPTER VIII

CONCLUSIONS AND RECOMMENDATION:

Conclusions

The time domain technique appears to be a useful method for linear system identification. It permits accurate determination of the system's parameters and overcomes most of the difficulties associated with frequency domain identification methods. The associated error propagation analysis proves to be a very powerful tool, both as a convergence criterion for actual execution of the identification and as a device for prediction of the degree of reliability of the identified model. The computational requirements are not severe, and the identification can be effectively carried out on a medium size digital computer. The computing time is insignificant, and except for the error analysis, could be carried out by hand.

Range of Applicability

From the results of the identification studies of simulated systems it appears that systems with a very wide range of natural frequency ratios can be identified. Second order systems with real poles, which have natural frequency ratios ranging from 1 to 1000, have been properly evaluated.

No particular difficulties have been encountered in identification of systems with complex poles. However, the identification of the imaginary parts of the poles for systems with high damping factors is slightly more sensitive than the identification of the corresponding real parts. Third order systems have been properly identified, and it appears that the method may be useful for identification of even higher order systems when the noise level is low.

Effect of Noise

The uncertainties introduced into the identification due to the presence of noise are generally nominal, i.e., of the same order of magnitude as the noise present. Low natural frequencies are usually less affected, whereas the high natural frequencies may be lost if the noise level is too high. Low frequency disturbances such as changes or drifts of steady state operation level cannot be tolerated for this identification method.

Determination of the System's Order

Although this identification method does not explicitly determine the order of the system, it can be evaluated through proper execution of the identification. The order of the system is indicated through the error propagation analysis.

Comparison to Frequency Domain Identification Methods

This investigation indicates superiority of the present time domain identification method over frequency

domain techniques from various aspects:

(1) It permits exact evaluation of the numerical values of the system's parameters which cannot be easily done with frequency domain methods.

(2) It is applicable over much wider ranges of natural frequency ratios than frequency domain methods.

(3) It optimizes the identification in the sense that the error of identification is minimized, thereby reducing the effects of noise and error in the data.

(4) It permits estimation of the quality of the identified model.

Limitations

No particular limitations of the time domain technique have been discovered in this investigation. It is, however, possible that for higher order systems the method may become unstable. Nonlinearities in the system, the effect of which has not been studied in this investigation, may also introduce some unexpected difficulties.

Recommendations for Future Work

The present investigation focused mainly on the determination of the general model of linear systems. Several aspects of the proposed identification method have yet to be studied.

(1) The theoretical principles for identification of the Particular Model have been outlined in Chapter II. This problem has to be investigated both experimentally and

computationally. An error evaluation procedure will have to be established for this phase of the identification. Various types of input functions should be investigated.

(2) The effect of nonlinearities in the systems should be studied with particular emphasis on the approximation of nonlinear systems by linear models. Input forcing functions must be generated which will either not affect the system parameters or will affect them in some predetermined way.

(3) An extensive comparison of this method with frequency domain techniques should be made.

(4) A rigorous experimental program designed to elucidate difficulties with this and other identification techniques should be performed. In particular various common types of chemical processes should be investigated.

(5) With regard to the error analysis an optimal search scheme for the determination of the minimum error, and correspondingly the "correct" value of the time constant, should be established.

REFERENCES

1. Aris, R., Optimal Design of Chemical Reactors, Academic Press, New York, 1961.
2. _____ and Amundson, N. R., "An Analysis of Chemical Reactor Stability and Control," Chem. Engr. Sci., Vol. 7, No. 3, p. 121, March 1958.
3. Bendat, J. S., Principles and Application of Random Noise Theory, John Wiley & Sons, Inc., New York, 1958.
4. Berger, J. S. and Perlmutter, D. D., "Chemical Reactor Stability by Liapunov's Direct Method," A.I.Ch.E. Journal, 10, No. 2, p. 233, March 1964.
5. Bishop, K. A., and Sims, R. A., "Electronic Instrumentation for Research in Process Control," Paper Presented at 1962 Mid-America Electronics Conference, Kansas City, Mo., Nov. 20, 1962.
6. Bollinger, R. E. and Lamb, D. E., "The Design of Combined Feed Forward-Feedback Control System," Preprint XVIII-4, 1963 Joint Automatic Control Conference, University of Minnesota, June 19, 1963.
7. Chang, S. S. L., Synthesis of Optimum Control Systems, McGraw Hill Book Co., Inc., New York, 1961.
8. Chestnut, H., "Control Gap: Is It a Vector?" Control Engr., 2, No. 2, p. 53, Feb. 1964.
9. Clements, W. C., Jr. and Schnelle, K. B., Jr., "Pulse Testing for Dynamic Analysis," I. & E. C. Process Design and Development Quarterly, 2, No. 2, p. 94, April 1963.
10. Coddington, E. A. and Levinson, N., Theory of Ordinary Differential Equations, McGraw Hill Book Co., Inc., 1955.
11. Cohen, W. C., and Johnson, E. F., Ind. Engr. Chem., 48, p. 1031, 1956.

12. Dreike, G. E., "Effects of Input Pulse Shape and Width on Accuracy of Dynamic System Analysis from Experimental Pulse Data," D.Sc. Dissertation, Washington University, Sever. Inst. of Tech., St. Louis, 1961.
13. _____, Hougen, J. O. and Mesmer, G., "Effects of Truncation on Time to Frequency Domain Conversions," I. S. A. Trans. 1, No. 4, p. 353, 1962.
14. _____, and Hougen, J. O., "Experimental Determination of System Dynamics by Pulse Methods," Preprint XXII-3, 1963, Joint Automatic Control Conference, University of Minnesota, June 19, 1963.
15. Dudnikov, E. E., "Determining the Transfer Function Coefficients of a Linear System from the Initial Portion of an Experimental Amplitude Phase Characteristic," Automation and Remote Control, 20, No. 5, 1959.
16. Dwyer, P. S., "On Errors in Matrix Inversion," American Statistical Association Journal, 48, p. 289, June 1953.
17. Faddeev, D. K. and Faddeeva, V. N., Computational Methods of Linear Algebra, W. H. Freeman and Co., San Francisco, 1963.
18. Fanning, R. J. and Sliepceвич, C. M., "The Dynamics of Heat Removal From a Continuous Agitated Tank Reactor," A.I.Ch.E. Journal, 5, No. 2, p. 240, June 1959.
19. Gallier, P. W., "Identification of System Frequency Response by Statistical Correlation and Spectral Analysis," Ph.D. Dissertation, The University of Oklahoma, 1962.
20. _____, Sliepceвич, C. M. and Puckett, T. H., "Some Practical Limitations of Correlation Techniques in Determining Process Frequency Response," Chem. Engr. Prog. Symposium Series, 57, No. 36, p. 59, 1961.
21. Golant, A. I. and Dudnikov, E. E., "Methods of Determining Linear System Parameters from the High Frequency Part of an Experimental Amplitude - Phase Characteristic," Automation and Remote Control, 25, No. 9, p. 1243, Sept. 1964.
22. Haskins, D. E., "Synthesis of Invariance Principle Control Systems for Chemical Processes," Ph.D. Dissertation, The University of Oklahoma, 1964.
23. Hearn, J. S., "The Dynamic Response of a Heat Exchanger to Shell Side Flow Rate Disturbances," Masters Thesis, St. Louis University, 1959.

24. Hougen, J. O. and Lees, S., "Pulse Testing a Model Heat Exchange Process," Ind. Eng. Chem., 48, p. 1064, June 1956.
25. _____, and Walsh, R. A., "Pulse Testing Methods," Chem. Engr. Prog., 57, No. 3, p. 69, March 1961.
26. Kalman, R. E., "When Is a Control System Optimal?" Preprint I-1, 1963, Joint Automatic Control Conference, University of Minnesota, June 19, 1963.
27. Laning, J. H. and Battin, R. H., Random Processes in Automatic Control, McGraw Hill Book Co., Inc., New York, 1956.
28. Lawson, J. L. and Uhlenbeck, G. E., Threshold Signals, McGraw Hill Book Co., Inc., New York, p. 58, 1950.
29. Longuet-Higgins, M. S., "On the Intervals Between Successive Zeros of a Random Function," Proc. Roy. Soc. A. 246, p. 99, 1958.
30. _____, "The Distribution Intervals Between Zeros of a Stationary Random Function," Phil. Trans. Roy. Soc. (London) A, 254, p. 557, May 1962.
31. Masubuchi, M., "Dynamic Response and Control of Multi-pass Heat Exchangers," A.S.M.E. Trans. (J. Basic Engr.) 82D, No. 1, p. 51, March 1960.
32. McFadden, J. A., "The Axis-Crossing Intervals of Random Functions," I.R.E. Trans. on Information Theory, IT-2, p. 146, Dec. 1956.
33. _____, "The Axis-Crossing Intervals of Random Functions," ibid., IT-4, p. 14, March 1958.
34. McGuire, M. L. and Lapidus, L., "On the Stability of a Detailed Packed-Bed Reactor," A.I.Ch.E. Journal, 11, No. 1, p. 85, Jan. 1965.
35. McKnight, G. W. and Worley, C. W., Nat'l. Instrument Soc. Am., Paper 53-6-1, (1953).
36. Miller, K. S., "Properties of Impulsive Response and Green's Functions," I.R.E. Transactions Circuit Theory, CT-2, p. 26, March 1955.
37. Morris, J. H., "The Dynamic Response of a Heat Exchanger to Inlet Temperature Disturbances," Masters Thesis, St. Louis University, 1959.

38. Mozley, J. M., Ind. Engr. Chem., 48, p. 1035, 1956.
39. Naur, P. (Editor), "Revised Report on the Algorithmic Language Algol 60," Communication of the A.C.M. 6, No. 1, p. 1, Jan. 1963.
40. Perlis, S., Theory of Matrices, Addison Wesley, 1952.
41. Preminger, J. and Rootenberg, J., "Some Considerations Relating to Control Systems Employing the Invariance Principle," I.E.E.E. Trans. on Automatic Control, AC-9, No. 3, p. 209, July 1964.
42. Raniel, A. J., "Zero Crossing Intervals of Gaussian Processes," I.R.E. Trans. on Information Theory, IT-8, p. 372, Oct. 1962.
43. Rice, S. O., "Mathematical Analysis of Random Noise," Bell Sys. Tech. J., 23, p. 282, July 1944; 24, p. 46, Jan. 1945.
44. Ridings, R. V., "Errors in Experimental Evaluation of System Frequency Response," Masters Thesis, St. Louis University, 1961.
45. Rosenbrock, H. H., "Distinctive Problems of Process Control," Chem. Engr. Prog., 59, No. 9, p. 43, Sept. 1962.
46. Solodovnikov, V. V. and Uskov, A. S., "A frequency Method for Determining Characteristics of Objects of Automatic Control from Data on Their Normal Usage," Automation and Remote Control, 20, No. 12, p. 1533, June 1960.
47. Stainthorp, F. P. and Axom, A. G., "The Dynamic Behavior of Thick-Walled Jacketed Pans," Chem. Engr. Sci., 20, p. 1, 1965.
48. _____, "The Dynamic Behavior of Multipass Steam Heated Exchanges--I, Response to Steam Temperature and Steam Flow Perturbations," ibid., 20, p. 107, 1965.
49. Stermole, F. J. and Larson, M. A., "Dynamic Response of Heat Exchangers to Flow Rate Changes," I. & E. C. Fundamentals Quarterly, 2, No. 1, p. 62, Feb. 1963.
50. Stewart, W. S., "Dynamics of Heat Removal from a Jacketed Agitated Vessel," Ph.D. Dissertation, The University of Oklahoma, 1960.
51. Tierney, J. W. and Homan, C. W., "Determination of Process Dynamics in Presence of Random Disturbances," Chem. Engr. Sci., 12, p. 153, 1960.

52. Watjen, J. W. and Hubbard, M., "The Dynamic Behavior of a Pulsed-Plate Extraction Column," A.I.Ch.E. Journal, 9, No. 5, p. 614, Sept. 1963.
53. Waugh, F. V., "Inversion of the Leontief Matrix by Power Series," Economica, 18, p. 143, 1950.
54. Wescott, J. H., "The Parameter Estimation Problem," Automatic and Remote Control Proceedings of the First International Congress of the International Federation of Automatic Control, Moscow, U.S.S.R., 1960; Butterworth's, London, 1961, 2, p. 779.
55. Wolf, A. A. and Dietz, J. H., "A Statistical Theory for Parameter Identification in Physical Systems," J. Franklin Inst., 247, No. 5, p. 369, Nov. 1962.
56. Worrell, R. B., "The Osage Algorithmic Language - Osage Algol," M.S. Thesis, The University of Oklahoma, 1964.
57. Zadeh, L. A., "Time Varying Networks, I," I.R.E. International Convention Record, Part 4, pp. 251-267, 1961.
58. Zadeh, L. A., "An Introduction to State-Space Techniques," Preprint, 1962 Joint Automatic Control Conference, June 1962, American Institute of Electrical Engineers, Paper 10-1, pp. 1-5.

APPENDIX A

NOMENCLATURE*

$A(t)$	= $n \times n$ coefficient matrix of a linear time varying system (2-2)
$\bar{A}(t)$	= coefficient matrix of equivalent linear system (2-38)
$\bar{A}$	= coefficient matrix of equivalent constant parameter linear system (2-51)
$\underline{A}$	= constant
a_1	= constant coefficient of $(n-1)^{th}$ derivative of the dependent variable in an n^{th} order scalar linear differential system (2-49)
$a_{1j}(t)$	= $1j^{th}$ element of matrix $A(t)$ (2-3)
$a_1(t)$	= coefficient of the $(n-1)^{th}$ derivative of the dependent variable in an n^{th} order scalar differential system (2-4)
$B(t)$	= $n \times m$ coefficient matrix of a linear time varying system (2-2)
$b_{1j}(t)$	= $1j^{th}$ element of matrix $B(t)$ (2-3)
C	= general constant matrix

*The numbers at the end of the definitions are the equation numbers in which the symbols are either defined or first used.

C_j	= arbitrary constant
$\underline{C}_1$	= distortion coefficient for matrix D_h corresponding to eigenvalue v_1 (3-33a)
c	= arbitrary constant
c	= parameter in Exponential-Cosine autocorrelation function (4-36)
D	= error (discrepancy) matrix of inverse matrix $\bar{\Phi}^{-1}$ (3-5)
D_1	= first order approximation of matrix D (3-6)
D_h	= error matrix associated with the matrix H (3-22)
d_{1j}	= $1j^{\text{th}}$ element of matrix D
$d_{h:1j}$	= $1j^{\text{th}}$ element of matrix D_h
$d_{1:kl}$	= kl^{th} element of matrix D_1
$\det.[\]$	= determinant
E	= error bound matrix for $\bar{\Phi}$ (3-1)
E_1	= error bound matrix for fundamental matrix $\bar{\Phi}_1$ (3-21)
$E [\]$	= statistical expectation (4-1)
e	= index
e_{1j}	= $1j^{\text{th}}$ element of error matrix E
$e_{1:1j}$	= $1j^{\text{th}}$ element of error matrix E_1
$f(t)$	= general time function (general forcing function of a process)
$\bar{F}(t)$	= forcing function vector of equivalent linear system (2-37)

$\bar{F}_n(t)$	= n^{th} forcing function of equivalent linear system (2-36)
G	= nonsingular $n \times n$ constant matrix
$G_{xx}(\omega)$	= power spectral density function of random function $x(t)$ (4-14)
$G_{yy}(\omega)$	= power spectral density function of random function $y(t)$ (4-40)
$g_{jk}(t)$	= coefficient of the $(n-k)^{\text{th}}$ derivative of the j^{th} independent variable of a scalar differential system (2-4)
H	= constant matrix defined by the product $\frac{\phi_{n-1}(t_{n-1}) \cdot \phi_{n-1}(t_{n-1}-\delta)^{-1}}$
H^T	= transpose of matrix H
I	= identity matrix
$\text{Im}(\)$	= imaginary part of a complex number
i	= subscript
J	= Jordan canonical matrix (2-56), (2-57)
j	= subscript
J	= $\sqrt{-I}$ (3-34)
K	= constant (parameter of Exponential-Cosine autocorrelation function) (4-36)
$\underline{K}$	= parameter of band-pass noise model (4-38)
$\bar{K}$	= spectral density of white noise
k	= subscript
k	= integer
k_j	= order of differential operator (2-5)

$L^{(n)} []$	= linear differential operator of order n (2-5)
$L^{(n)*} []$	= adjoint linear differential operator of order n (2-28)
l	= index
$l_1^{(k_j)}$	= 1 th experimental zero crossing interval
$M_j^{(k_j)} []$	= linear differential operator of order k_j of the j^{th} forcing function (2-5)
$M_j^{(k_j)*} []$	= adjoint linear operator of $M_j^{(k_j)} []$
m	= index
m	= parameter of band-pass noise model (4-38)
$N()$	= norm of a matrix (3-7)
$N_T()$	= Turing norm of a matrix (3-8)
$N_H()$	= Hotteling norm of a matrix (3-9)
n	= order of the system
P	= general $n \times n$ matrix
P_1	= general $n \times n$ matrix
P_2	= general $n \times n$ matrix
$P(\tau)$	= probability density function for zero crossing intervals (4-19)
$P_n(\tau)$	= probability density function of n successive axis crossing intervals
$Pr\{ \}$	= probability
$P(t)$	= periodic nonsingular matrix (B-5)
P_{ij}	= ij^{th} element of matrix P (3-8)
$p(n, \tau)$	= probability that there are exactly n zeros in an interval τ of a random process
Q	= nonsingular constant matrix (2-55)

q	= index
R	= constant $n \times n$ matrix (B-5)
$\text{Re}()$	= real part of a complex number
$R_{yy}(\tau)$	= autocorrelation function of random process $y(t)$ (4-7)
$R_{xx}(\tau)$	= autocorrelation function of random process $x(t)$
r	= index
$r(\tau)$	= autocorrelation function of clipped Gaussian random process (4-18)
s	= Laplace operator
s	= subscript
T	= true value matrix for $\bar{\Phi}$ (3-1)
T	= interval of random process
T	= sampling interval of discrete random process
T_0	= period of time variation
t	= time
t_1	= limit of i^{th} integration of response record (2-72) - (2-76)
t_0	= length of a finite segment of a random process
$U(t-\lambda)$	= unit step function applied at $t = \lambda$
$u(t)$	= particular solution matrix function (2-41)
v_1	= i^{th} eigenvector of matrix H (associated with eigenvalue v_1)
$\underline{v}_1$	= i^{th} eigenvector of the transpose H^T of matrix H
$\underline{v}_1^T$	= transpose of eigenvector $\underline{v}_1$
$W(t, \lambda)$	= unit impulse response (general weighting

	function) of linear system (2-12)
$W_j(t, \lambda)$	= particular weighting function of a linear system associated with input $x_j(t)$ (2-34)
$\bar{W}(t, \lambda)$	= weighting function matrix of linear system (2-47)
$\bar{W}(t - \lambda)$	= weighting function matrix for constant parameter linear system (2-53)
$W(t - \lambda)$	= weighting function for constant parameter linear system (2-54)
$W(\varphi_1, \dots, \varphi_n)(t)$	= Wronskian of a differential system (2-30)
$W_{rj}(\lambda)$	= (2-29)
x_1	= 1 th sample point of random process
$x(t)$	= random process (4-2)
$\bar{x}(t)$	= clipped Gaussian random process (4-17)
$\bar{x}(t)$	= discrete random process
$\bar{x}_j(t)$	= j th total input variable (2-83)
x_{jss}	= j th steady state input (2-83)
$x_{jp}(t)$	= input perturbation variable (2-83)
$x_1(t)$	= 1 th input function to a chemical process (2-3)
$\bar{x}(t)$	= input vector of a linear process (2-2)
$y(t)$	= dependent (output) variable of a scalar linear process (2-7)
$y(t)$	= random process (4-2)
$y_1(t)$	= 1 th dependent variable of linear process (2-3)
$\bar{y}(t)$	= output vector of a linear process (2-2)
$\bar{y}(t)$	= total output variable (2-83)
y_{ss}	= steady state output (2-83)
$y_p(t)$	= output perturbation variable (2-83)

- $z_1(t)$ = i^{th} dependent variable of equivalent linear system (2-36)
- $\mathbf{z}(t)$ = output (dependent variable) vector of equivalent linear system (2-37)

Greek Letters

- β = mean number of zero crossing intervals of a random process per unit time
- γ = initial time of random process (4-3)
- γ = mean zero crossing interval of Gaussian random process (4-20)
- γ_1 = fractional distance between sample points of random process, Figure 16
- γ_2 = fractional distance between sample points of random process, Figure 16
- Δ = difference
- δ = time interval of integration of response records (2-70a)
- $\delta(t)$ = unit impulse function (2-9)
- ϵ = small increment
- ϵ_{ij} = small constants (3-28)
- $\epsilon(t)$ = random function (4-1)
- η = (4-45)
- η_{ij} = upper bound for the ij^{th} element of the error matrices E and E_1
- Θ = sign matrix associated with matrix E
- Θ_1 = sign matrix associated with matrix E_1

θ_{1j}	= $1j^{\text{th}}$ element of sign matrix Θ (3-14)
$\theta_{1:1k}$	= $1k^{\text{th}}$ element of matrix Θ_1
λ	= time parameter at which impulse is applied (2-10)
λ	= particular value of time, parameter of $W(t, \lambda)$
ν_1	= 1^{th} eigenvalue of matrix H
$\xi(t)$	= Gaussian random process
ξ	= damping factor of second order linear system
ρ_1	= 1^{th} characteristic root or pole of linear system
$\rho(\tau)$	= normalized autocorrelation function of Gaussian random process
Σ	= summation symbol
σ	= standard deviation of statistical function
τ	= particular value of time
τ_1	= variable of integration (4-6)
τ_2	= variable of integration (4-6)
$\Phi(t)$	= fundamental solution matrix for linear system of equation (2-39), (2-40)
$\bar{\Phi}$	= abbreviated notation for $\overline{\Phi_{n-1}(t_{n-1}-\delta)}$
$\bar{\Phi}_1$	= abbreviated notation for $\overline{\Phi_{n-1}(t_{n-1})}$
$\overline{\Phi_{n-1}(t_{n-1})}$	= higher integrated fundamental solution matrix (2-75)
$\overline{\Phi_{n-1}(t_{n-1}-\delta)}$	= lower integrated fundamental solution matrix (2-76)
$\bar{\Phi}_{1j}$	= $1j^{\text{th}}$ element of matrix $\bar{\Phi}$ (3-13)

$\bar{\phi}_{1:1k}$	= $1k^{\text{th}}$ element of matrix $\bar{\Phi}_1$
$\phi_j(t)$	= j^{th} response function (or solution) of linear system
$\phi_j^*(t)$	= j^{th} solution of the adjoint linear system (2-27)
$\phi_{1j}(t)$	= $1j^{\text{th}}$ element of fundamental solution matrix $\Phi(t)$
$\overline{\phi_{1j}(t_{n-1})}$	= $1j^{\text{th}}$ element of $\overline{\Phi_{n-1}(t_{n-1})}$
$\overline{\phi_{1j}(t_{n-1}-\delta)}$	= $1j^{\text{th}}$ element of $\overline{\Phi_{n-1}(t_{n-1}-\delta)}$
$\phi_j^{(i-1)}(t)$	= $(i-1)^{\text{th}}$ derivative with respect to time of j^{th} response of linear system
$\phi_{j,k}(t)$	= k^{th} integral of j^{th} response function evaluated at time t
$\Psi(t)$	= particular solution (or response function) of linear system
$\underline{\Psi}(t)$	= complete solution matrix of linear system (2-41)
ω	= frequency
ω_0	= frequency parameter of band pass noise model (4-38)
$\bar{\omega}$	= natural oscillation frequency of second order linear system

APPENDIX B

GENERAL MATHEMATICAL RELATIONSHIPS FOR IDENTIFICATION OF LINEAR SYSTEMS WITH PERIODICALLY TIME VARYING PARAMETERS

Consider the linear system

$$\dot{\tilde{y}}(t) = A(t)\tilde{y}(t) + B(t)\tilde{x}(t) , \quad (B-1)$$

where $A(t)$ is a $(n \times n)$ matrix with periodic elements of period T_0 and $B(t)$ is a $(n \times m)$ matrix with periodic elements of period T_0 such that

$$\begin{aligned} A(t) &= A(t + T_0) \\ B(t) &= B(t + T_0) \end{aligned} \quad T_0 \neq 0 \quad (B-2)$$

Let $\Phi(t)$ be a fundamental solution matrix to the homogeneous system

$$\dot{\tilde{y}}(t) = A(t)\tilde{y}(t) \quad (B-3)$$

such that each column of $\Phi(t)$ is a solution vector $\varphi_j(t) = \text{col. } [\varphi_{j1}(t), \dots, \varphi_{jn}(t)]$, and $\varphi_j(t)$, $j = 1, \dots, n$, are linearly independent. Then $\Psi(t)$ is also a fundamental solution matrix to (B-3) if it is related to $\Phi(t)$ by

$$\Psi(t) = \Phi(t + T_0) . \quad (B-4)$$

Correspondingly to every such matrix $\Phi(t)$ there exists a periodic nonsingular matrix P with period T_0 , and a constant matrix R such that

$$\Phi(t) = P(t)e^{tR}. \quad (B-5)$$

The proof of (B-5) will be outlined here to clarify the above relationship. Since

$$\dot{\Phi}(t) = A(t)\Phi(t)$$

it follows that

$$\begin{aligned} \dot{\Psi}(t) &= \dot{\Phi}(t + T_0) = A(t + T_0)\Phi(t + T_0) \\ &= A(t)\Psi(t). \end{aligned}$$

Thus $\Psi(t)$ is a solution matrix to (B-3), and it is a fundamental matrix (nonsingular) since $\det. \Psi(t) = \det. \Phi(t + T_0) \neq 0$.

Since every fundamental matrix can be obtained from another fundamental matrix by the relation

$$\Psi(t) = \Phi(t) \cdot C \quad (B-6)$$

where C is a constant, nonsingular matrix, there exists a particular matrix C such that

$$\Phi(t + T_0) = \Psi(t) = \Phi(t) \cdot C. \quad (B-7)$$

Moreover, there exists a constant matrix R such that

$$C = e^{T_0 R}. \quad (B-8)$$

Thus by combining (B-7) and (B-8)

$$\Phi(t + T_0) = \Phi(t)e^{T_0 R} . \quad (B-9)$$

Let $P(t)$ be defined as

$$P(t) = \Phi(t)e^{-tR} . \quad (B-10)$$

Then using (B-9) it follows that

$$\begin{aligned} P(t + T_0) &= \Phi(t + T_0)e^{-(t + T_0)R} \\ &= \Phi(t)e^{T_0 R} e^{-(t + T_0)R} \\ &= \Phi(t)e^{-tR} = P(t) . \end{aligned} \quad (B-11)$$

Thus, $P(t)$ is a periodic matrix, and it is nonsingular since $\Phi(t)$ and e^{-tR} are both nonsingular.

The significance of the above relations is that the determination of a fundamental matrix $\Phi(t)$ over an interval of length T_0 leads at once to the determination of $\Phi(t)$ over $(-\infty, \infty)$.

Thus C in (B-8) is given, in terms of (B-7), by

$$\Phi^{-1}(\tau)\Phi(\tau + T_0) = \Phi^{-1}(0)\Phi(\omega) = C \quad (B-12)$$

and from this relation R is given by $(\log C)/T_0$. The matrix $P(t)$ is determined by (B-10) over an interval $(\tau, \tau + T_0)$. Since $P(t)$ is periodic it is at once determined over an interval $(-\infty, \infty)$. Thus $\Phi(t)$ is determined over $(-\infty, \infty)$ by Equation (B-5).

Every solution to (B-1) is given in terms of $\Phi(t)$ by

means of the equation

$$\begin{aligned} y(t) &= \Phi(t) \int_{-\infty}^t \Phi^{-1}(\lambda) B(\lambda) x(\lambda) d\lambda \\ &= \int_{-\infty}^t [\Phi(t) \Phi^{-1}(\lambda)] B(\lambda) x(\lambda) d\lambda \end{aligned} \quad (B-13)$$

where $\Phi(t) \Phi^{-1}(\lambda) = W(t, \lambda)$ is the weighting function for (B-1).

By means of Equation (B-5) it follows that

$$\begin{aligned} W(t, \lambda) &= \Phi(t) \Phi^{-1}(\lambda) = [P(t) e^{tR}] [P(\lambda) e^{\lambda R}]^{-1} \\ &= P(t) e^{(t-\lambda)R} P^{-1}(\lambda) \end{aligned} \quad (B-14)$$

Since:

$$P(t) = P(t + T_0) = P(t + nT_0), \quad n = 1, 2, 3, \dots$$

one obtains

$$\begin{aligned} P(t) &= P(\delta) \\ P^{-1}(\lambda) &= P^{-1}(\tau), \end{aligned}$$

where

$$\begin{aligned} \delta &= t - mT_0 & 0 \leq \delta \leq T_0 \\ \tau &= t - nT_0 & 0 \leq \tau \leq T_0. \end{aligned}$$

Thus, Equation (B-14) can be written as

$$W(t, \lambda) = P(\delta) e^{(t-\lambda)R} P^{-1}(\tau). \quad (B-15)$$

It is interesting to examine some of the properties

of the constant matrix R . If $\Phi_1(t)$ is another fundamental solution matrix of (B-3), then:

$$\Phi(t) = \Phi_1(t) \cdot G$$

where G is a nonsingular constant matrix. Substitution of the above relation into Equation (B-9) yields

$$\Phi_1(t + T_0)G = \Phi_1(t)Ge^{T_0 R}, \quad (B-16)$$

and after postmultiplication of both sides of the equation by G^{-1} :

$$\begin{aligned} \Phi_1(t + T_0) &= \Phi_1(t)Ge^{T_0 R}G^{-1} \\ &= \Phi_1(t)e^{T_0 GRG^{-1}} \end{aligned} \quad (B-17)$$

Thus, every matrix $\Phi_1(t)$ determines a matrix GRG^{-1} which is similar to R . It should, however, be noted that although R is not determined uniquely, all the characteristic quantities associated with R , which are invariant under similarity transformations, are unique. It is clear from the above discussion that the matrix $P(t)$ is also not unique and depends on $\Phi(t)$.

There exists a matrix G that transforms R into the Jordan canonical form, Equation (2-55), such that $GRG^{-1} = J$.

Equation (B-14) then becomes

$$\begin{aligned} W(t, \lambda) &= P(t)G^{-1} e^{(t-\lambda)GRG^{-1}} GP^{-1}(\lambda) \\ &= P(t)G^{-1} e^{(t-\lambda)J} GP^{-1}(\lambda). \end{aligned} \quad (B-18)$$

If $P(t)G^{-1}$ is denoted by $P_1(t)$ then:

$$W(t, \lambda) = P_1(t) e^{(t-\lambda)J} P_1^{-1}(\lambda) \quad (B-19)$$

The matrices J and $P_1(t)$ are uniquely described by the system and are not dependent on any particular fundamental solution matrix.

Equation (B-19) reduces to

$$W(t, \lambda) = P_1(\delta) e^{(t-\lambda)J} P_1^{-1}(\tau) \quad (B-20)$$

The uniqueness of the matrices $P_1(t)$ and J is important when analyzing real systems by means of experimental data.

APPENDIX C

DERIVATION OF EQUATION (4-42) FOR THE MEAN ZERO CROSSING INTERVAL OF BAND PASS NOISE

The mean zero-crossing interval of a Gaussian random process is given by Equation (4-27) as:

$$\gamma = \frac{1}{\beta} = \pi \left[-\frac{1}{\rho''(0)} \right] . \quad (4-27)$$

The normalized correlation function $\rho(\tau)$ of a band-pass noise model is described by Equation (4-41) as:

$$\rho(\tau) = \frac{\sin m\omega_0\tau - \sin \omega_0\tau}{(m-1)\omega_0\tau} . \quad (4-41)$$

Differentiation of (4-41) with respect to τ yields

$$\begin{aligned} \rho'(\tau) = & \frac{1}{(1-m)\omega_0} \left[\frac{1}{\tau} (\omega_0 \cos \omega_0\tau - m\omega_0 \cos m\omega_0\tau) \right. \\ & \left. - \frac{1}{\tau^2} (\sin \omega_0\tau - \sin m\omega_0\tau) \right] \end{aligned} \quad (C-1)$$

For the existence of (4-27) it is necessary to show that $\rho'(0) = 0$. Substitution of (4-41) into (C-1) and rearrangement of the equation yields

$$\begin{aligned} \rho'(\tau) = & \frac{1}{(1-m)\omega_0} \frac{1}{\tau} (\omega_0 \cos \omega_0\tau - m\omega_0 \cos m\omega_0\tau) - \frac{1}{\tau} \rho(\tau) \\ = & \frac{1}{\tau} \left[\frac{\omega_0 \cos \omega_0\tau - m\omega_0 \cos m\omega_0\tau}{(1-m)\omega_0} - \rho(\tau) \right] \end{aligned} \quad (C-2)$$

In the limit, as τ approaches zero, Equation (C-2) becomes (after utilization of L'Hospital's rule)

$$\lim_{\tau \rightarrow 0} \rho'(\tau) = \frac{0}{0} = \lim_{\tau \rightarrow 0} \left[\frac{(m\omega_0)^2 \sin m\omega_0\tau - \omega_0^2 \sin \omega_0\tau}{(1-m)\omega_0} - \rho'(\tau) \right], \quad (C-3)$$

from which it follows that

$$2 \lim_{\tau \rightarrow 0} \rho'(\tau) = \lim_{\tau \rightarrow 0} \frac{(m\omega_0)^2 \sin m\omega_0\tau - \omega_0^2 \sin \omega_0\tau}{(1-m)\omega_0} = 0. \quad (C-4)$$

and finally

$$\rho'(0) = 0 \quad (C-5)$$

as required. The second derivative of $\rho(\tau)$ can be evaluated as:

$$\begin{aligned} \rho''(\tau) &= \frac{1}{\tau} \left[\frac{(m\omega_0)^2 \sin m\omega_0\tau - \omega_0^2 \sin \omega_0\tau}{(1-m)\omega_0} \right] \\ &\quad - \frac{1}{\tau^2} \left[\frac{\omega_0 \cos \omega_0\tau - m\omega_0 \cos m\omega_0\tau}{(1-m)\omega_0} \right] - \frac{\rho'(\tau)}{\tau} + \frac{\rho(\tau)}{\tau^2} \\ &= \frac{1}{\tau} \left[\frac{(m\omega_0)^2 \sin m\omega_0\tau - \omega_0^2 \sin \omega_0\tau}{(1-m)\omega_0} - 2\rho'(\tau) \right] \quad (C-6) \end{aligned}$$

The last step in the derivation of Equation (C-6) results from substitution of Equation (C-2) into it, in order to eliminate $\rho(\tau)$. In the limit as τ approaches zero:

$$\begin{aligned} \rho''(0) = \lim_{\tau \rightarrow 0} \rho''(\tau) &= \frac{0}{0} = \lim_{\tau \rightarrow 0} \left[\frac{(m\omega_0)^3 \cos m\omega_0\tau - \omega_0^3 \cos \omega_0\tau}{(1-m)\omega_0} \right. \\ &\quad \left. - 2\rho''(\tau) \right] \quad (C-7) \end{aligned}$$

Thus

$$\rho''(0) = \frac{1}{3} \frac{\omega_0^2(m^3 - 1)}{1 - m} = -\frac{1}{3} \omega_0^2(m^2 + m + 1) . \quad (C-8)$$

The final result is obtained by substitution of (C-8) into (4-27) as:

$$\gamma = \frac{\pi}{\omega_0} \sqrt{\frac{3}{m^2 + m + 1}} \quad (4-42)$$

APPENDIX D

DERIVATION OF EQUATION (4-45) FOR THE RATIO BETWEEN THE STANDARD DEVIATION AND THE MEAN ZERO CROSSING INTERVAL OF BAND-PASS NOISE

The variance of zero crossing intervals is expressed by Equation (4-44) as

$$\sigma^2 = 2 \int_0^Y (\tau - \gamma)^2 P(\tau) d\tau . \quad (4-44)$$

The zero crossing probability distribution function was expressed by Equation (4-34) as

$$P_0(\tau) = \frac{r''(\tau)}{4\beta} = \frac{\gamma}{4} r''(\tau) . \quad (4-34)$$

Substitution of Equation (4-34) into (4-44) yields

$$\begin{aligned} \frac{\sigma^2}{2} &= \frac{\gamma}{4} \int_0^Y (\tau - \gamma)^2 r''(\tau) d\tau \\ &= \frac{\gamma}{4} \int_0^Y (\tau^2 + \gamma^2 - 2\gamma\tau) r''(\tau) d\tau \\ &= \frac{\gamma}{4} \int_0^Y \tau^2 r''(\tau) d\tau + \frac{\gamma^3}{4} \int_0^Y r''(\tau) d\tau - \frac{\gamma^2}{2} \int_0^Y \tau r''(\tau) d\tau \end{aligned} \quad (D-1)$$

To evaluate Equation (D-1) it is necessary to solve each term on the right of the equation individually. Thus the first term yields

$$\begin{aligned}
\int_0^Y \tau^2 r''(\tau) d\tau &= \left[\tau^2 r'(\tau) \right]_0^Y - 2 \int_0^Y \tau r'(\tau) d\tau \\
&= \left[\tau^2 r'(\tau) \right]_0^Y - \left[2\tau r(\tau) \right]_0^Y + 2 \int_0^Y r(\tau) d\tau \\
&= Y^2 r'(Y) - 2Yr(Y) + 2 \int_0^Y r(\tau) d\tau . \quad (D-2)
\end{aligned}$$

The second term yields:

$$\int_0^Y r''(\tau) d\tau = \left[r'(\tau) \right]_0^Y = r'(Y) - r'(0) . \quad (D-3)$$

And for the third term one obtains

$$\begin{aligned}
\int_0^Y \tau r''(\tau) d\tau &= \left[\tau r'(\tau) \right]_0^Y - \int_0^Y r'(\tau) d\tau \\
&= \left[\tau r'(\tau) \right]_0^Y - \left[r(\tau) \right]_0^Y \\
&= Yr'(Y) - r(Y) + r(0) . \quad (D-4)
\end{aligned}$$

After substitution of (D-2), (D-3) and (D-4) into (D-1):

$$\begin{aligned}
\frac{\sigma^2}{2} &= \frac{Y}{4} \left[Y^2 r'(Y) - 2Yr(Y) + 2 \int_0^Y r(\tau) d\tau \right] + \frac{Y^3}{4} [r'(Y) - r'(0)] \\
&\quad - \frac{Y^2}{2} [Yr'(Y) - r(Y) + r(0)] \\
&= \frac{Y}{2} \int_0^Y r(\tau) d\tau - Y^2 \left[\frac{r(0)}{2} + Y \frac{r'(0)}{4} \right] \\
&= \frac{Y}{2} \int_0^Y r(\tau) d\tau - Y^2 \left[\frac{1}{2} + \frac{Y}{2} \left(-\frac{1}{Y} \right) \right] \\
&= \frac{Y}{2} \int_0^Y r(\tau) d\tau \quad (D-5)
\end{aligned}$$

The second to the last step in the derivation of Equation (D-5) follows since $r(0) = 1$ by definition, Equation (4-18), and since $r'(0) = -2/\gamma'$, Equation (4-24). Division of Equation (D-5) by γ^2 yields

$$\left(\frac{\sigma}{\gamma}\right)^2 = \frac{1}{\gamma} \int_0^{\gamma} r(\tau) d\tau. \quad (D-6)$$

For a band-pass noise model

$$r(\tau) = \frac{2}{\pi} \sin^{-1} \left[\frac{\sin m\omega_0\tau - \sin \omega_0\tau}{(m-1)\omega_0\tau} \right], \quad (D-7)$$

and γ is given by Equation (4-42) as

$$\gamma = \frac{\pi}{\omega_0} \sqrt{\frac{3}{m^2 + m + 1}} \quad (4-42)$$

Substitution of the above two equations into (D-6)

yields:

$$\left(\frac{\sigma}{\gamma}\right)^2 = \frac{\omega_0}{\pi} \sqrt{\frac{m^2 + m + 1}{3}} \cdot \frac{\pi}{\omega_0} \sqrt{\frac{3}{m^2 + m + 1}} \int_0^{\frac{\pi}{\omega_0} \sqrt{\frac{3}{m^2 + m + 1}}} \sin^{-1} \left[\frac{\sin m\omega_0\tau - \sin \omega_0\tau}{(m-1)\tau\omega_0} \right] d\tau \quad (D-7)$$

Through changing the variable of integration in Equation (D-7) from τ to $\omega_0\tau$, Equation (D-7) becomes

$$\left(\frac{\sigma}{\gamma}\right)^2 = \frac{2}{\pi^2} \sqrt{\frac{m^2 + m + 1}{3}} \int_0^{\pi \sqrt{\frac{3}{m^2 + m + 1}}} \sin^{-1} \left[\frac{\sin m\tau - \sin \tau}{(m-1)\tau} \right] d\tau \quad (D-8)$$

After denoting $\pi \sqrt{\frac{3}{m^2 + m + 1}}$ by η the required result is obtained:

$$\left(\frac{\sigma}{\gamma}\right)^2 = \frac{2}{\pi\eta} \int_0^{\eta} \sin^{-1} \left[\frac{\sin m\tau - \sin \tau}{(m-1)\tau} \right] d\tau \quad (4-45)$$

APPENDIX E

COMPUTER PROGRAM LISTINGS FOR IDENTIFICATION PROCESS

In this appendix the listings of the computer programs that were used for the identification procedure, written in the Osage Algol language [56], are presented.

The meaning of the various input variables and proper explanations of the various computational procedures are given in comments throughout the programs.

Programs for the preparation of the fundamental matrices $\Phi_{n-1}(t_{n-1})$ and $\Phi_{n-1}(t_{n-1}-\delta)$ are not listed because of their specialized nature, their dependence on the particular recording procedures and on the recording facilities that are available.

For utilization of the identification programs the following comments may be helpful:

1. Care should be taken to read-in the various input variables according to the proper formats as indicated by the declarations at the beginning of Program 170. Thus variables that are declared as Real (or Real Array) should have a decimal point, and variables that are declared as Integer (or Integer Array) may not have a decimal point.

2. The series identification number (identified as

sys) should have at most four digits.

3. If only the final results at each value of δ and the final identification results are desired, only $t1^4$ and $t1^5$ should be read-in as 1; all other t 's should be read-in as zero.

4. For identification of second or third order systems, values of 0.0000001 for ϵ s and the convergence criteria ϵ s 1, ϵ s 2, ϵ s 3 and ϵ s 4 are generally satisfactory.

Program 170;

Begin Integer n,K,t1,t2,t3,t4,t5,t6,t7,t8,t9,t10,t11,t12,t13,t14,t15;
Real eps,eps0,eps1,eps2,eps3;

Comment This program carries out the identification of linear systems of order n. K is the maximum number of

iterations desired for the bairsto procedure.

The parameters t1,...,t15 are switches for optional intermediate outputs. If the particular output is

desired the value for the corresponding t is read

in as i . eps is the criterion for singularity of the

fundamental solution matrix in the inversion routine.

eps0,eps1,eps2 and eps3 are convergence criteria

for the bairsto procedure;

READPT(DECIMAL,n,K,t1,t2,t3,t4,t5,t6,t7,t8,t9,t10,t11,t12,t13,t14,
t15,eps,eps0,eps1,eps2,eps3);

Begin Integer i,j,last,sys,alarm,r1,m,ldm;

Real T1,delta,discrep,dism,d1,TC,EMTC,norm1,norm2,norm3,norm4;

Real Array A,a,a1,discrep,prod,P[1:n,1:n],boundda,x,z,rz,law,r1,rv1,

law1,ATC,Root,Conf[1:n],b[0:n],pe1,pe2[1:n,1:n];

Real NP,TCL,ENVT;

Real Array Confid,Con,TT2,delta2,delta1,TT1,Co,Ce[1:n];

Integer Array ex,nat[1:n];

Real Procedure norm(a,n); Value n; Integer n; Array a;

Begin Integer i,j; Real f; f = 0;

For i = 1 Step 1 Until n Do

For j = 1 Step 1 Until n Do

f = f + a[i,j] * a[i,j]; norm = SQRT(f) End;

Mc Procedure invert(200,1,4);

Mc Procedure multip(202,1,4);

Mc Procedure charec(206,1,3);

Mc Procedure bairsto(207,1,12);

Mc Procedure inveter(210,1,8);

Mc Procedure inverro(213,1,5);

Mc Procedure multerr(214,1,6);

Mc Procedure errorr(215,1,9);

Format out(1,'Matrix',S2,'a',S2,'is',S2,'singular',J2);

Format out(1,'System',S2,'Identification',S5,'order',S2,I2,S5,'Series',

S2,I4,S5,'(m-1)',R5,S5,'delta=',R5,J2);

Format out(2,'Matrix',S2,'a',J1);

Format out(3,'J1,S5,6(R10,S3),J1);

Format out(4,'Matrix',S2,'a1',J1);

Format out(5,'J1',Inverse',S2,'of',S2,'a',J1);

Format skip(J1);

Format out(6,'J1',Error',S2,'Bounds',S2,'of',S2,'Inverse',J1,'alarm',I2,

J1);

Format out(7,'J1',Error',S2,'Bounds',S2,'of',S2,'Data',J1);

Format out(8,'J1',Error',S2,'Bounds',S2,'of',S2,'Weight',S2,'Matrix',J1);

Format out(9,'J1',Weight',S2,'Matrix',J1);

Format out(10,'J1','Coefficients',S2,'of',S2,'characteristic',S2,'Equation',

J1,S5,6(R10,S3),J2);

Format out(11,'J1','Characteristic',S2,'Roots',J1);

Format out(12,S5,2(R10,S3),2(I2,S11),J1);

Format out(13,'J1','Eigenvalues',S2,'Unreliable',J1);

Format out(14,'J1','Eigenvector',S2,I2,S2,'Singular',J1);

Format out(15,'J1','Eigenvector',S2,I2,J1);

Format out(16,'J1','Eigenvector',S2,'Transpose',S2,I2,J1);

Format out(17,'J1','Eigenvalue',S2,I2,S2,'-',S2,R10,S2,'+',S2,'J1',R10,J1);

Format out(18,'J2','Natural',S2,'Frequency',S1,'-',R10,S5,'error',R10,J1);

Format out(19,'J1','Negative',S2,'Argument',S2,'of',S2,'Log',J1);

Format out(20,'J1','norm',S1,'a',S8,R10,S5,'norm',S1,'b',S3,R10,

J1,'norm',S1,'a',S1,'Inverse',R10,S5,'norm',S1,'prod',

R10,J1);

Format out(21,'J1','Relative',S2,'error',S2,'of',S2,'a',J1);

Format out(22,'J1','Relative',S2,'error',S2,'of',S2,'a1',J1);

Format out(23,'J1','Relative',S2,'error',S2,'bounds',S2,'of',S2,'Inverse',J1);

Format out(24,'J1','Relative',S2,'error',S2,'bounds',S2,'of',S2,

'weight',S2,'matrix',J1);

```

RETURN out25(j), 'Error', s2, 'of', s2, 'eigenvalue', s2, i2, j1, s5,
    'Real', s2, 'Part', s3, 'Imaginary', s2, 'part', j1);
RETURN out26(j), 'Relative', s2, 'error', s2, 'of', s2, 'eigenvalue', s2, i2,
    j1, s5, 'Real', s2, 'Part', s3, 'Imaginary', s2, 'Part', j1);
RETURN out27(j), 'Relative', s2, 'Error', s2, 'of', s2, 'Natural', s2,
    'Frequency', j1, 'computed', s2, '-', s2, s10, s5, 'Predicted',
    s2, '-', s10, j1);
RETURN out28(j), 'System', s2, 'Identification', s2, 'of', s2, 'TEST',
    s2, 'Series', s3, i4, s3, 'ORDER', s3, i2, j1);)
RETURN out29(j), 'NATURAL', s2, 'FREQUENCY', s4, '-', s4, s10, s5,
    'PREDICTED', s2, 'RELATIVE', s2, 'ERROR', s4, '-', s4, s10, s5, j1);)
RETURN out30(j),
    'out31(j), 'NATURAL FREQUENCY', j1, s2, 'REAL PART (DAMPING) =', s10, s13,
    'ERROR =', s10, j1, s2, 'COMPLEX PART (beta*delta) =',
    s10, s10, 'ERROR =', s10, j1);)
RETURN out32(j), 'RELATIVE ERROR OF NATURAL FREQUENCY', j1,
    'REAL PART', s60, 'COMPLEX PART', j1, 'COMPUTED =', s10, s2, 'PREDICTED',
    s2, '-', s2, s10, s9, 'PREDICTED =', s10, j1);)
RETURN out33(j), 'r(n-1)', s2, '-', s10, s8, 'delta', s2, '-', s10, j1);)
RETURN out34(j), 'NATURAL FREQUENCY', j1, 'REAL PART =', s10, s8,
    'PREDICTED RELATIVE ERROR =', s10, j1, 'COMPLEX PART =', s10,
    s5, 'PREDICTED RELATIVE ERROR =', s10, j1);)
FOR i = 1 STEP 1 UNTIL n DO BEGIN
    COMMENT read in the error bounds bounds[i] of the data.
    i is the ith row of the fundamental matrix;
    READP(DECIMAL, bounds[i]);
    FOR j = 1 STEP 1 UNTIL n DO BEGIN
    IF t=1 THEN BEGIN PRINT(out8);
    FOR i = 1 STEP 1 UNTIL n DO
    PRINT(out3, a[i, 1], ..., a[i, n]); PRINT(ackp) END;
    i=n = 0;
    COMMENT Input of data. a[i, j] and a[i, j] are the elements of
    the matrices  $\phi_{n-1}(t_{n-1}-t^0)$  and  $\phi_{n-1}(t_n-t^0)$  respectively.
    t1 is the value of  $t_{n-1}$ , last is an
    indication whether the record has been exhausted. If last=0
    then the record has been exhausted. srs is an identification
    number of the system;
    Initial: FOR i = 1 STEP 1 UNTIL n DO
    FOR j = 1 STEP 1 UNTIL n DO
    READP(REAL, a[i, j], s1[i, j]);
    READP(REAL, s1, delta, last, srs);
    COMMENT When the record is exhausted the final identification output
    will be executed;
    BEGIN IF i=n=0 THEN RETURN ELSE IF -(i=n=srs) THEN BEGIN
    PRINT(out28, i, n); FOR i = 1 STEP 1 UNTIL n DO BEGIN
    IF Conf[i] < Conf[i] THEN BEGIN PRINT(out29, Root[i], Conf[i]);
    PRINT(out33, t1[i], delta[i]); END
    ELSE BEGIN j=i; ALTER PRINT(out34, Co[i], Conf[j], Co[j], Conf[j]);
    PRINT(out33, t12[j], delta2[j]); j=i+1;
    IF j=2*INTEGER(j/2)=0 THEN BEGIN i=j; GETE ALTER END END;
    PRINT(out30);
    RETURN FOR i = 1 STEP 1 UNTIL n DO BEGIN
    Root[i] = Conf[i]; t1[i] = delta[i]; Co[i] = 0;
    t12[i] = delta2[i]; Co[i] = Conf[i];
    Co[i] = Co[i] = 1 END; i=n = srs END END;
    BEGIN INTERVAL new;
    COMMENT Optional input of expected values of time constants.
    (useful and necessary for comparative purposes only);
    IF t=i=1 THEN BEGIN
    FOR i = 1 STEP 1 UNTIL n DO
    READP(REAL, a1c[i]) END;
    PRINT(out1, n, srs, t1, delta);
    IF t=2=1 THEN BEGIN PRINT(out2);
    FOR i = 1 STEP 1 UNTIL n DO
    PRINT(out3, a[i, 1], ..., a[i, n]); PRINT(ackp);
    PRINT(out4);
    FOR i = 1 STEP 1 UNTIL n DO
    PRINT(out3, a[i, 1], ..., a[i, n]); PRINT(ackp) END;

```

```

If t9=1 Then Begin
  For i = 1 Step 1 Until n Do
    For j = 1 Step 1 Until n Do Begin
      po[i,j] = ((A[i,j]/A[i,j]))';
      po[i,j] = ((A[i,j]/A[i,j]))'; End;
    PRINT(out8); For i = 1 Step 1 Until n Do
      PRINT(out3,po[i,1],...,po[i,n]); PRINT(akip);
    PRINT(out8); For i = 1 Step 1 Until n Do
      PRINT(out3,po[i,1],...,po[i,n]); PRINT(akip) End;
    If t10=1 Then Begin norm1 = norm(a,n); norm2 = norm(a1,n) End;
    Invert(a,n,po,alingul); Goto next;
  next: If t3=1 Then Begin PRINT(out5);
    For i = 1 Step 1 Until n Do
      PRINT(out3,a[i,1],...,a[i,n]); PRINT(akip) End;
    alarm = 0;
    For i = 1 Step 1 Until n Do Begin
      For j = 1 Step 1 Until n Do Begin
        pred[i,j] = disc[i,j] - 0; p[i,j] = 0 End;
        ex[i] = 0; b[i] = 0; x[i] = x[i] - 0; mat[i] = 0 End; b[i] = 1;
        Invert(a,n,diso,A,alarm);
        If t4=1 Then Begin PRINT(out6,alarm);
          For i = 1 Step 1 Until n Do
            For j = 1 Step 1 Until n Do
              po[i,j] = ((disc[i,j]/A[i,j]))';
              PRINT(out8); For i = 1 Step 1 Until n Do
                PRINT(out3,po[i,1],...,po[i,n]); PRINT(akip) End;
                multip(a,pred,a1,a);
                If t10=1 Then Begin norm3 = norm(a,n); norm4 = norm(pred,n);
                PRINT(out8,norm1,norm2,norm3,norm4) End;
                multerr(a,a1,a,diso,A,p);
                If t5=1 Then Begin PRINT(out9);
                  For i = 1 Step 1 Until n Do
                    PRINT(out3,pred[i,1],...,pred[i,n]);
                    PRINT(akip); PRINT(out8);
                  For i = 1 Step 1 Until n Do
                    PRINT(out3,p[i,1],...,p[i,n]); PRINT(akip);
                    Body: obareo(pred,n,b); m = 0;
                    balrest(a,b,opso,ops1,opsc,ops3,K,m,x,z,mat,ex);
                    If t12=1 Then Begin
                      For i = 1 Step 1 Until n Do
                        For j = 1 Step 1 Until n Do
                          po[i,j] = ((p[i,j]/pred[i,j]))';
                          PRINT(out2b); For i = 1 Step 1 Until n Do
                            PRINT(out3,po[i,1],...,po[i,n]); PRINT(akip) End;
                            If t6=1 Then Begin
                              PRINT(out10,b[i,0],...,b[i,n]);
                              PRINT(out11);
                              For i = 1 Step 1 Until n Do
                                PRINT(out12,x[i],x[i],mat[i],ex[i]) End;
                                x1 = 0;
                              End;
                            End;
                          Begin Inverse norm;
                          For i = 1 Step 1 Until n Do Begin
                            If ex[i]=4 Then Begin
                              PRINT(out13); Goto advance End; If mat[i]=1 Then Begin
                                1-4: If n=2 THEN (n/2)=1-1-1-1 Then
                                  x[i]=x[i] < |x[i]| Then x[i]
                                Else x[i];
                                If n=2 THEN (n/2)=1-1-1-1 Then Goto All;
                                x1 = x1+i; x2[x1] = x[i]; x3[x1] = 0;
                                x1 = x1+i; x2[x1] = x[i]; x3[x1] = -x[i]; End End;
                                All:
                                x1 = x1+i; x2[x1] = x[i]; x3[x1] = 0 End Else Begin
                                  x1 = x1+i; x2[x1] = x[i]; x3[x1] = x[i];
                                  x1 = x1+i; x2[x1] = x[i]; x3[x1] = -x[i]; End End;

```

```

For i = 1 Step 1 Until n Do
For j = 1 Step 1 Until n Do
a[i,j] = pred[j,i];
For i = 1 Step 1 Until r1 Do Begin
For j = 1 Step 1 Until n Do Begin
rv[j] = inv[j] - 0; rvi[j] = inv[i] - 0 End;
ivector(pred,rv,inv,rx[i],imx[i],n,eps,singul);
Goto instead; single: PRINT(out14,i); Goto advance;
instead: ivector(a,rvi,inv1,rx[i],imx[i],n,eps,singul);
If t7=1 Then Begin PRINT(out15,i);
PRINT(out3,rv[1],...,rv[n]);
PRINT(out3,inv[1],...,inv[n]); PRINT(skip);
PRINT(out16,i);
PRINT(out3,rvi[1],...,rvi[n]);
PRINT(out3,inv1[1],...,inv1[n]); PRINT(skip) End;
discrep = disim = 0;
error(n,imx[i],discrep,disim,F,rv,inv,rvi,inv1);
If t8=1 Then Begin
PRINT(out17,i,rx[i],imx[i]); PRINT(out25,i);
PRINT(out3,discrep,disim) End;
If t13=1 Then Begin norm1 = discrep/rx[i]; norm2 = disim/imx[i];
PRINT(out26,i); PRINT(out3,norm1,norm2) End;
Comment Calculation of system poles. At this stage also the testing
and minimization of error is carried out;
If imx[i]=0 Then Begin If rx[i] < 0 Then PRINT(out19) Else Begin
d1 = LN(rx[i]);
TC = d1/delta; ERTC = discrep/(delta*rx[i]);
PRINT(out18,TC,ERTC); If t15=1 Then Begin
norm1 = |((TC-ATC[i])/ATC[i])|;
norm2 = |(ERTC/TC)|;
If norm2 < Conf[i] Then Begin Conf[i] = norm2;
TT1[i]=T1; delta1[i] = delta;
Root[i] = TC End;
PRINT(out27,norm1,norm2) End End End
Else Begin d1= rx[i]*rx[i] +imx[i]*imx[i];
TC=LN(d1)/(2*delta);
ERTC = (|discrep*rx[i]| + |disim*imx[i]|)/(delta*d1);
TCI = rx[i]/SQRT(d1);
NF=ARCTAN(SQRT(1-TCI*TCI)/TCI);
ERNF= (|discrep| + |rx[i]*delta*ERTC|)/imx[i];
PRINT(out31,TC,ERTC,NF,ERNF);
norm3= |(ERTC/TC)|; norm4=|(ERNF/NF)|;
norm1=|((TC-ATC[i])/ATC[i])|;
If norm3<Conf[i] Then Begin Conf[i]=norm3;
TT2[i]=T1; delta2[i]=delta;
Confid[i]=norm4; Co[i]=TC; Co[i]=NF End;
PRINT(out32,norm1,norm3,norm4)EndEnd;
End;
advance: If last=0 Then Goto final;
Goto initial;
final: End End;

```

Program 200;

Procedure invert (a,n,eps,singul); Value n,eps;

Real Array a; Integer n; Real eps; Label singul;

Comment inverts a matrix in it's own space using the

Gauss-Jordan method with complete matrix pivoting.

I.e. at each stage the pivot has the largest absolute value of any element in the remaining matrix. The

coordinates of the successive matrix pivots used at

each stage of the reduction are recorded in the successive

element positions of the row and column index vectors r

and c. These are later called upon by the procedure permute

which rearranges the rows and columns of the matrix. If

the matrix is singular the procedure exits to an appropriate

label in the main program;

```

Begin Integer i,j,k,l,pivi,pivj,p; Real pivot, h;
Integer Array r,c[1:n];
Mo Procedure permute (201,1,8);
Comment set row and column index vectors;
For i = 1 Step 1 Until n Do r[i] = c[i] - 1;
Comment find initial pivot; pivi = pivj = 1;
For i = 1 Step 1 Until n Do
For j = 1 Step 1 Until n Do
If |(a[pivi,pivj])'| < |(a[i,j])'| Then Begin
pivi = i; pivj = j End;
Comment start reduction;
For i = 1 Step 1 Until n Do Begin
l = r[i]; r[i] = r[pivi]; r[pivi] = l; l = c[i]; c[i] = c[pivj];
c[pivj] = l; If |(a[r[i], c[i]])'| < eps Then Begin
Comment here include an appropriate output procedure to record
1 and the current values of r[1:n] and c[1:n];
Goto singul End;
For j = n Step -1 Until i+1, i-1 Step -1 Until 1 Do
a[r[i], c[j]] = a[r[i], c[j]]/a[r[i], c[i]];
a[r[i], c[i]] = 1/a[r[i], c[i]]; pivot = 0;
For k = 1 Step 1 Until i-1, i+1 Step 1 Until n Do Begin
For j = n Step -1 Until i+1, i-1 Step -1 Until 1 Do Begin
h = a[r[i], c[j]] * a[r[k], c[i]];
a[r[k], c[j]] = a[r[k], c[j]] - h;

If i < k & i < j & |(pivot)'| < |(a[r[k], c[j]])'|
Then Begin pivi = k; pivj = j; pivot = a[r[k], c[j]] End conditional
End j loop;
a[r[k], c[i]] = -a[r[i], c[i]]*a[r[k], c[i]] End k loop End i loop and reduction;
Comment rearrange rows; permute(a[j,p],a[k,p],j,k,r,c,n,p);
Comment rearrange columns; permute(a[p,j],a[p,k],j,k,c,r,n,p)
End invert;

```

```

Program 201;
Procedure permute(a,b,j,k,s,d,n,p); Value n;
Real a,b; Integer Array s,d; Integer j,k,n,p;
Comment a procedure using Jensen's device which exchanges
rows or columns of a matrix to achieve a rearrangement
specified by the permutation vectors s,d[1:n]. Elements
of s specify the original source locations while elements
of d specify the desired destination locations. Normally
a and b will be called as subscripted variables of the
same array. The parameters j,k nominate the subscripts of
the dimension affected by the permutation, p is the Jensen
parameter. As an example of the use of this procedure
suppose r,c[1:n] to contain the row and column subscripts
of the successive matrix pivots used in a matrix inversion
of an array a[1:n,1:n] i.e. r[1], c[1] are the relative
subscripts of the first pivot r[2], c[2] those of the second
pivot and so on. The two calls permute(a[j,p],a[k,p],j,k,
r,c,n,p) and permute(a[p,j],a[p,k],j,k,c,r,n,p) will perform
the required rearrangement of rows and column respectively;
Begin Integer Array tag, loc[1:n]; Integer i,t;
Real w;
Comment set up initial vector tag number and address arrays;
For i = 1 Step 1 Until n Do tag[i] = loc[i] = 1;
Comment start permutation;
For i = 1 Step 1 Until n Do Begin
t = s[i]; j = loc[t]; k = d[i];
If -(j=k) Then Begin For p = 1 Step 1 Until n Do Begin
w = a; a = b; b = w End; tag[j] = tag[k]; tag[k] = t;
loc[t] = loc[tag[j]]; loc[tag[j]] = j End jk conditional
End i loop
End permute;

```

```

Program 202;
Procedure multip(n,a,b,c); Value n;
Array a,b,c; Integer n;
Comment This procedure multiplies two matrices b and c such
that  $a[i,j] = b[i,k]*c[k,j]$  and the result stored in a;
Begin Integer i,j,k;
For i = 1 Step 1 Until n Do
For j = 1 Step 1 Until n Do Begin
a[i,j] = 0;
For k = 1 Step 1 Until n Do
a[i,j] = a[i,j]+b[i,k]*c[k,j] End
End multip;

```

```

Program 206;
Procedure charac(a,n,b); Value n; Integer n;
Array a,b;
Comment finds the coefficients of the characteristic
equation of a matrix b[i] of the form  $b[0]*r^n +$ 
 $b[1]*r^{n-1} + \dots + b[n-1]*r + b[n] = 0$ . the procedure
uses Hessenberg's method and if it fails it
switches to Leverrier's method;
Begin Integer i,j,l,k,m; Real fu,fun,func,beta;
Array q,p[1:n,1:n], S[1:n];
Real Procedure f(a,q,z,s,n); Value z,s,n;
Integer z,s,n; Array a,q;
Begin Integer k; Real t;
t = 0; For k = s-1 Step 1 Until n Do
t = t+a[z,k] * q[k,s-1]; t = t End;
Real Procedure g(a,p,z,s,m); Value z,s,m;
Integer z,s,m; Array p,q;
Begin Integer k; Real t;
t = 0; For k = 2 Step 1 Until m Do
t = t+q[z,k] * p[k,s-1]; g = t End;
For i = 1 Step 1 Until n Do If i=1 Then Begin
q[i,1] = 1; q[i,2] = 0; p[i,1] = -a[i,1] End
Else If i=2 Then Begin
q[i,1] = 0; q[i,2] = a[2,1]; p[i,1] = -1 End
Else Begin q[i,1] = 0; q[i,2] = a[i,1]; p[i,1] = 0 End;
If q[2,2]=0 Then Goto next;
For j = 3 Step 1 Until n Do Begin
For i = 1 Step 1 Until n Do If i<j Then Begin
q[i,j] = 0; fun = f(a,q,i,j,n) + g(q,p,i,j,i-1); p[i,j-1] =
-fun/q[i,1] End
Else If i=j Then Begin
p[i,j-1] = -1; fun = f(a,q,i,j,n) + g(q,p,i,j,j-1);
q[i,j] = fun End
Else Begin p[i,j-1] = 0; fun = f(a,q,i,j,n) + g(q,p,i,j,j-1);
q[i,j] = fun End;
If q[j,j] = 0 Then Goto next End;
For i = 1 Step 1 Until n Do If i=1 Then
p[i,n] = -a[i,n] * q[n,n] Else Begin
fun = g(q,p,i,n+1,i-1) + a[i,n] * q[n,n]; p[i,n] = -fun/q[i,1] End;
Comment Now that the matrix p[i,j] has been evaluated the following
procedure will evaluate the coefficients b[j] in the characteristic
equation;
For i = 1 Step 1 Until n Do Begin b[i] = 0;
For j = 1 Step 1 Until n Do q[i,j] = 0 End;
q[i,1] = 1; b[n] = b[n] + q[i,1] * p[i,n];
For j = 2 Step 1 Until n Do Begin fun = 0;
For l = 1 Step 1 Until j-1 Do fun = fun+p[l,j-1]*q[i,l]; q[i,j] = fun;
For l = 2 Step 1 Until j Do If l < j Then Begin fun = 0;
For l = 1 Step 1 Until j-1 Do
fun = fun+q[i,l] * p[l,j-1]; q[i,j] = fun+q[i-1,j-1] End
Else q[i,j] = q[i-1,j-1];
For l = 0 Step 1 Until j-1 Do
b[n-1] = b[n-1] + p[j,n]*q[i+1,j] End; b[0] = 0;

```

```

For i = 0 Step 1 Until n-1 Do
  b[i] = b[i] + q[n-i,n] ; Goto final;
next: Comment The following is an alternative procedure for
      evaluation of the coefficients of the characteristic
      equation. It is based on Leverrier's method;

For i = 1 Step 1 Until n Do
  For j = 1 Step 1 Until n Do
    If i=j Then p[i,j] = 1 Else p[i,j] = 0;
  For l = 1 Step 1 Until n Do Begin
    For i = 1 Step 1 Until n Do
      For j = 1 Step 1 Until n Do Begin fu = 0;
      For k = 1 Step 1 Until n Do
        fu = fu + a[i,k] * p[k,j]; q[i,j] = fu End;
      For l = 1 Step 1 Until n Do
        For j = 1 Step 1 Until n Do
          p[i,j] = q[i,j]; S[i] = 0;
        For i = 1 Step 1 Until n Do
          S[i] = S[i] + p[i,i] End;
        b[0] = 1;
      For j = 1 Step 1 Until n Do Begin
        func = 0; fun = j;
        For i = 1 Step 1 Until j Do Begin
          beta = (-1)i-1; func = func + b[j-i] * S[i] * beta End;
        b[j] = func/func End next;
      For j = 1 Step 1 Until n Do
        b[j] = b[j]*((-1)j;
      final: End charac;

```

Program 207;

Procedure bairsto(n,a,eps0,eps1,eps2,eps3,K,m,x,y,nat,ex);

Value n; Integer n,K,m; Real eps0,eps1,eps2,eps3;

Integer Array ex,nat; Real Array x,y;

Comment This Birstow Hitchcock iteration is used to find successively pairs of roots of a polynomial equation of degree n with coefficients a[i] (i=0,1,...,n) where a[n] is the constant term. On exit from the procedure, m is the number of pairs of roots found, x[i] and y[i] (i=1,...,m) are a pair of real roots if nat[i]=1, the real and imaginary parts of a complex pair if nat[i]=-1, and ex[i] indicates which of the following conditions was met to exit from the iteration loop in finding this pair

- (1) Remainders r1,r0, become absolutely less than eps1.
- (2) Corrections, incrp, incrq, become absolutely less than eps2
- (3) The ratios, incrp/p, incrq/q, become absolutely less than eps3
- (4) The number of iterations become K. In the last case, the pair of roots found is not reliable and no further effort to find additional roots is made. The quantity eps0 is used as a lower bound for the denominator in the expressions from which incrp and incrq are found;

Begin Integer i,j,k,n1,n2,m1;

Array b,c[0:n+1]; Real p,q,r0,r1,s0,s1,v0,v1,det0,det1,det2,incrp,incrq,S,T;

For i = 0 Step 1 Until n Do
b[i] = a[i]; b[n+1] = 0; n2 = ENTIER((n+1)/2); n1 = 2*n2;

For m1 = 1 Step 1 Until n2 Do Begin

p = 0; q = 0;

For k = 1 Step 1 Until K Do Begin

step: For i = 0 Step 1 Until n1 Do

c[i] = b[i];

For j = n1-2, n1-4 Do Begin

For i = 0 Step 1 Until j Do Begin

c[i+1] = c[i+1]-p*c[i]; c[i+2] = c[i+2]-q*c[i] End End;

```

r0 = c[n1]; r1 = c[n1-1]; s0 = c[n1-2]; s1 = c[n1-3];
v0 = -q*s1; v1 = s0-s1*p; det0 = v1*s0-v0*s1;
If |det0| < eps0 Then Begin p = p+1; q = q+1;
  Goto step End;
det1 = s0*r1-s1*r0; det2 = r0*v1-r1*v0; incrp = det1/det0;
incrq = det2/det0; p = p+incrp; q = q+incrq;
If |r0| < eps1 Then Begin
  If |r1| < eps1 Then Begin ex[m1] = 1; Goto next End End;
  If |incrp| < eps2 Then Begin
    If |incrq| < eps2 Then Begin ex[m1] = 2; Goto next End End;
    If |incrp/p| < eps3 Then Begin
      If |incrq/q| < eps3 Then Begin ex[m1] = 3; Goto next End End End;
    ex[m1] = 4;
  next: S = -p/2; T = S2-q;
  If T < 0 Then Begin
    T = SQR(T); nat[m1] = 1; x[m1] = S+T; y[m1] = S-T End;
  If T > 0 Then Begin
    nat[m1] = -1; x[m1] = S; y[m1] = SQR(-T) End;
  If ex[m1] = 4 Then Goto out;
  For j = 0 Step 1 Until n1-2 Do Begin
    b[j+1] = b[j+1]-p*b[j]; b[j+2] = b[j+2]-q*b[j] End;
  n1 = n1-2;
  If n1 < 1 Then Begin
    out: m = m1; Goto final End;
  If n1 < 3 Then Begin
    m1 = m1+1; ex[m1] = 1; p = b[1]/b[0]; q = b[2]/b[0]; Goto next End End;
  i = i+1;
final: End;

```

Program 210;

```

Procedure ivector(a,rx,ix,rl,iml,n,eps,single);
Value n,a,rl,iml,eps; Integer n; Array a,ix,rx; Real rl,iml,eps;
Label single;
Comment This procedure evaluates the components of the
eigenvector of matrix a[i,j] assuming that the
first component is equal to 1. rl and iml are
the real and imaginary parts of the eigenvalue
for which the eigenvector is evaluated. rx[i]
and ix[i] are the real and imaginary parts
respectively of the ith component of the eigenvector;
Begin Integer i,j,m; Array d[1:2*n-1],b[1:2*n-1,1:2*n-1],c[1:2*n,1:2*n];
Mc Procedure invert(200,1,4);
If ~(iml=0) Then Goto other;
rx[1] = 1;
For i = 2 Step 1 Until n Do Begin
  d[i-1] = -a[i,1];
  For j = 2 Step 1 Until n Do
    b[i-1,j-1] = a[i,j] End;
  For i = 1 Step 1 Until n Do ix[i] = 0;
  For i = 1 Step 1 Until n-1 Do
    b[i,i] = b[i,i]-rl; invert(b,n-1,eps,singul);
  Goto exit; singul: Goto single;
exit: For i = 2 Step 1 Until n Do Begin rx[i] = 0;
  For j = 1 Step 1 Until n-1 Do
    rx[i] = rx[i] + b[i-1,j] * d[j] End; Goto final;
other: For i = 1 Step 1 Until n Do Begin
  rx[i] = 1; d[i] = iml;
  For j = 1 Step 1 Until n Do b[i,j] = a[i,j];
  b[i,i] = b[i,i] - rl End;
  invert(b,n,eps,singul);
  For i = 1 Step 1 Until n Do Begin
    ix[i] = 0; For j = 1 Step 1 Until n Do
      ix[i] = ix[i] + b[i,j]*d[j] End;
  final: End;

```

```

Program 213;
Procedure inverre(n,c,d,A,alarm); Value n,A,c; Integer n,alarm;
Array c,d,A;
Comment This procedure finds the extreme discrepancy matrix
d[i,j] for the inverse c[i,j] of a matrix knowing
the magnitude of the error bounds A[i,j] of the original
matrix. The extreme discrepancies are based on the element
which attains the largest absolute error. If alarm=0
at the exit from the procedure then the computation has
been regularly performed. If alarm=1 then the error was
too large and only the first order discrepancy is given;
Begin Integer i,j,k,l,i1,j1; Real f,s,f1,norm,norm1;
Array g,g1,g2,g3,g4,g5,i,T,T1,d[1:n,1:n];
No Procedure multip(202,1,4); f = 0; alarm = 0;
For i = 1 Step 1 Until n Do
  For j = 1 Step 1 Until n Do Begin
    g4[i,j] = g2[i,j] - g2[i,j] - 0; g1[i,j] = g[i,j] - I[i,j] - 0;
    T[i,j] = d[i,j] - d[i,j] - 0;
    For k = 1 Step 1 Until n Do
      For l = 1 Step 1 Until n Do Begin
        s = c[i,k]*c[l,j]; T1[k,l] = If s<0 Then -1 Else If s=0 Then 0 Else 1;
        d[i,j] = d[i,j] + T1[k,l] * A[k,l] * s End; f1 = d[i,j];
        If 'f' < 'f1' Then Begin f = f1;
        For i1 = 1 Step 1 Until n Do
          For j1 = 1 Step 1 Until n Do
            T[i1,j1] = T1[i1,j1] End End;
        For l = 1 Step 1 Until n Do
          For j = 1 Step 1 Until n Do Begin d[i,j] = 0;
          For k = 1 Step 1 Until n Do
            For l = 1 Step 1 Until n Do
              d[i,j] = c[i,k] * A[k,l] * T[k,l] * c[l,j] + d[i,j] End;
            For l = 1 Step 1 Until n Do
              For j = 1 Step 1 Until n Do
                A[i,j] = A[i,j] * T[i,j]; norm = 0; multip(n,g,A,c);
                For l = 1 Step 1 Until n Do Begin norm1 = 0;
                For j = 1 Step 1 Until n Do
                  norm1 = norm1 + |g[i,j]|; If norm < norm1 Then norm = norm1 End;
                If norm < 0.1 Then Goto final
                Else If 0.95 < norm Then Begin alarm = 1; Goto final End;
                multip(n,g1,g,g); multip(n,g2,g1,g); multip(n,g3,g2,g);
                multip(n,g4,g3,g);
                For l = 1 Step 1 Until n Do I[i,l] = 1;
                For l = 1 Step 1 Until n Do
                  For j = 1 Step 1 Until n Do
                    g5[i,j] = I[i,j] - g[i,j] + g1[i,j] - g2[i,j] + g3[i,j] - g4[i,j];
                    multip(n,d1,d,g5);
                For l = 1 Step 1 Until n Do
                  For j = 1 Step 1 Until n Do d[i,j] = d[i,j];
                final: End;

```

```

Program 214;
Procedure multerr(n,A,c,d,b,F); Value n,A,c,d,b; Integer n;
Array A,c,d,b,F;
Comment This procedure finds the bounds of error F[i,j] of the
product of two matrices A[i,j] * c[j,k] given the bounds of
error d[i,j] of c[i,j] and b[i,j] of A[i,j]. The extreme
discrepancy is found such as to maximize the largest element of F;
Begin Integer i,j,k,i1,j1,mr,mt;
Real f,f1; Integer Array r,t[1:n];
Array T,R,F1,T1,R1[1:n,1:n],AR,AT[1:n];
For i = 1 Step 1 Until n Do Begin
  For j = 1 Step 1 Until n Do
    T[i,j] = R[i,j] - F[i,j] - F1[i,j] - 0;
    r[i] = t[i] - 0 End; f = 0;
  For j1 = 1 Step 1 Until n Do Begin
    For l = 1 Step 1 Until n Do

```

```

For j = 1 Step 1 Until n Do Begin
  If t[j]=0 And r[i]=0 Then Begin
    For k = 1 Step 1 Until n Do Begin
      T[k,j] = If A[i,k] < 0 Then -1 Else If A[i,k]=0 Then 0 Else 1;
      R[i,k] = If c[k,j] < 0 Then -1 Else If c[k,j]=0 Then 0 Else 1;
      P[i,j] = P[i,j] + A[i,k] * T[k,j] * d[k,j] + b[i,k] * R[i,k] * c[k,j] *
      c[k,j] End End;
      f1 = P[i,j]; P[i,j] = 0; If f < f1 Then Begin
        f = f1; mt = j; mr = i;
      End End;
    For i1 = 1 Step 1 Until n Do Begin
      AT[i1] = T[i1,mt]; AR[i1] = R[mr,i1] End End End;
    For i = 1 Step 1 Until n Do Begin
      T[i,mt] = AT[i]; R[mr,i] = AR[i] End;
      t[mt] = r[mr] - 1; f = 0 End;
    For i = 1 Step 1 Until n Do
      For j = 1 Step 1 Until n Do Begin
        For k = 1 Step 1 Until n Do
          P[i,j] = P[i,j] + A[i,k] * T[k,j] * d[k,j] + b[i,k] * R[i,k] * c[k,j]
        End End;
      End End;

```

Program 215;

Procedure error(n,beta,discr,discim,da,g,d,e,m);

Value n; Integer n; Real beta,discr,discim;

Array da,g,d,e,m;

Comment This procedure evaluates the error bounds of the eigenvalues of a matrix of order n, given the error bounds da[i,j] of the matrix. beta is the value of the imaginary part of the eigenvalue for which the bounds are computed. discr and discim are the bounds of the real and imaginary parts of the eigenvalue respectively (always given in absolute value). g[i] and d[i] are the real and imaginary parts of the ith component of the eigenvector corresponding to the eigenvalue in question. e[i] and m[i] are the real and imaginary parts of the ith component of the eigenvector of the transpose matrix;

Begin Integer i,j; Real h,p,r,s,t,u,v,w,den,num1,num2;

Array o,r[1:n];

If beta=0 Then Goto other;

For i = 1 Step 1 Until n Do Begin o[i] = r[i] = 0;

For j = 1 Step 1 Until n Do Begin

o[i] = o[i] + da[i,j] * g[j];

r[i] = r[i] + da[i,j] * d[j] End End;

h = p = r = s = t = u = v = w = 0;

For i = 1 Step 1 Until n Do Begin

h = h + g[i] * e[i]; p = p + d[i] * m[i];

r = r + o[i] * e[i]; s = s + r[i] * m[i];

t = t + g[i] * m[i]; u = u + d[i] * e[i];

v = v + o[i] * m[i]; w = w + r[i] * e[i] End;

den = (h-p) * (h-p) + (t+u) * (t+u);

num1 = (h-p) * (r-s) + (t+u) * (v+w);

num2 = (h-p) * (v+w) - (t+u) * (r-s);

discr = |num1 / den|; discim = |num2 / den|;

Goto final;

other : t = 0; s = 0;

For i = 1 Step 1 Until n Do Begin o[i] = 0;

For j = 1 Step 1 Until n Do

o[i] = o[i] + da[i,j] * g[j] End;

For i = 1 Step 1 Until n Do Begin

t = t + o[i] * e[i]; s = s + g[i] * e[i] End;

discr = |t/s|; discim = 0;

final: End;