

**SIGNAL SPACE APPROACH
TO
LINEAR SYSTEM IDENTIFICATION**

by

H. B. Kekre, B. E. Hons; M. Tech.

**A thesis submitted to the Faculty of Science
in partial fulfilment of the requirement for the
degree of Master of Science.**

**Department of Electrical Engineering,
Faculty of Pure and Applied Science,
The University of Ottawa,
Ottawa, Canada,
1965.**

ABSTRACT

In signal space, the system function $h(t)$ is visualized as a "signal vector" and is identified by its projections onto a suitably chosen orthonormal basis. The system function representation on orthonormal basis in signal space satisfies the integral squared error criteria.

Instrumentation, by using analog computer, is suggested to evaluate the system function coefficients for linear time invariant systems. The procedure can be easily modified to cover the case of time varying systems.

Behavior of the linear system in signal space is analogous to the linear operator in vector algebra. The numerical representative of the "operator" is shown to be a matrix. Instrumentation to evaluate the elements of the system matrix is suggested for linear time invariant and time varying systems.

ACKNOWLEDGEMENTS

The author expresses his profound gratitude to Professor G. S. Glinski who suggested and supervised this project.

The author would like to acknowledge the financial assistance, for graduate studies in Canada, from the Commonwealth Scholarships and Fellowships program of the Government of Canada, and the Indian Institute of Technology, Bombay, India.

TABLE OF CONTENTS

	Page
Abstract	i
Acknowledgements	ii
<u>Chapter I</u>	
1.1 Formulation of Problem	1
1.2 Impulse Response	2
1.3 Step Response	3
1.4 Frequency Response	4
1.5 Signal Space Concepts	5
1.6 Differential Equation Approach	10
1.7 Damped Exponentials	12
1.8 Example	14
<u>Chapter II</u>	
2.1 Choice of an Error Criterion	17
2.2 Integral Squared Error Criterion	17
2.3 Orthonormal Components	23
2.4 Kautz Procedure	26
2.5 Construction of Basis	28
(a) Real Poles	28
(b) Complex Poles	30
2.6 Analog Simulation of Orthonormal Filters	35
<u>Chapter III</u>	
3.1 Construction of System Model	39
3.2 Generation of Growing Exponentials	44

3.3	System Representation in Signal Space	52
3.4	Identification of Matrix Elements of System Matrix	54
3.5	Modifications for Time Varying Systems	59
<u>Chapter IV</u>		
4.1	Computational Error	61
4.2	Examples	63
	Example - 1	63
	Example - 2	70
	Example - 3	72
	Example - 4	74
	Example - 5	76
	Conclusion	81
	References	83
	Vita	86

CHAPTER I

1.1 Formulation of Problem

The problem of identifying a black box - that is determining its input-output relationships by experimental means - occurs under different guises in various branches of science. Wiener¹, Bose² and Lee³ refer to this as a characterization problem. Some others call it the evaluation problem, whereas Zadeh⁴ and others^{5,6} refer to it as identification problem. The terms characterization and identification seem to be more popular. In this work we shall refer to it as "Identification Problem". The formulation of the problem can be stated as,

- Given -
- (1) A black box B whose input-output relationship is not known a priori.
 - (2) The input space (domain) of B i. e. the class of all time functions on which B operates is defined.
 - (3) A class of Black boxes U which on the basis of a priori information about B is known to contain B.
- Aim -
- By observing the response of B to inputs determine a member of U which is equivalent to B in the sense that its responses to all the time functions in the input space of B, within a certain degree of accuracy under some specified error criteria, are the same.

From the above statement it is obvious that the system identification is mainly concerned with the study of input

and output signals. This study is mainly dominated by the concepts associated with the time and frequency domains and the transforms from one domain to another. Some of the well known techniques are summarized below. In the following discussion only the linear time invariant systems are considered. System is assumed to be initially relaxed i. e. initial conditions are assumed to be zero.

1.2 Impulse Response

The input of a linear time-invariant system $r(t)$ and output $c(t)$ are related to each other by

$$c(t) = \int_{-\infty}^t h(t-\tau) r(\tau) d\tau \quad (1.1)$$

where $h(t)$ is the system function.

If $r(t) = \delta(t)$

where $\delta(t)$ is a unit impulse at time $t = 0$

$$\begin{aligned} \text{then } c(t) &= \int_{-\infty}^t h(t-\tau) \cdot \delta(\tau) d\tau \\ &= h(t) \end{aligned} \quad (1.2)$$

Thus the response of the system to the unit impulse (Dirac delta function) is the system function itself. Hence the unit impulse is a probing signal since the response to any signal can be calculated from the response to the unit impulse. The response to the unit impulse, therefore, completely identifies a linear time-invariant system.

Although this seems to be an excellent probe, it has many disadvantages:

Unit impulse is a mathematical concept and it cannot be generated in practice. At most a very rough approximation to it is feasible. A large amplitude pulse for a very short duration may be used as an approximation to it. Although this kind of approximation works well with some electrical networks it has severe practical limitations with non electrical systems. For example if input is a flow of liquid to some chemical plant then the flow control valve has a limited play - and creating even a very rough approximation to an impulse is not possible.

1.3 Step Response

In equation (1.1), if the input to the system is a unit step $u(t)$

$$\begin{aligned} \text{where} \quad u(t) &= 1 && \text{for } t \geq 0 \\ &= 0 && \text{for } t < 0 \end{aligned}$$

substituting $r(t) = u(t)$ in equation (1.1) we get,

$$c(t) = \int_{-\infty}^t h(t-\tau) u(\tau) d\tau \quad (1.3)$$

$$\text{or} \quad c(t) = \int_0^t h(t-\tau) d\tau \quad (1.4)$$

This can be transformed to

$$c(t) = \int_0^t h(\tau) d\tau = A(t) \quad (1.5)$$

which shows that step response $A(t)$ is the integral of the system function $h(t)$. The system response to any other input signal $r(t)$ is given by⁷

$$c(t) = \int_{-\infty}^t A'(t-\tau) r(\tau) d\tau + A(0) r(t) \quad (1.6)$$

where $A'(t) = \frac{d}{dt} A(t)$

and $A(0) = \lim_{t \rightarrow 0^+} A(t)$

As is well known, the integration tends to smooth out high frequency components. Hence for practical purpose a vital information is lost. The step response has a great advantage of simplicity of generation of step and this method works very well with the systems of the order not more than two. For higher order systems, the response usually shows a dead zone before it starts growing up. The popular approach is to approximate this by a pure delay and a second order system, which sometimes leads to erroneous results.

1.4 Frequency Response

The most common method of identifying the transmittance dynamics of the process is frequency response: the gain and phase characteristics of the transfer function as a function of angular frequency ω measured in radians per second. The transfer function is given by

$$H(j\omega) = \frac{C(j\omega)}{R(j\omega)} \quad (1.7)$$

Where $R(j\omega)$ is a constant amplitude sinusoidal signal and relative amplitude and phase of resulting $C(j\omega)$ are measured. This is nothing but tracking the system transfer function $H(s)$

along $j\omega$ -axis of complex frequency s -plane. The ratio of the amplitudes of $C(j\omega)$ and $R(j\omega)$ is $|H(j\omega)|$ and the angle $\angle H(j\omega)$ measures the phase of $C(j\omega)$ with respect to $R(j\omega)$. The plot of $|H(j\omega)|$ and $\angle H(j\omega)$ against ω can be used to construct the system model by well known approximation techniques in frequency domain.

Although this method can give satisfactory results in many cases, it has a great disadvantage of being laborious and time consuming, since frequency response by its basic character is a steady state measurement. For slow systems a considerable time may be involved to attain a steady state. For some systems these signals themselves could be unacceptable because of practical difficulties of generating them. And finally Guillemin⁸ has shown that what seems to be a good approximation in the frequency domain can lead to totally unacceptable results in the time domain.

1.5 Signal Space Concepts

Zadeh^{9, 10} has pointed out that the transmission and filtering of signals are much more general concepts than they are appreciated in usual practice. The signals can be visualized in more general way through the concepts of geometry of orthogonal spaces, extensively studied by mathematicians. Furthermore, the spectral theory of functions and operators together with the notions of orthogonal vector spaces imbedded in Hilbert space, have been widely used by physicists to provide a practical solution to some of the problems arising in quantum mechanics. These powerful concepts have been introduced to the signal theory for practical solutions to common engineering problems such as signal analysis and system identification.

A general engineering problem of system identification is illustrated in Fig. 1.1, in which the operations of measurements and specifications bridge the gap between the physical reality and mathematical theory.

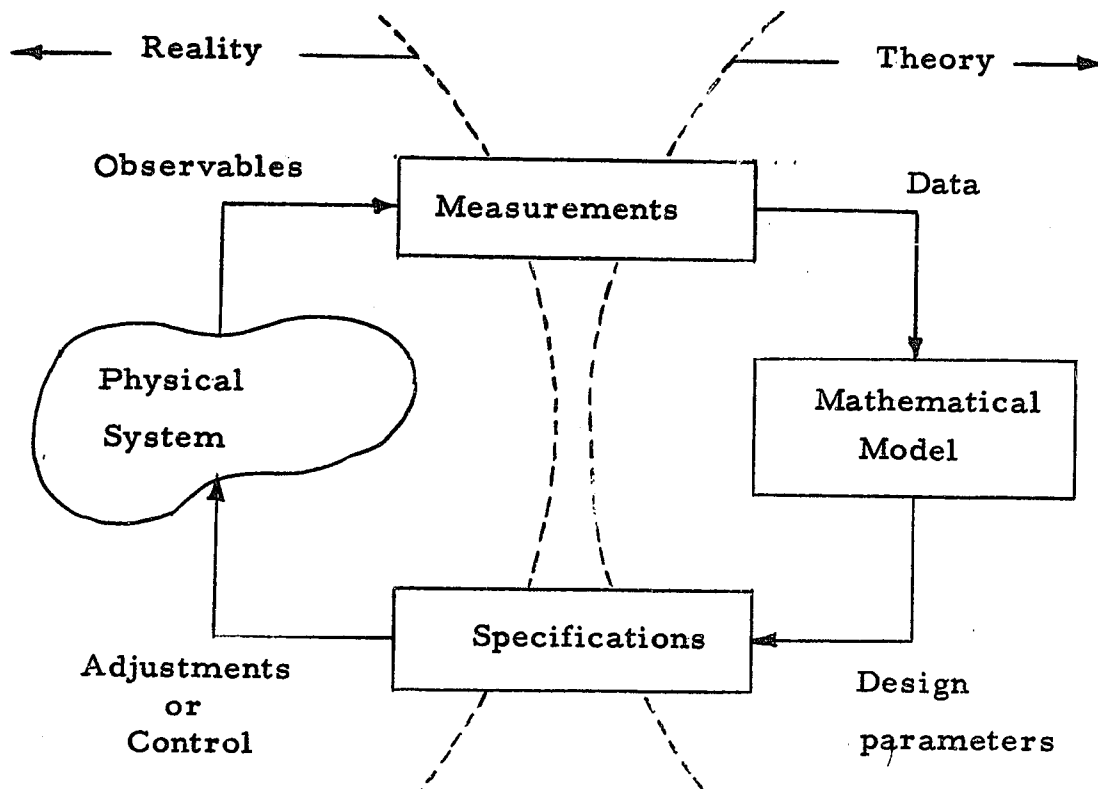


Fig. 1.1. - Physical System and Mathematical Model.

Lai¹¹ and Ross²⁶ have developed most appropriate notation for signal space, by making use of Dirac's notation extensively used in quantum mechanics. The notation has distinctive algebraic symbols to denote explicitly each of the following kinds of entities.

- 1) A physical observable (i. e. signal),
- 2) A physical measurement process
(i. e. pattern),

3) A physical system (i. e. operator).

The notation also provides the numerical representatives (i. e. scalar, column, row, matrix, function, etc.) for the above three basic entities. This establishes the correspondance between physical signals, linear transducers and measuring process in the real world and abstract vectors, linear operators and linear functionals of the mathematicians finite dimensional vector spaces. The meaning of the terms signal, pattern and operator is briefly presented here.

Signal Vectors

A signal is an entity which exists in the physical world. It may be a node voltage, branch current, shaft rotation, temperature, pressure, magnetic field intensity etc. The term "signal" refers to the entire behavior of some particular physical entity throughout one complete operation of the system of which the signal is a part i. e. throughout one complete operation of the signal generator. From this point of view the best way to describe the signal is to specify (or to identify) a generator of that signal. The generator is described as a device which converts magnitudes into signals.

Pattern Vectors

To incorporate physical measurement process in signal theory "pattern vectors" are considered. The term "measurement" refers to a device or process which produces a set of numbers or magnitudes at its output in response to a signal at its input. A common example of a measuring apparatus is a four terminal network followed by a sampler and perhaps also by a digitizer. Sometimes the term "siftor" is also used in place of "measuring apparatus" in order to have a term paralleling "generator".

Thus the generator converts magnitudes (or numbers) into the signals and a sifter converts signals into magnitudes (or numbers).

Each physical sifter in signal analysis problem is identified with some linear functional in the world of mathematics. In vector algebra a linear functional is defined as a linear mapping of the vectors, in some given vector space, into the field of scalars. Furthermore, the set of all such linear functionals is shown to be a vector space. This space of linear functionals is called the dual of the given vector space. More often than not the linear functional is referred as a vector. To avoid confusion a term pattern vector is used to distinguish it from signal vector.

From the above discussion it should be clear that, although patterns and signals are both vectors they are not at all the same. They do not even belong to the same space.

The symbol $| F \rangle$ is used to denote the signal vector in signal space, whereas a symbol of the form $\langle \tilde{G} |$ is used to denote the linear functional corresponding to some measuring apparatus. The tilde "~" over the letter G is meant to indicate transpose. That is if a numerical representative of signal $| F \rangle$ is a set of numbers arranged in a column then the numerical representation of $\langle \tilde{G} |$ will be a set of numbers arranged in a row.

The result obtained at the output of measuring apparatus $\langle \tilde{G} |$ in response to the signal $| F \rangle$ is denoted by a number $\langle \tilde{G} | F \rangle$. The process is indicated in the block diagram of Fig. 1.2.

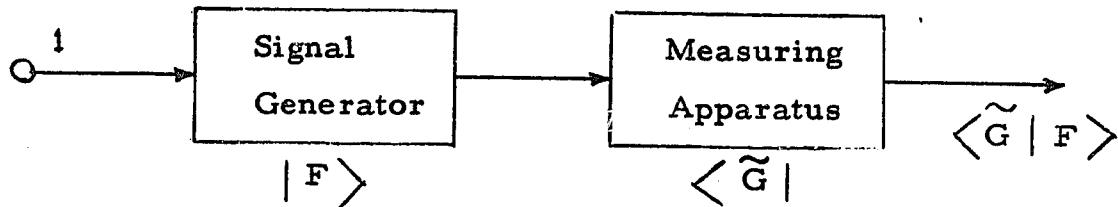


Fig. 1.2. - Correspondance between measurements process in signal theory and the inner product in vector algebra.

Operators

Corresponding to various linear physical devices and processes which yield an output signal in response to the input signal, e. g. RLC network, electromechanical transducer, linear distortion in communication channels, etc., the mathematical terms such as linear transformation, linear operator, or just "operator" are used. The operators are denoted by the symbol $|H|$. Thus the operator $|H|$ transforms the input signal vector $|R\rangle$ into a new output signal vector $|C\rangle$ whose magnitude and direction will be in general different from input signal vector. The operation is shown in Fig. 1.3.

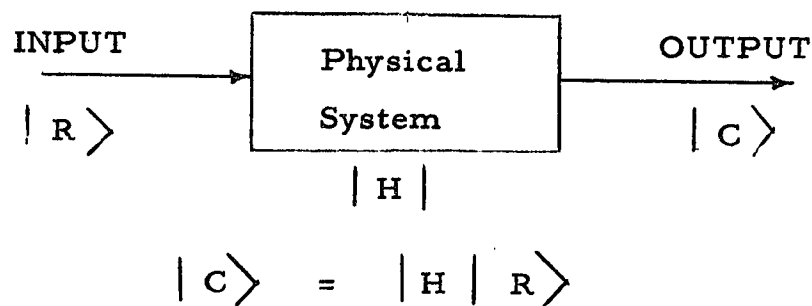


Fig. 1.3. - Physical System corresponds to "operator" in signal space.

Correspondance between the different entities in the reality considered by signal theory and the different entities in linear vector algebra is shown in the following table.

<u>Linear Algebra</u>	<u>Signal Theory</u>
Vector	Signal
Scalar	Magnitude
Norm square	Energy
Direction	Wave form (or spectrum)
Operator	Transducer or filter
Linear functional	Measuring apparatus
Coordinates	Measured results
Inner product	Measurement process

In general, signal space is infinite dimensional i. e. it has one to one correspondance with the Hilbert space. This means a complete representation of the signal can only be achieved by a complete basis which must contain an infinite number of members. But with real apparatus we can only measure a finite number of components and thus can represent a signal approximately allowing some latitude of error in some sense.

1.6 Differential Equation Approach

If $r(t)$ is the input to the system whose output is $c(t)$ then the differential equation representing the operation of the system assumes the form of a linear relation among the derivatives of $c(t)$ and $r(t)$. The equation can be written as

$$\sum_{n=0}^N a_n \frac{d^n}{dt^n} c(t) = \sum_{m=0}^M b_m \frac{d^m}{dt^m} r(t) \quad (1.12)$$

The order of the system is determined by the order of the highest derivative of $c(t)$ i. e. N and the realizability condition requires that $N \geq M$ and $a_N \neq 0$. Evaluation of the co-efficients can be broken down into two parts: evaluation of a 's and evaluation of b 's. For several reasons¹² evaluation of a 's is far more important and difficult aspect of the identification. More significantly, once a 's are known, the calculation of b 's is relatively a straight forward task, since b 's determine only the relative amplitudes of the terms in system response, while a 's determine the natural frequencies or the forms of time variations. The a 's are the coefficients in the characteristic polynomial of the system function. The zeros of the polynomial

$$a_n s^n + a_{n-1} s^{n-1} + \dots + a_1 s + a_0 = 0 \quad (1.13)$$

are the natural frequencies of the systems. If the zeros are denoted by $s_1, s_2, \dots, s_N$ the transient response of the system contains the terms of the form

$$e^{s_1 t}, e^{s_2 t}, \dots, e^{s_N t}$$

Thus the system function $h(t)$ takes the form

$$h(t) = \sum_{i=0}^N A_i e^{s_i t} \quad (1.14)$$

The fundamental parameters that best identify the system function $h(t)$ are the complex natural frequencies $\{s_i\}$. Unfortunately, experimental determination of these frequencies is

extremely difficult since the complex natural frequency parameters $\{s_i\}$ appear as exponents. When the damping factor is small, the usual frequency analysis may be used with some degree of accuracy, but with large damping factor, which is usually the case, the unique solution for s_i 's cannot be obtained in practice, because of the large correlation existing between the different exponent terms. Once s_i 's are known, the determination of the amplitude factors A_i 's is relatively easy and simple.

1.7 Damped Exponentials

At this point we may pause and consider the whole issue of whether it is really essential that the complex natural frequencies be known with great precision. If not, then we may select another set of exponentials with complex frequencies $\{\lambda_k\}$ which will provide an acceptable approximation to the original set. This is possible, because a collection of highly damped exponentials exhibit a very strong correlation with one another. For instance a correlation coefficient φ_{fg} between two time functions $f(t)$ and $g(t)$ over the time interval $[0, \infty)$ is defined as¹³

$$\varphi_{fg} = \frac{\int_0^{\infty} f(t) \cdot g^*(t) dt}{\left(\int_0^{\infty} f(t) \cdot f^*(t) dt \int_0^{\infty} g(t) \cdot g^*(t) dt \right)^{\frac{1}{2}}} \quad (1.15)$$

Where "*" denotes complex conjugate.

If $f(t)$ and $g(t)$ are considered as vectors $|f\rangle$ and $|g\rangle$ in signal space, and their inner product defined

$$\langle \tilde{g} | f \rangle = \int_0^{\infty} f(t) g^*(t) dt \quad (1.16)$$

then

$$\gamma_{fg} = \frac{\langle \tilde{g} | f \rangle}{[\langle \tilde{f} | f \rangle \cdot \langle \tilde{g} | g \rangle]^{\frac{1}{2}}}$$

$$= \cos \theta \quad (1.17)$$

where θ is the angle between the two vectors $|f\rangle$ and $|g\rangle$. When there is no correlation between two time functions $f(t)$ and $g(t)$, the correlation coefficient $\gamma_{fg} = 0$. In this case $\cos \theta = 0$ i. e. $\theta = 90^\circ$. The two time functions $f(t)$ and $g(t)$ are said to be orthogonal to each other over the semi infinite period and in this case one function is useless to approximate the other or vice versa. Considering the vectors $|f\rangle$ and $|g\rangle$ in signal space, their orthogonality means that the projection of one vector onto the other is zero and hence no component exists in that direction. In the case of damped exponentials, e^{pt} and e^{st} , the correlation coefficient between them is given by

$$\gamma_{ps} = \frac{\int_0^\infty e^{pt} \cdot e^{s^*t} dt}{\left[\int_0^\infty e^{(p+p^*)t} dt \cdot \int_0^\infty e^{(s+s^*)t} dt \right]^{\frac{1}{2}}}$$

$$= \frac{[(p+p^*)(s+s^*)]^{\frac{1}{2}}}{p + s^*}$$

$$= \frac{2 [\operatorname{Re}(p) \cdot \operatorname{Re}(s)]^{\frac{1}{2}}}{p + s^*} \quad (1.18)$$

If either of the complex frequencies p or s happens to have the real part equal to zero then the correlation between the two

functions is zero and hence we conclude that pure sine wave is useless for approximating a damped wave over a semi infinite interval, and vice versa.

On the other hand when both components have negative real parts, large correlation exists. For example, properly chosen linear combination of $e^{-\alpha t}$ and $e^{-2\alpha t}$ could be used to approximate $e^{-1.5\alpha t}$ with extremely small integral square error of one part in 1200. In fact these two exponentials could be used to represent $e^{-\lambda\alpha t}$ for values of λ ranging from 0.66 to 3.0 with integral square error less than one part in one hundred, which is better than the slide rule accuracy.

1.8 Example

In fact the example cited above has been solved in a general way in the following manner.

Let a pole at $s = -\gamma$ be approximated by two poles at $s = -\alpha_1$, and $s = -\alpha_2$

and let $\gamma = n\alpha_1$ where $0 < n < \infty$

and $\alpha_2 = m\alpha_1$ where $0 < m < \infty$

Then the normalized integral square error is given by

$$\epsilon = \left[\frac{(n-1)(n-m)}{(n+1)(n+m)} \right]^2 \quad (1.19)$$

This gives $\epsilon = 1$ for $n = 0$,
and $n \neq m$

When $n = 0$ the pole is at $s = 0$ and hence damped exponentials are useless to approximate it, which proves our earlier contention. The second case i. e. $n \rightarrow \infty$, is a trivial case, since the function being approximated is

$$f(t) = \lim_{n \rightarrow \infty} e^{-n\alpha t}$$

which in the limit becomes zero everywhere except at $t = 0$, where it is equal to one. This function has zero energy which means that if it is the impulse response of a system, then its response will be zero everywhere for all the inputs.

The integral square error ϵ is plotted against n for various values of m in Fig. 1.4. In the example cited above it is not really necessary that there should be only one pole in the region spanned by preselected two exponentials. Even if there are two or more poles within the region, the same two exponentials should be sufficient to give a fairly good approximation.

Referring to Fig. 1.4, we may conclude at this point that a fairly small number of preselected damped exponentials which "span" the region in the left side of the complex frequency plane may provide a more or less good approximation to any set of exponentials having the frequencies within that range.

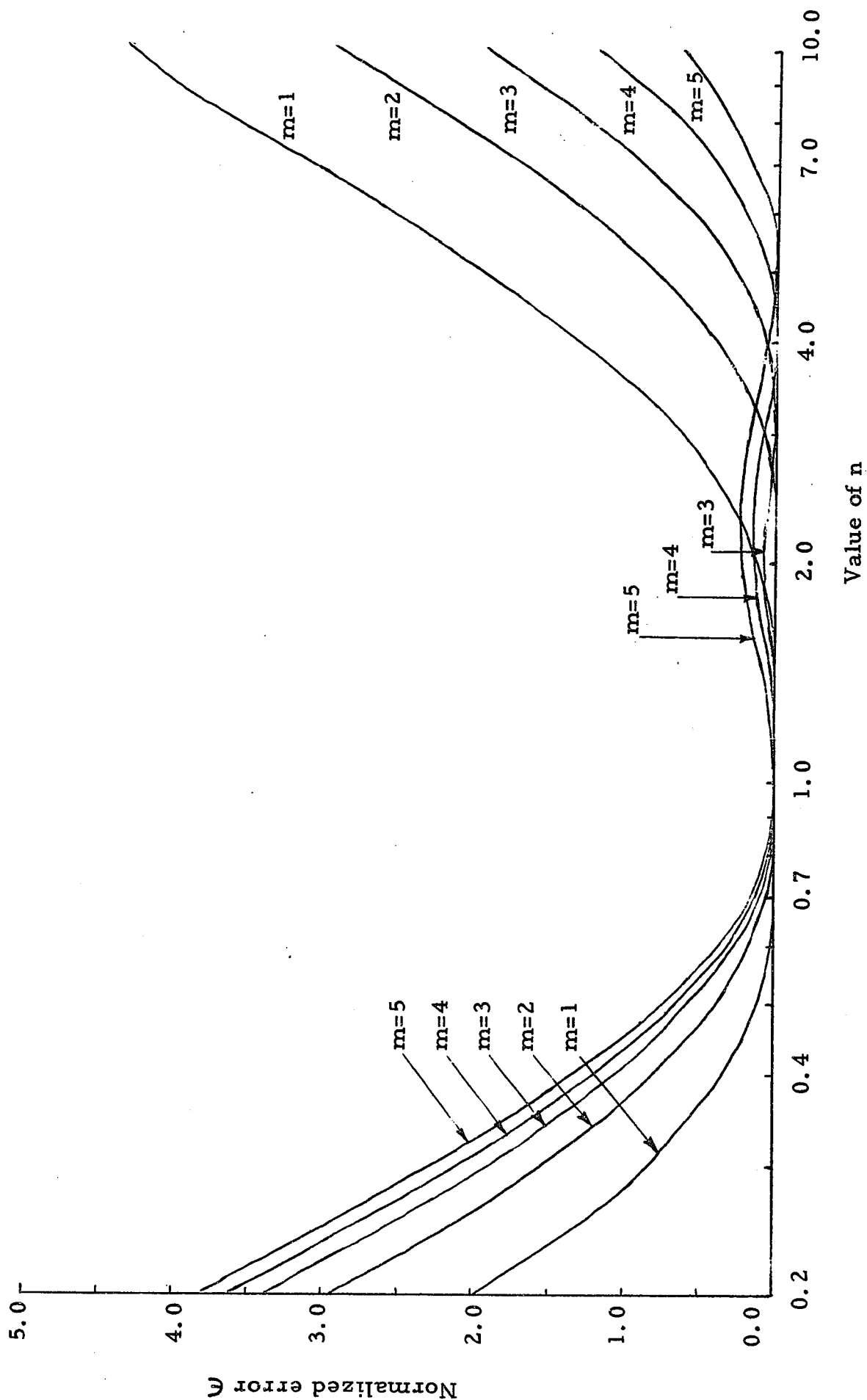


Fig. 1.4. Normalized error ϵ vs n for various values of m .

CHAPTER II

2.1 Choice of an Error Criterion

At the end of the Chapter I we have indicated that a set of preselected damped exponentials is sufficient for approximation of the linear system function. The approximation problem of the linear system synthesis may be stated as the determination of the realizable system function which closely approximates the prescribed system function. The approximation is to be achieved by minimizing the error specified by a certain error criterion. Though various procedures have been developed for solving the approximation problem, none of them yields simple calculations when the prescribed system is not specified in an elementary analytic form. This is particularly true when the prescribed system function is the impulse response and the problem is of time domain approximation¹⁴.

Several^{6, 15} different error criteria have been evolved to measure the goodness of approximation, each of which is appropriate in a certain class of problems. The criterion used here is the minimization of the integral of the squared error between the system function and a set of approximating functions. This criterion weighs large errors heavily and minimizes the area between the prescribed function and approximating functions. This criterion also makes it possible to represent the time functions as vectors in signal space which further simplifies the calculations and instrumentation.

2.2 Integral Squared Error Criterion

If the system function $h(t)$ is to be approximated by the finite sum, $h_a(t)$, of predetermined approximating functions $\{ a_i \phi_i(t) \}$

then

$$h(t) \cong h_a(t) = \sum_{i=1}^n a_i \phi_i(t) \quad (2.1)$$

where n is a finite number.

The approximation problem here has two aspects: the choice of functions $\phi_i(t)$ should be made to obtain the desired degree of accuracy with smallest number of terms and having once selected the functions $\phi_i(t)$, the coefficients a_i 's must be chosen to achieve the best approximation so that the integral

$$\epsilon = \int_0^{\infty} [h(t) - h_a(t)]^2 dt \quad (2.2)$$

is minimized. The infinite integration allows to define $h(t)$ over the whole range of time t . Substituting the value of $h_a(t)$, as defined in (2.1)

$$\epsilon = \int_0^{\infty} \left[h(t) - \sum_{i=1}^n a_i \phi_i(t) \right]^2 dt \quad (2.3)$$

or

$$\begin{aligned} &= \int_0^{\infty} [h(t)]^2 dt - 2 \int_0^{\infty} h(t) \sum_{i=1}^n a_i \phi_i(t) dt \\ &\quad + \int_0^{\infty} \left[\sum_{i=1}^n a_i \phi_i(t) \right]^2 dt \end{aligned} \quad (2.4)$$

Minimization of ϵ demands that ϵ should be finite.

This, in turn, restricts $h(t)$ and $\phi_i(t)$ to the condition that these functions should be squared integrable and the integration should be finite over the semi-infinite interval. In other words $h(t)$ and $\phi_i(t)$ should belong to L^2 space.

$$\begin{aligned} \text{i. e.} \quad & h(t) \in L^2 \\ \text{and} \quad & \phi_i(t) \in L^2 \quad i = 1, 2, \dots, n \end{aligned} \quad (2.5)$$

This is not a very severe restriction. All that it demands is that the system should be a finite energy system. It should contain some dissipative elements. Most of the practical systems fulfil this requirement.

The realizability condition demands that $\phi_i(t)$ should be zero for negative time and, in turn, the same restriction will hold good for $h(t)$.

$$\begin{aligned} \text{Thus} \quad & \left. \begin{aligned} h(t) &= 0 \\ \phi_i(t) &= 0 \end{aligned} \right\} \quad \text{for } t < 0 \end{aligned} \quad (2.6)$$

As already mentioned, the minimization of error has two aspects, first being a selection of $\{ \phi_i(t) \}$ such that the desired accuracy is obtained with a minimum number of terms. In general this cannot be done without prior knowledge of $h(t)$. In identification, very little knowledge of $h(t)$ is assumed and hence the selection of $\{ \phi_i(t) \}$ is more or less arbitrary. With some knowledge about $h(t)$, such as maximum and minimum time constants the selection of $\{ \phi_i(t) \}$ is greatly facilitated. Once the selection of $\{ \phi_i(t) \}$ is made, then the best $\{ a_i \}$ is selected to minimize ϵ in equation (2.3). The set $\{ a_i \}$ is to be calculated by determining stationary points of ϵ given by the conditions

$$\frac{\partial \epsilon}{\partial a_i} = 0 \quad i = 1, 2, \dots, n \quad (2.7)$$

Applying these conditions to equation (2.3) we get

$$\frac{\partial \epsilon}{\partial a_i} = \frac{\partial}{\partial a_i} \int_0^\infty \left[h(t) - \sum_{i=1}^n a_i \phi_i(t) \right]^2 dt \quad (2.8)$$

By the principle of variational calculus²⁴.

$$\begin{aligned} \frac{\partial \epsilon}{\partial a_i} &= \int_0^\infty \frac{\partial}{\partial a_i} \left[h(t) - \sum_{i=1}^n a_i \phi_i(t) \right]^2 dt \\ &= \int_0^\infty \left[h(t) - \sum_{i=1}^n a_i \phi_i(t) \right] \left[-\phi_i(t) \right] dt \\ &= - \int_0^\infty h(t) \cdot \phi_i(t) dt + \int_0^\infty \left[\sum_{i=1}^n a_i \phi_i(t) \right] \phi_i(t) dt \\ &\quad i = 1, 2, \dots, n \end{aligned} \quad (2.9)$$

substituting this in equation (2.7)

$$\int_0^\infty h(t) \cdot \phi_i(t) dt = \int_0^\infty \phi_i(t) \sum_{i=1}^n a_i \phi_i(t) dt \quad i = 1, 2, \dots, n \quad (2.10)$$

Equation (2.10) shows that one has to solve n simultaneous equations to work out the values of a_i 's and secondly, any change in the total number n , demands the whole procedure to be repeated again, i. e. addition or subtraction of one term in $\{ \phi_i(t) \}$ will lead, in general, to a different set $\{ a_i \}$. To simplify the matter, let us assume that $\{ \phi_i(t) \}$ is an orthonormal set. If it is not, it can always be orthonormalized, if it is a set

of independent functions. The orthonormality condition is:

$$\int_0^{\infty} \phi_i(t) \cdot \phi_j(t) dt = \delta_{ij} \quad (2.11)$$

where $\delta_{ij} = 0$ for $i \neq j$

and $\delta_{ij} = 1$ for $i = j$

Applying this condition to equation (2.10) gives

$$a_i = \int_0^{\infty} h(t) \cdot \phi_i(t) dt \quad (2.12)$$

The coefficients a_i 's are called generalized Fourier coefficients of $h(t)$ with respect to the orthonormal set $\{\phi_i(t)\}$.

Substituting the values of a_i 's, as obtained from equation (2.12), into equation (2.3) we get the minimized error

$$\epsilon_{\min} = \int_0^{\infty} [h(t)]^2 dt - \sum_{i=1}^n a_i^2 \quad (2.13)$$

Since by definition ϵ is a positive quantity and similarly both quantities on R.H.S. of equation (2.13), this leads to a very important and useful inequality

$$\int_0^{\infty} [h(t)]^2 dt \geq \sum_{i=1}^n a_i^2$$

This indicates that every added term in $\{\phi_i(t)\}$ further reduces the error and, when $\epsilon = 0$, the representation of $h(t)$ on $\{\phi_i(t)\}$ is said to be complete. Hence the remaining coefficients will be zero.

At this stage let us define the inner product of two time functions $g(t)$ and $f(t)$ by

$$\langle \tilde{g}(t) | f(t) \rangle = \int_0^{\infty} g(t) \cdot f(t) dt \quad (2.15)$$

Hence for real time functions

$$\langle \tilde{g}(t) | f(t) \rangle = \langle f(t) | g(t) \rangle \quad (2.16)$$

Using this notation, equation (2.12) becomes

$$a_i = \langle h(t) | \phi_i(t) \rangle \quad (2.17)$$

$$\text{or simply} \quad a_i = \langle h | \phi_i \rangle \quad (2.18)$$

and equation (2.1) can be written as:

$$\begin{aligned} h(t) &= h_a(t) + e_{\min}(t) \\ &= \sum_{i=1}^n a_i \phi_i(t) + e_{\min}(t) \end{aligned} \quad (2.19)$$

Where $e_{\min}(t)$ is minimized instantaneous error.

In Dirac notation

$$|h\rangle = \sum_{i=1}^n a_i |\phi_i\rangle + |e_{\min}\rangle \quad (2.20)$$

where $e_{\min}$ is very small we can write

$$|h\rangle \cong \sum_{i=1}^n a_i |\phi_i\rangle \quad (2.21)$$

Comparing equations (2.18) and (2.21) it is obvious that $|h\rangle$ is a vector in orthonormal signal space defined by the base vectors $\{| \phi_i \rangle\}$ and a_i 's are the projections of vector $|h\rangle$ along the directions of the base vectors $\{| \phi_i \rangle\}$. Since $| \phi_i \rangle$'s are preselected base vectors, they can be so chosen as to have an easy instrumentation. Knowledge of a_i 's will give the best approximate system satisfying the integral squared error criteria, and thus the system is identified.

The error $e_{\min}$ as in equation (2.20) is orthogonal to the space spanned by $\{| \phi_i \rangle\}$. This can be shown by operating the pattern $\langle \tilde{\phi}_k |$ on both sides of equation (2.20)

$$\langle \tilde{\phi}_k | h \rangle = \langle \tilde{\phi}_k | \sum_{i=1}^n a_i | \phi_i \rangle + \langle \tilde{\phi}_k | e_{\min} \rangle$$

$$k = 1, 2, \dots, n \quad (2.22)$$

this gives

$$\langle \tilde{\phi}_k | e_{\min} \rangle = 0 \quad \text{for } k = 1, 2, \dots, n \quad (2.23)$$

This shows that to reduce the error further, every new added element must be orthogonal to all previous vectors.

2.3 Orthonormal Components

As already indicated at the end of Chapter I, a linear combination of properly chosen damped sinusoids which spans the signal space can provide the best approximation to the signals which are bounded and contain a finite energy, i. e. they satisfy the condition,

$$\int_0^{\infty} [h(t)]^2 dt < \infty \quad (2.24)$$

In Dirac's notation the same condition can be written as

$$\langle \tilde{h} | h \rangle = (\| h \|)^2 \quad (2.25)$$

Thus the square of the norm of the vector in signal space corresponds to the energy contents of the signal.

We have already seen that damped sinusoids exhibit a high correlation and hence any change in the amplitude of one of them may be more or less compensated by suitable changes in the other components. This results in mathematical equations whose solutions are "soft"¹³ and are excessively sensitive to slight numerical errors. To eliminate these difficulties we form the linear combination of these elements, so that no correlation exists between the elements of a new set. Referring to equation (1.17), correlation coefficient between two signals $|g\rangle$ and $|f\rangle$ is given by

$$\begin{aligned} \psi_{fg} &= \frac{\langle \tilde{g} | f \rangle}{[\langle \tilde{f} | f \rangle \cdot \langle \tilde{g} | g \rangle]^{\frac{1}{2}}} \\ &= \cos \theta \end{aligned} \quad (2.26)$$

where $\psi_{fg} = 0$, $\theta = 90^\circ$ and two vectors are orthogonal to each other if $\langle \tilde{g} | f \rangle = 0$.

The process of constructing an orthonormal set from the set of linearly independent functions is known as "Gram-Schmidt Orthogonalization"^{16, 17} and is as follows.

Let $f_1, f_2, \dots, f_n$ be the set of linearly independent functions satisfying the condition

$$\sum_{i=1}^n c_i f_i = 0$$

iff each and every element of the set of scalars c_i is zero¹⁶. A new orthonormal set $\phi_1, \phi_2, \dots, \phi_n$ can be formed from these functions as shown

$$\begin{aligned}\phi_1 &= u_1 f_1 \\ \phi_2 &= u_2 (f_2 - a_{21} \phi_1) \\ \phi_3 &= u_3 (f_3 - a_{31} \phi_1 - a_{32} \phi_2) \\ &\vdots \\ \phi_j &= u_j (f_j - \sum_{k=1}^{j-1} a_{jk} \phi_k)\end{aligned}\tag{2.27}$$

The successive a_{jk} 's are chosen so as to make ϕ_j orthogonal to all previously constructed ϕ_k 's i. e.

$$\langle \tilde{\phi}_j | \phi_k \rangle = 0 \quad k = 1, 2, \dots, j-1\tag{2.28}$$

This gives

$$a_{jk} = \langle \tilde{f}_j | \phi_k \rangle\tag{2.29}$$

whereas u_j is chosen to normalize ϕ_j

$$\text{i. e.} \quad \langle \phi_i | \phi_j \rangle = 1\tag{2.30}$$

As it is, this procedure is fairly straightforward and general, but it is rather tedious. For our special case of interest, when the functions $\{f_i\}$ are damped exponentials, much simpler procedure is available by using frequency domain representation of the signals, which is due to Kautz¹⁸.

2.4 Kautz Procedure

Since this procedure is in s-domain, we will evaluate equivalent frequency domain representation of the inner product of two time functions $f(t)$ and $g(t)$.

Let $F(s)$ and $G(s)$ be the Laplace transforms of $f(t)$ and $g(t)$. From the complex convolution theorem¹⁹, Laplace transform of the product of two time functions is given by,

$$\int_0^{\infty} f(t) \cdot g(t) \cdot e^{-st} dt = \int_{-j\infty}^{j\infty} F(s-\lambda) \cdot G(\lambda) \frac{d\lambda}{2\pi j} \quad (2.31)$$

Setting $s = 0$ and changing λ to s

$$\int_0^{\infty} f(t) \cdot g(t) dt = \int_{-j\infty}^{j\infty} F(-s) G(s) \frac{ds}{2\pi j} \quad (2.32)$$

The L.H.S. of this equation is the inner product of two time functions $f(t)$ and $g(t)$. Let us define the expression on R.H.S. of this equation as the inner product of two complex frequency functions $F(s)$ and $G(s)$.

$$\text{Then} \quad \langle \widetilde{f(t)} \mid g(t) \rangle = \langle \widetilde{F(s)} \mid G(s) \rangle \quad (2.33)$$

$$\text{or simply} \quad \langle \widetilde{f} \mid g \rangle = \langle \widetilde{F} \mid G \rangle \quad (2.34)$$

$$\text{where} \quad \langle \widetilde{f} \mid g \rangle = \int_0^{\infty} f(t) \cdot g(t) dt \quad (2.35)$$

$$\text{and} \quad \langle \widetilde{F} \mid G \rangle = \int_{-j\infty}^{j\infty} F(-s) \cdot G(s) \frac{ds}{2\pi j} \quad (2.36)$$

It is easy to show that

$$\int_{-j\infty}^{j\infty} F(-s) G(s) \frac{ds}{2\pi j} = \int_{-j\infty}^{j\infty} G(-s) F(s) \frac{ds}{2\pi j} \quad (2.37)$$

thereby giving

$$\langle \tilde{F} | G \rangle = \langle \tilde{G} | F \rangle \quad (2.38)$$

Equation (2.38) is nothing but the frequency domain representation of

$$\langle \tilde{f} | g \rangle = \langle \tilde{g} | f \rangle \quad (2.39)$$

Let us use this result for construction of orthonormal basis $\{ \phi_k(t) \}$ whose Laplace transforms are given by $\{ \phi_k(s) \}$. Substituting $\phi_k(s)$ and $\phi_j(s)$ for $F(s)$ and $G(s)$ in equation (2.36) we get

$$\langle \tilde{\phi}_k | \phi_j \rangle = \int_{-j\infty}^{j\infty} \phi_k(-s) \cdot \phi_j(s) \frac{ds}{2\pi j} \quad (2.40)$$

This integral may be evaluated in terms of residues¹⁹ at the poles contained in an appropriate contour C that includes $j\omega$ -axis and a semicircle at infinity.

$$\begin{aligned} \text{Hence} \quad \langle \tilde{\phi}_k | \phi_j \rangle &= \int_C \phi_k(-s) \cdot \phi_j(s) \frac{ds}{2\pi j} \\ &= \sum \phi_k(-s) \phi_j(s) \end{aligned} \quad (2.41)$$

residues at poles
enclosed in C

From equation (2.41) it is clear that if the poles and zeros of $\phi_k(-s)$ and $\phi_j(s)$ are such that the product function becomes analytic in either left half or the right half of s -plane, then the integral must vanish and consequently $\phi_k(t)$ will be orthogonal to $\phi_j(t)$ in semi infinite interval $[0, \infty)$. Since $\phi_k(t)$ is real and has

finite energy, $\phi_k(s)$ will have its poles in the left half of s-plane and the poles will occur in complex conjugate pairs, in case they are complex. Hence $\phi_k(-s)$ will have its poles in the right half of s-plane at the mirror image of the poles of $\phi_k(s)$, reflected across jw-axis. Obviously if we select zeros of $\phi_j(s)$ at the poles of $\phi_k(-s)$, then the product function $\phi_k(-s) \cdot \phi_j(s)$ will have no poles in the right half of s-plane. This will make the product function analytic in the right half of s-plane and the integral in (2.41) will vanish.

2.5 Construction of Basis

(a) Real Poles

The construction of orthonormal basis hinted in the previous section will be clearer with the specific cases considered. Let us consider the construction procedure from the set of N poles located on the negative real axis of s-plane, as

$$F_n = \frac{1}{s + \alpha_n} \quad n=1, 2, \dots, N \quad (2.42)$$

where α_n is positive real constant.

The first component in the sequence will be

$$\phi_1(s) = u_1 \frac{1}{s + \alpha_1} \quad (2.43)$$

where $u_1 = \sqrt{2\alpha_1}$ is normalizing coefficient and is given by equation,

$$\langle \tilde{\phi}_1 | \phi_1 \rangle = 1$$

The second component must have a zero at α_1 to cancel the pole of $\phi_1(-s)$ in positive half of s-plane, and also must have one more pole from the set.

Hence,

$$\phi_2(s) = u_2 \cdot \frac{(s - \alpha_1)}{(s + \alpha_1)} \cdot \frac{1}{(s + \alpha_2)} \quad (2.44)$$

This is obviously orthogonal to $\phi_1(s)$ and $u_2 = \sqrt{2 \alpha_2}$ is given by

$$\langle \tilde{\phi}_2 | \phi_2 \rangle = 1$$

continuing the process, the n^{th} component is

$$\phi_n(s) = \frac{(s - \alpha_1) \dots (s - \alpha_{n-1})}{(s + \alpha_1)(s + \alpha_2) \dots (s + \alpha_{n-1})} \cdot \frac{1}{(s + \alpha_n)} \quad (2.45)$$

and normalizing coefficient $u_n = \sqrt{2 \alpha_n}$ is given by

$$\langle \tilde{\phi}_n | \phi_n \rangle = 1$$

This is orthogonal to all previous components,

$$| \phi_1 \rangle, | \phi_2 \rangle, \dots, | \phi_{n-1} \rangle$$

From this, the formation of every successive orthonormal component is obvious. The zeros of each $\phi_n(s)$ are located at the mirror image of the poles of the previous component $\phi_{n-1}(s)$. All the poles of $\phi_{n-1}(s)$ are retained. An additional pole at $s = -\alpha_n$ is added. The normalizing coefficient $u_n = \sqrt{2 \alpha_n}$.

It is obvious that

$$\langle \tilde{\phi}_j | \phi_k \rangle = \delta_{jk} \begin{cases} 0 & \text{for } j \neq k \\ 1 & \text{for } j = k \end{cases}$$

It is interesting to note that by selecting the poles on negative real axis at a regular interval such as $-\alpha, -3\alpha, -5\alpha, \dots$ we can generate complete orthonormal set of Legendre polynomials³, modified for the interval $[0, \infty)$.

This procedure does not restrict the selection of the pole different from the one used for the previous member. In fact, selecting the same pole again, and again, a set of Laguerre functions^{3, 20} can be generated. In this case only the order of the pole is increased. The n^{th} component of Laguerre functions for pole at $s = -\alpha$ is given by,

$$|\phi_n\rangle = \sqrt{2\alpha} \left(\frac{s-\alpha}{s+\alpha} \right)^{n-1} \cdot \frac{1}{s+\alpha} \quad (2.46)$$

(b) Complex Poles

Kautz procedure¹⁸ can be further extended to the case of complex poles. In case of real time functions, the complex poles can be introduced only in complex conjugate pairs. For completeness of the set, every time a complex conjugate pole pair is introduced, two real time functions must be generated.

If the poles are at $s_1, s_2 = -\alpha \pm j\beta$, Kautz¹⁸ has suggested the following pair.

$$|\phi_1\rangle = \sqrt{2\alpha} \frac{s + \sqrt{\alpha^2 + \beta^2}}{(s+\alpha)^2 + \beta^2} \quad (2.47a)$$

$$|\phi_2\rangle = \sqrt{2\alpha} \frac{s - \sqrt{\alpha^2 + \beta^2}}{(s+\alpha)^2 + \beta^2} \quad (2.47b)$$

It can be easily proved that $|\phi_1\rangle$ and $|\phi_2\rangle$ are orthogonal

to each other. This has an advantage of maintaining the same normalizing coefficients.

A pair of complex orthonormal exponentials which are more easily instrumented with analog equipment has been suggested by Huggins¹³ as follows.

$$|\psi_1\rangle = 2\sqrt{\alpha} \frac{s}{(s+\alpha)^2 + \beta^2} \quad (2.48a)$$

$$|\psi_2\rangle = 2\sqrt{\alpha(\alpha^2 + \beta^2)} \frac{1}{(s+\alpha)^2 + \beta^2} \quad (2.48b)$$

The plane in signal space spanned by $|\psi_1\rangle$ and $|\psi_2\rangle$ is the same as the plane spanned by $|\phi_1\rangle$ and $|\phi_2\rangle$. This will be more obvious if we look at their time domain transforms.

$$|\phi_1\rangle = R_1 e^{-\alpha t} \cos(\beta t - \theta_1) \quad (2.49a)$$

$$|\phi_2\rangle = R_2 e^{-\alpha t} \cos(\beta t + \theta_2) \quad (2.49b)$$

where

$$R_1, R_2 = \frac{\sqrt{8\alpha}}{\beta} \left[\alpha^2 + \beta^2 \mp \alpha\sqrt{\alpha^2 + \beta^2} \right]^{\frac{1}{2}}$$

and

$$\theta_1, \theta_2 = \tan^{-1} \frac{\sqrt{\alpha^2 + \beta^2} \mp \alpha}{\beta}$$

Considering the other pair as given in equation (2.48) we get

$$|\psi_1\rangle = \frac{2\sqrt{\alpha(\alpha^2 + \beta^2)}}{\beta} \cdot e^{-\alpha t} \cos(\beta t + \theta) \quad (2.50a)$$

$$|\psi_2\rangle = \frac{2\sqrt{\alpha(\alpha^2 + \beta^2)}}{\beta} \cdot e^{-\alpha t} \sin \beta t \quad (2.50b)$$

where

$$\theta = \tan^{-1} \frac{\alpha}{\beta}$$

Obviously the second pair shows a simplified time domain form. It is interesting to note that these signals are not 90° out of phase in time. This deviation is caused by a damping factor α . For $\alpha = 0$ obviously they are 90° out of phase. Furthermore, it can be easily seen that both these pairs are the particular cases of the set of damped sinusoids of the following form:

$$|f_1\rangle = R_1 e^{-\alpha t} \cos(\beta t + \theta_1) \quad (2.51a)$$

$$|f_2\rangle = R_2 e^{-\alpha t} \cos(\beta t + \theta_2) \quad (2.51b)$$

Where R_1 and R_2 are the normalizing coefficients and out of θ_1 and θ_2 we can select only one, the other being fixed by the orthogonality condition. Since variation in θ only changes the location of zero in s-plane, a general representation in frequency domain can be considered as follows:

$$|F_1\rangle = \frac{\rho_1 (s - \gamma_1)}{(s + \alpha)^2 + \beta^2} \quad (2.52a)$$

$$|F_2\rangle = \frac{\rho_2 (s - \gamma_2)}{(s + \alpha)^2 + \beta^2} \quad (2.52b)$$

where ρ_1, ρ_2 are normalizing coefficients. Then, inner product of $|F_1\rangle$ and $|F_2\rangle$ is given by

$$\langle \tilde{F}_1 | F_2 \rangle = \int_{-j\infty}^{j\infty} F_1(-s) F_2(s) \frac{ds}{2\pi j} \quad (2.53)$$

or

$$\langle \tilde{F}_1 | F_2 \rangle = \rho_1 \rho_2 \int_C \frac{(-s - \gamma_1)}{(-s + \alpha)^2 + \beta^2} \cdot \frac{(s - \gamma_2)}{(s + \alpha)^2 + \beta^2} \frac{ds}{2\pi j} \quad (2.54)$$

Solving this by residue theorem,

$$\langle \tilde{F}_1 | F_2 \rangle = \rho_1 \rho_2 \frac{\alpha^2 + \beta^2 + \gamma_1 \gamma_2}{4\alpha(\alpha^2 + \beta^2)} \quad (2.55)$$

Thus the orthogonality requires that

$$\rho_1 \rho_2 \cdot \frac{\alpha^2 + \beta^2 + \gamma_1 \gamma_2}{4\alpha(\alpha^2 + \beta^2)} = 0$$

giving the condition

$$\gamma_1 \gamma_2 = -(\alpha^2 + \beta^2) \quad (2.56)$$

Since α and β are real numbers and the imaginary values of γ_1 and γ_2 are not permissible for real time functions,

γ_1 and γ_2 must be real and have opposite signs, and hence they will not be in the same half of the complex frequency plane.

The normalizing coefficients ρ_1 and ρ_2 are given by

$$\rho_1 = \left[\frac{4(\alpha^2 + \beta^2)}{\alpha^2 + \beta^2 + \gamma_1^2} \right]^{\frac{1}{2}} \quad (2.57a)$$

$$\rho_2 = \left[\frac{4(\alpha^2 + \beta^2)}{\alpha^2 + \beta^2 + \gamma_2^2} \right]^{\frac{1}{2}} \quad (2.57b)$$

Kautz components, as given by equation (2.47 a, b), correspond to the choice $\gamma_1, \gamma_2 = \pm \sqrt{\alpha^2 + \beta^2}$ which makes

the normalizing factors of two functions equal, while the modification employed by Huggins, as given by (2.48 a, b), corresponds to the choice $\gamma_1 = 0$ and $\gamma_2 = \infty$ which makes one function, the time derivative of the other.

Now, supposing that we want to introduce a complex conjugate pole pair for $(n+1)^{\text{th}}$ and $(n+2)^{\text{nd}}$ components in the orthonormal set, then

$$\phi_{n+1} = \rho_1 \cdot \frac{(s - \alpha_1) \dots (s - \alpha_n)}{(s + \alpha_1) \dots (s + \alpha_n)} \cdot \frac{(s - \gamma_1)}{(s + \alpha_{n+1})^2 + \beta^2} \quad (2.58)$$

and

$$\phi_{n+2} = \rho_2 \cdot \frac{(s - \alpha_1) \dots (s - \alpha_n)}{(s + \alpha_1) \dots (s + \alpha_n)} \cdot \frac{(s - \gamma_2)}{(s + \alpha_{n+1})^2 + \beta^2} \quad (2.59)$$

where

$$\gamma_1 \cdot \gamma_2 = -(\alpha_{n+1}^2 + \beta^2)$$

and ρ_1, ρ_2 are given by equations (2.57 a, b). Next component in the sequence is formed by introducing a real pole at $s = -\alpha_{n+3}$ will be

$$\phi_{n+3} = \sqrt{2\alpha_{n+3}} \cdot \frac{(s - \alpha_1) \dots (s - \alpha_n)}{(s + \alpha_1) \dots (s + \alpha_n)} \cdot \frac{(s - \alpha_{n+1})^2 + \beta^2}{(s + \alpha_{n+1})^2 + \beta^2} \cdot \frac{1}{s + \alpha_{n+3}} \quad (2.60)$$

Ross²¹ has proved that the generalized components obtained in the case of complex conjugate pole pair are in fact generalized and similar process can be applied to obtain generalized components from a pair of real poles.

2.6 Analog Simulation of Orthonormal Filters

Following the above mentioned procedure let us consider the construction of a set of orthonormal transfer functions $\{ \phi_n(s) \}$ to be obtained from the set of poles located on negative real axis at $s = -\alpha_1, -\alpha_2, \dots, -\alpha_n$. Since ϕ_1 has a single pole at $-\alpha_1$, ϕ_2 has two poles at $-\alpha_1$ and $-\alpha_2$, and ϕ_n has n poles at $-\alpha_1, -\alpha_2, \dots, -\alpha_n$; the analog computer simulation of the ϕ_n is most easily achieved by cascade connections of sections having transfer functions of the form

$$G_k = - \left(\frac{s - \alpha_k}{s + \alpha_k} \right) \quad k = 1, 2, \dots, n \quad (2.61)$$

If $E_o(s)$ is the transform of the input as shown in Fig. 2.1A, then the $E_o(s) \cdot \phi_k(s) (-1)^k$ are the transforms of the outputs, taken from the output terminals of the integrators inside the blocks G_k .

The outputs of the blocks are related to input by

$$E_k(s) = (-1)^k \cdot \frac{(s - \alpha_1) \dots (s - \alpha_k)}{(s + \alpha_1) \dots (s + \alpha_k)} E_o(s) \quad (2.62)$$

The computer circuits for blocks of Fig. 2.1A are shown in Fig. 2.1B.

For the case of complex poles the sections G_{n+1} , and G_{n+2} do not appear separately, since the output $e_{n+1}(t)$ is a complex function of time.

The two sections G_{n+1} and G_{n+2} are combined to give the transfer function

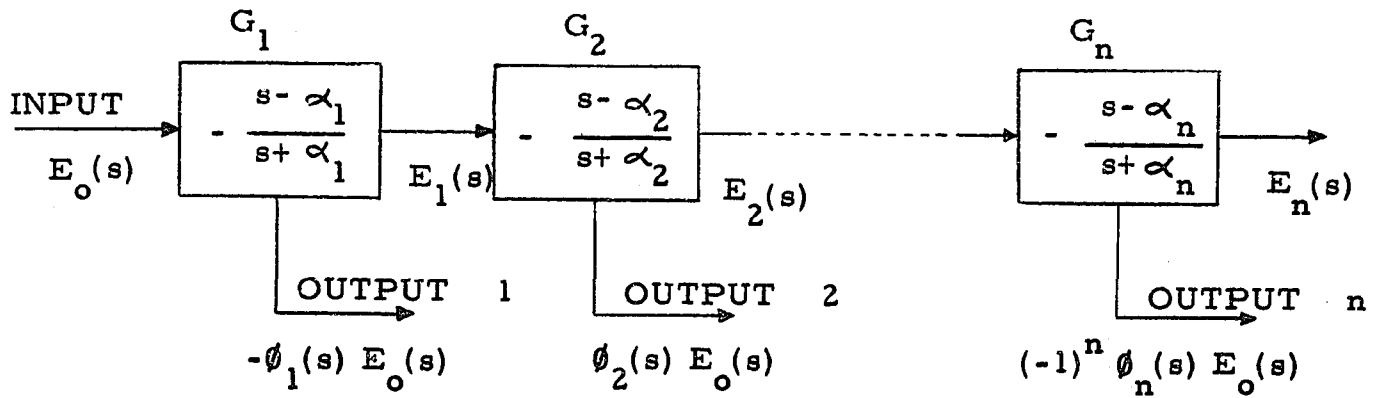


Fig. 2.1A. - General Scheme for realization of orthonormal filters $\{\phi_n(s)\}$

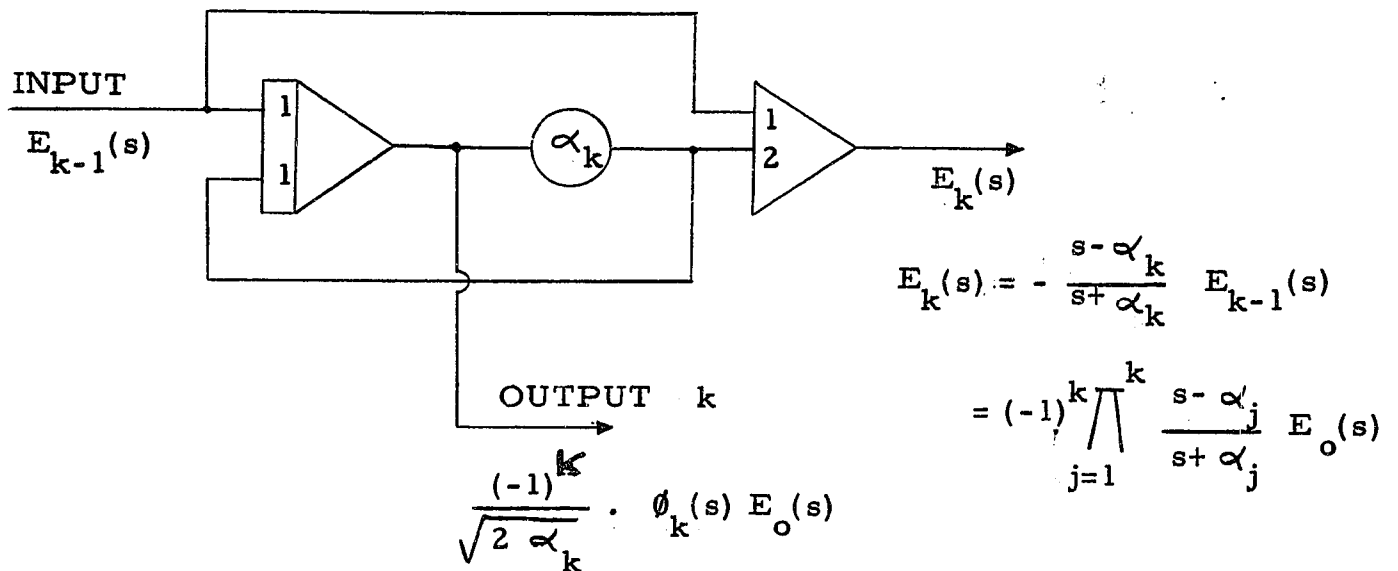


Fig. 2.1B. - Analog simulation of orthonormal filters $\{\phi_n(s)\}$ constructed from real poles. The output k multiplied by $(-1)^k \sqrt{2} \alpha_k$ will give the total transfer function between k INPUT and OUTPUT k terminals as $\phi_k(s)$.

$$G_{\substack{n+1 \\ n+2}} = \frac{(s - \alpha_{n+1})^2 + \beta^2}{(s + \alpha_{n+1})^2 + \beta^2} \quad (2.63)$$

as shown in Fig. 2.2A. The corresponding operational amplifier circuit is shown in Fig. 2.2B. The transfer functions obtained at the output terminals $(n+1)'$ and $(n+2)'$ are the components ψ_{n+1} and ψ_{n+2} given by equations (2.48 a, b) without the normalizing coefficients. From this the generalized components ϕ_{n+1} and ϕ_{n+2} as defined in equations (2.58) and (2.59), are obtained at terminals $(n+1)$ and $(n+2)$ as shown in Fig. 2.2B.

The next component ϕ_{n+3} obtained by adding a pole at $s = -\alpha_{n+3}$, can be introduced by adding similar section as shown in Fig. 2.1A and Fig. 2.1B.

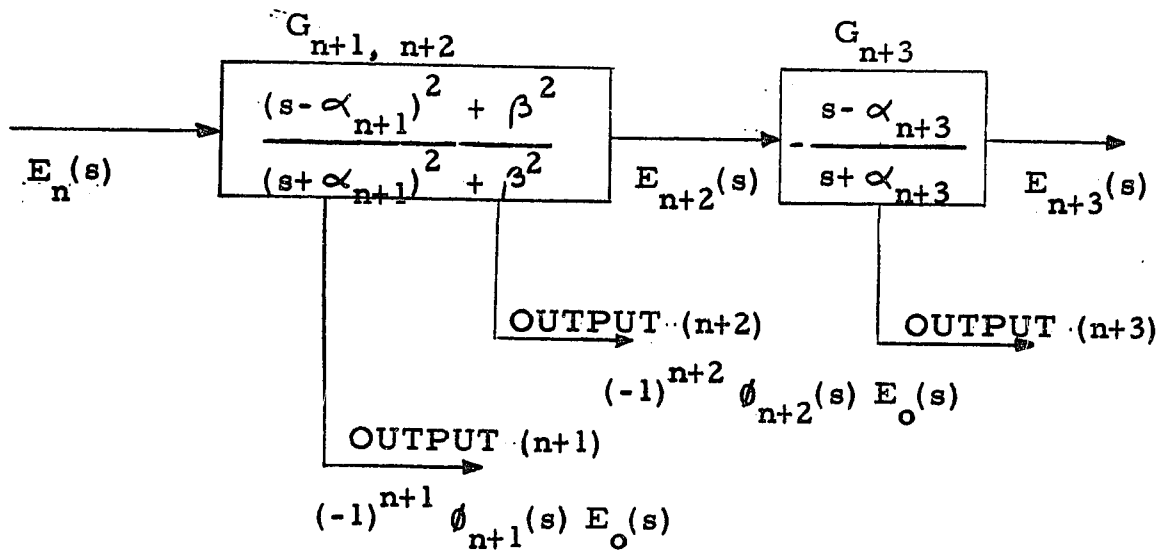


Fig. 2.2A. - Realization of Orthonormal filters with complex poles.

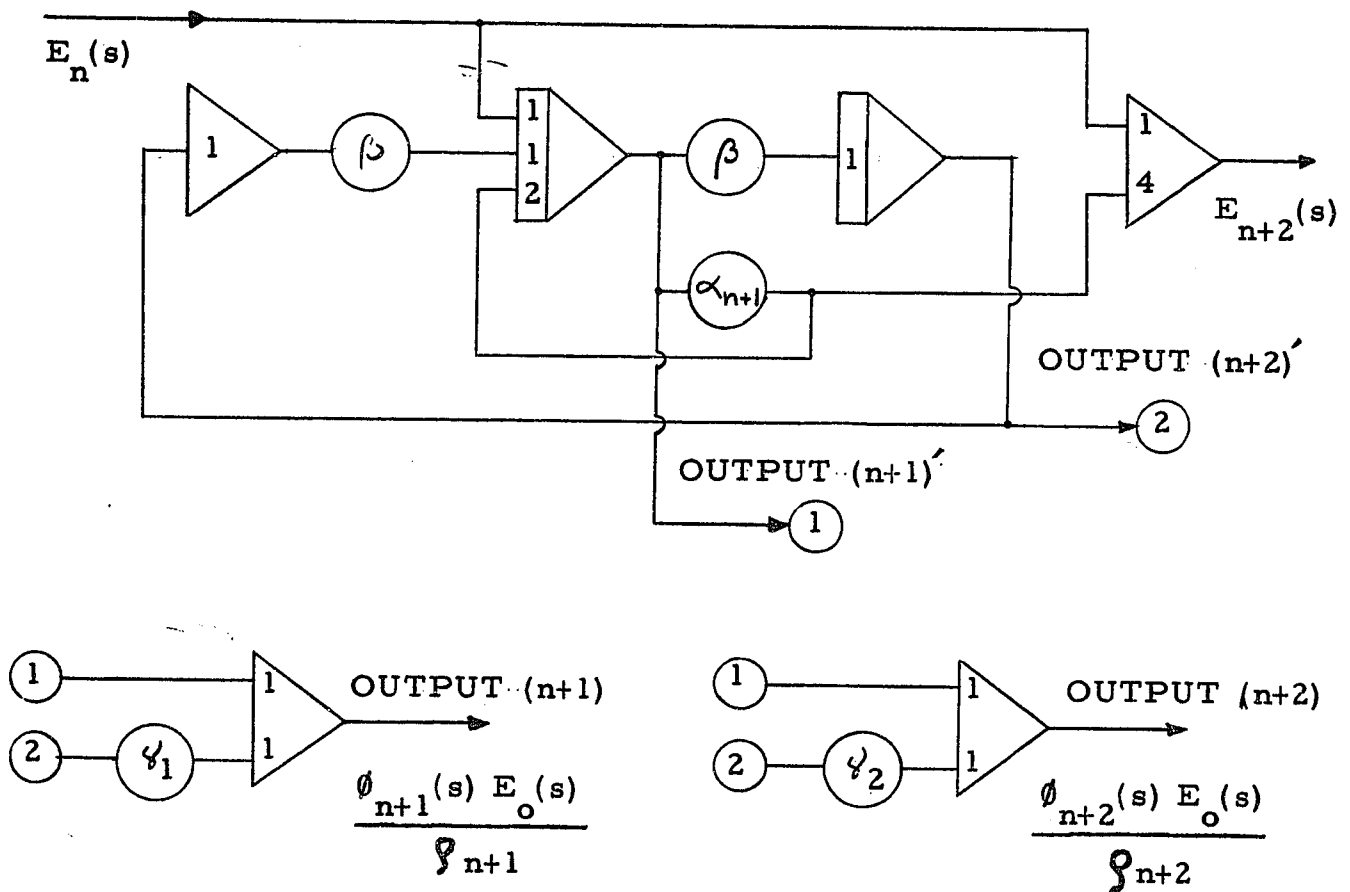


Fig. 2.2B. - Simulation of orthonormal filter with complex poles.

CHAPTER III

3.1 Construction of System Model

In Chapter I, it has been indicated that the system function $h(t)$, which is the response of the linear time invariant system to unit impulse, completely identifies the linear system. If we have prior knowledge of the system function $h(t)$, then the problem is that of synthesis. Thus, if $H(s)$ is the Laplace transform of $h(t)$, then $H(s)$ will be expressed as

$$H(s) = H_a(s) = \sum_{k=1}^n a_k \phi_k(s) \quad (3.1)$$

Thus if a_k 's are known, then the system model can be synthesized with orthonormal filters $\{ \phi_n(s) \}$ in parallel combination with their respective coefficients being set to a_k , where

$$a_k = \int_0^{\infty} h(t) \phi_k(t) dt \quad (3.2)$$

The arrangement of system model having the transfer function $H_a(s)$ is shown in Fig. 3.1. The realization of $\phi_k(s)$ was already indicated in previous chapters.

Now let us consider the scheme for computing a_k 's. Let $h(t)$ be the input applied to a filter having impulse response $\phi_k(t)$ as shown in Fig. 3.2.

The response $a_k(t)$ of the filter is given by a convolution integral

$$a_k(t) = \int_{-\infty}^t \phi_k(t - \tau) h(\tau) d\tau \quad (3.3)$$

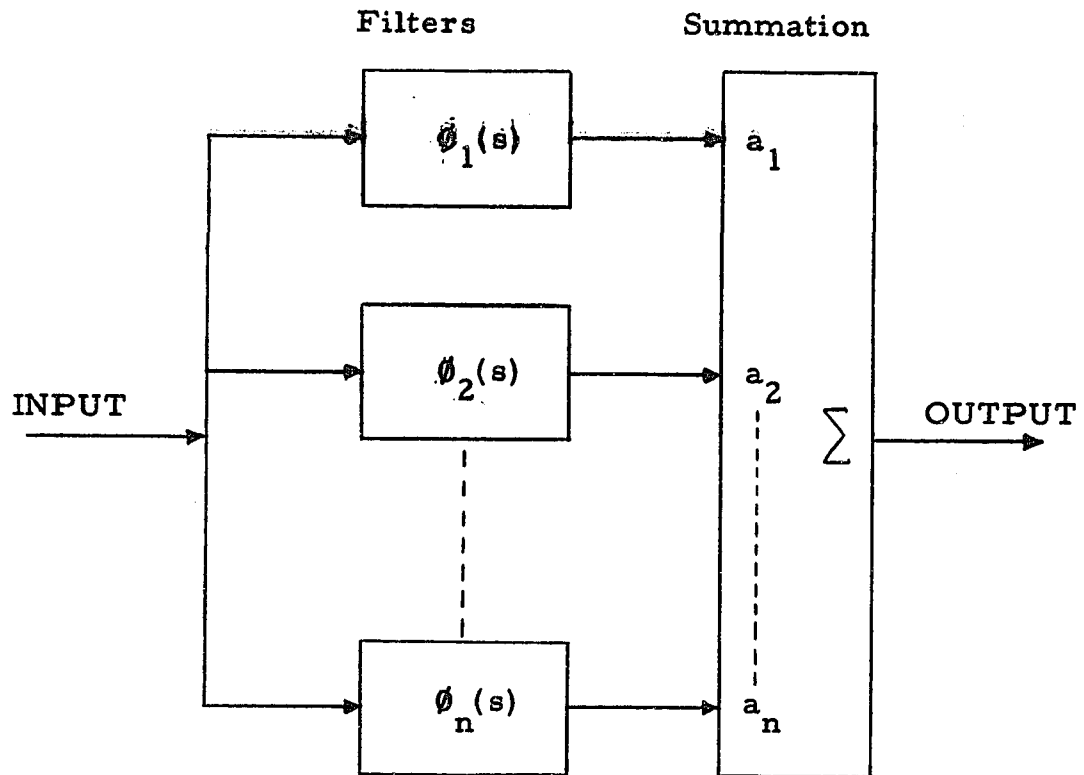


Fig. 3.1. - Realization of approximate system with the help of orthonormal filters.

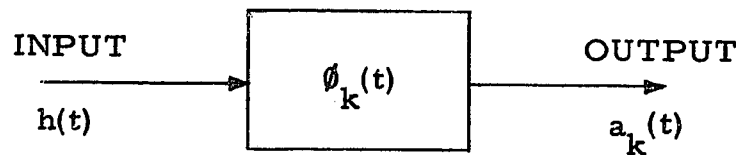


Fig. 3.2.

If we sample the output at time $t = 0$

$$\text{then} \quad a_k(0) = \int_{-\infty}^0 \phi_k(-\tau) h(\tau) d\tau \quad (3.4)$$

Changing the variable of integration τ to $-\tau$

$$a_k(0) = \int_0^{\infty} \phi_k(\tau) h(-\tau) d\tau \quad (3.5)$$

Comparing equation (3.5) with equation (3.2), it is seen that if in equation (3.5), $h(-\tau)$ is replaced by $h(\tau)$, then $a_k(0)$ is exactly the same as a_k . This means that we have to consider the time reversal of $h(t)$. The time reversed signal $h(-t)$ is applied to a filter with impulse response $\phi_k(t)$ and the output, sampled at time $t = 0$, will yield the component of vector $|h\rangle$ in the direction of the base vector $|\phi_k\rangle$. The procedure is indicated in Fig. 3.3.

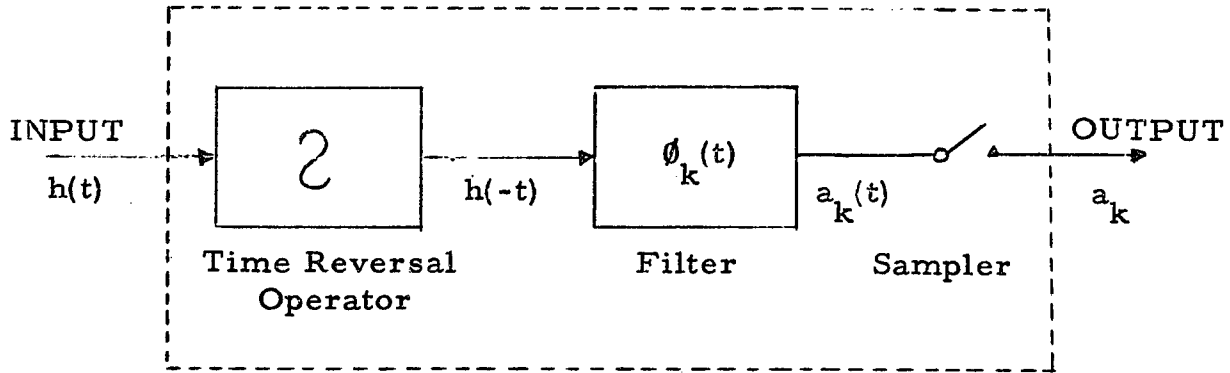


Fig. 3.3. Scheme to obtain the inner product of two signals $h(t)$ and $\phi_k(t)$.

Since a_k is the inner product of two vectors $|h\rangle$ and $|\phi_k\rangle$ in signal space, the combination of sections inside the dotted line in Fig. 3.3 can be called a pattern $\langle \phi_k |$ in pattern

space which corresponds to a vector $|\phi_k\rangle$ in signal space. When a vector $|h\rangle$ is sifted through the pattern $\langle\tilde{\phi}_k|$, it yields a number a_k which is the projection of vector $|h\rangle$ onto the base vector $|\phi_k\rangle$. The operation performed in Fig. 3.3 can be symbolically shown as in Fig. 3.4.

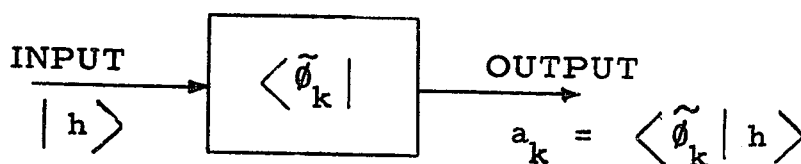


Fig. 3.4. Symbolic representation of instrumentation for inner product.

Now we are in a position to assign a proper meaning to the term "pattern". A measuring equipment designated by pattern $\langle\tilde{g}|$ performs the following operations on the input signal $|f\rangle = f(t)$.

- a) Time reversal of the input signal
 $|f\rangle = f(t)$
- b) Time reversed signal $f(-t)$ is transmitted through a filter having impulse response $g(t)$
- c) The output of the filter is sampled and recorded at time $t = 0$.

With the set of orthonormal filters $\{\phi_k(t)\}$ all the components of $|h\rangle$ onto the basis $\{|\phi_k\rangle\}$ can be evaluated provided the time reversal operation is possible. Huggins¹³ has

suggested a method of recording the signal $h(t)$ on the tape and running the tape in reversed direction. The signal read from this tape will be the time reversed signal $h(-t)$. But this can be done only if $h(t)$ is known. For identification, $h(t)$ is unknown, since the system is assumed to be a black box. In that case we can take the advantage of the fact that for real signals

$$\langle \tilde{\phi}_k | h \rangle = \langle \tilde{h} | \phi_k \rangle$$

That is, instead of performing time reversal of $h(t)$ we can perform the time reversal of $\{\phi_k(t)\}$ and then these time reversed signals $\{\phi_k(-t)\}$ can be applied as inputs to the system $h(t)$. Sampling the outputs at $t = 0$ will give the system coefficients a_k 's as shown in Fig. 3.5.

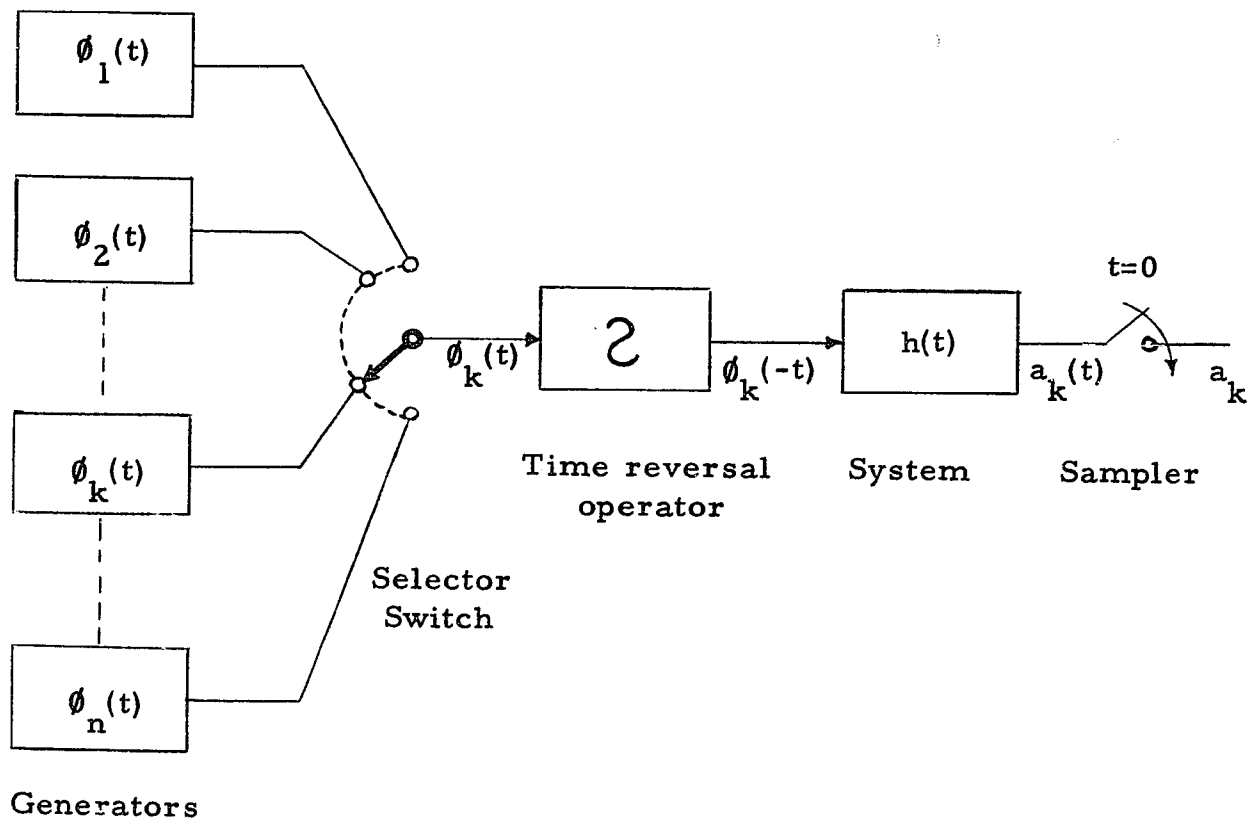


Fig. 3.5. Scheme for evaluating system coefficients a_k 's.

In any case, the time reversal operation seems to be unavoidable. The procedure suggested by Huggins is rather tedious. It will be shown in the next section that the time reversed signals $\phi_k(-t)$ can be easily generated on analog computer. Since $\phi_k(t)$ are the damped exponentials, the time reversed signals $\phi_k(-t)$ will be the growing exponentials.

3.2 Generation of Growing Exponentials

Let $f(t)$ be the time signal which is to be reversed in time. The function $f(t)$ satisfies the condition

$$f(t) = 0 \quad \text{for } t < 0 \quad (3.6)$$

Considering the bilateral Laplace transforms^{19, 22}, let $F(s)$ and $f(t)$ be the transform pair so that

$$F(s) = \int_{-\infty}^{\infty} f(t) \cdot e^{-st} dt \quad (3.7)$$

Applying the condition (3.6)

$$F(s) = \int_0^{\infty} f(t) \cdot e^{-st} dt \quad (3.8)$$

and

$$f(t) = \int_{-j\infty + \sigma}^{j\infty + \sigma} F(s) \cdot e^{st} \frac{ds}{2\pi j} \quad (3.9)$$

where σ is suitably chosen, so as to include all the poles of $F(s)$ to the left of Bromwich contour, as shown in Fig. 3.6A.

In order that the integrals (3.8) and (3.9) exist, in addition to the constraint that $f(t)$ be piecewise differentiable, it is sufficient that the function $f(t)$ belong to the class of "functions of exponential order"²²

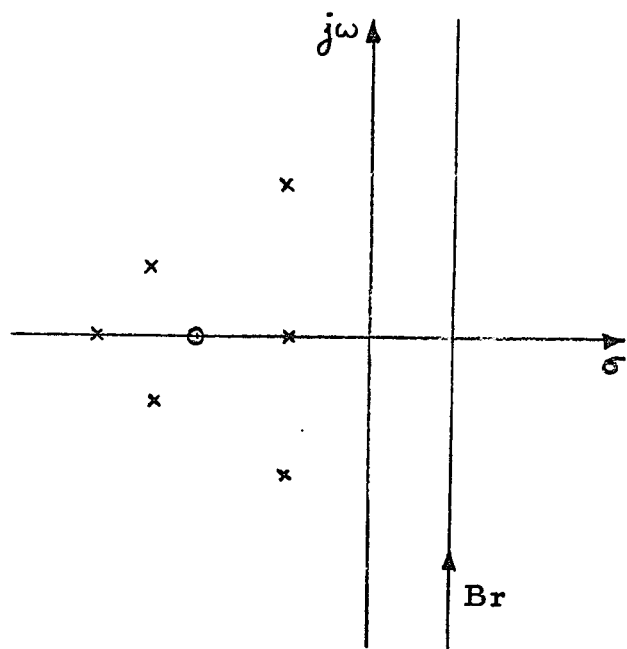


Fig. 3.6A. $F(s)$ in s -plane

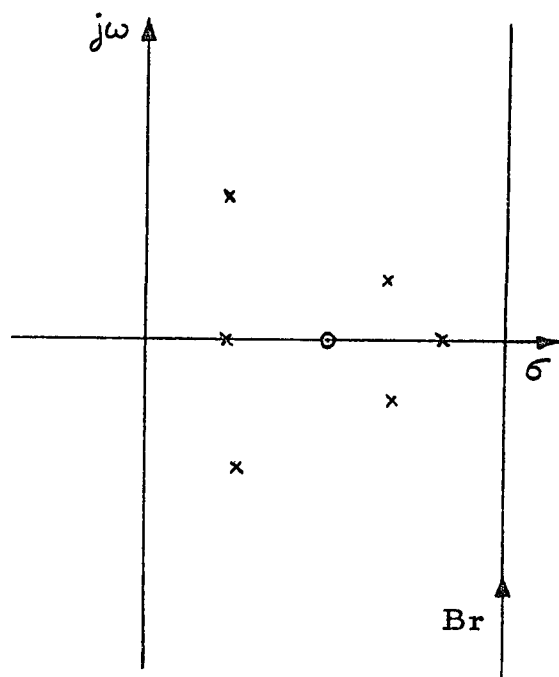


Fig. 3.6B. $F(-s)$ in s -plane

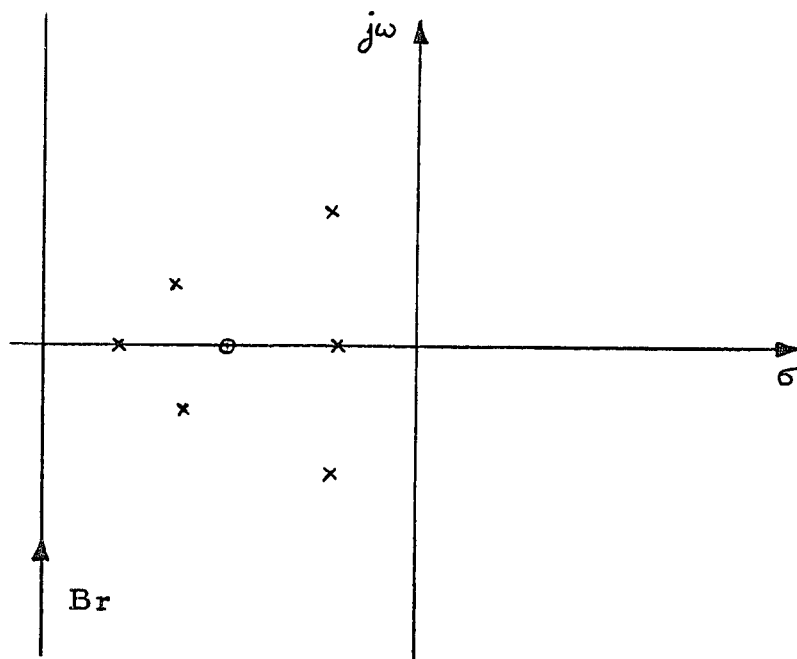


Fig. 3.6C. $F(\lambda)$ in λ -plane

and that $f(t)$ be bounded at $t \rightarrow 0^+$. A function is said to be of exponential order if

$$e^{-\sigma t} |f(t)| < M$$

where M is upper bound and σ is a real part of complex variable s . In our special case of damped exponentials all these conditions are easily satisfied.

Now let us derive the effect of changing the sign of complex variable s by finding the inverse transform of $F(-s)$. Using the equation (3.9) let

$$I = \int_{-j\infty + \sigma_1}^{j\infty + \sigma_1} F(s) e^{st} \frac{ds}{2\pi j} \quad (3.10)$$

The path of integration for evaluating I is defined by the Bromwich contour Br at $s = \sigma_1$, and is shown in Fig. 3.6B, σ_1 being selected so as to include all the poles of $F(-s)$. Changing the variable of integration s to $-\lambda$

so that $ds = -d\lambda$

and limits when $s \rightarrow -j\infty + \sigma_1$; $\lambda \rightarrow j\infty - \sigma_1$

and $s \rightarrow j\infty + \sigma_1$; $\lambda \rightarrow -j\infty - \sigma_1$

Substituting this in equation (3.10)

$$I = \int_{j\infty - \sigma_1}^{-j\infty - \sigma_1} F(\lambda) e^{-\lambda t} \frac{-d\lambda}{2\pi j} \quad (3.11)$$

which can be written as

$$I = \int_{-j\infty - \sigma_1}^{j\infty - \sigma_1} F(\lambda) \cdot e^{\lambda(-t)} \frac{d\lambda}{2\pi j} \quad (3.12)$$

The path of integration is shown in Fig. 3.6C. Comparing the Figures 3.6C to 3.6A it is seen that $F(\lambda)$ is the same as $F(s)$ except the Br-contour is shifted to the left of all the poles of $F(\lambda)$. Comparing equation (3.12) to equation (3.9) it is seen that they are the same except for the reversal of the sign of t . Hence we can write

$$I = -f(-t) \quad \text{for } t < 0 \quad (3.13)$$

This is a very important result which can be summarized as follows:

$$\begin{array}{lll} \text{if} & F(s) = \int_0^\infty f(t) e^{-st} dt & f(t) = 0 \quad \text{for } t < 0 \\ \text{then} & F(-s) = \int_0^\infty -f(-t) e^{-st} dt & f(-t) = 0 \quad \text{for } t > 0 \end{array}$$

This means that the image across the origin in time domain gives the image across the origin in frequency domain. This can be verified for our special case of damped exponentials.

$$\begin{aligned} \text{Let} \quad f(t) &= \sum_k A_k e^{-\alpha_k t} & \text{for } t \geq 0 \\ &= 0 & \text{for } t < 0 \end{aligned} \quad (3.14)$$

$$\text{then} \quad F(s) = \sum_k \frac{A_k}{s + \alpha_k} \quad (3.15)$$

$$\begin{aligned} \text{and} \quad F(-s) &= \sum_k \frac{A_k}{-s + \alpha_k} \\ \text{or} \quad F(-s) &= \sum_k \frac{-A_k}{s - \alpha_k} \end{aligned} \quad (3.16)$$

Taking the inverse transform

$$\begin{aligned} \mathcal{L}^{-1} F(-s) &= - \sum_k A_k e^{\alpha_k t} \\ &= -f(-t) \quad \text{for } t \leq 0 \end{aligned} \quad (3.17)$$

This proves that the growing exponentials can be generated provided we can get the computer set up for $\{\phi_k(-s)\}$. It has already been seen that the simulation of $\{\phi_k(s)\}$ is possible on analog computer. On analog computer set up s comes only in integrator units. All other units are independent of s . So, changing s to $-s$ means only an insertion of an inverter in series with every integrator at the specified time of the solution. Hence the time reversal operator is nothing but a switch which introduces inverters in the circuit in series with all integrators of the computer set up.

It is important to note here that if the time reversal is effected right at the beginning of the solution, then the time reversed function obtained will be the reflection of the negative past of the original function in which we are not interested. This is shown in Fig. 3.7A. To get the proper portion of the curve, the solution has to be reversed in time at a later time say $t = t_1$ as shown in Fig. 3.7B. Then the signal generated will be a useful

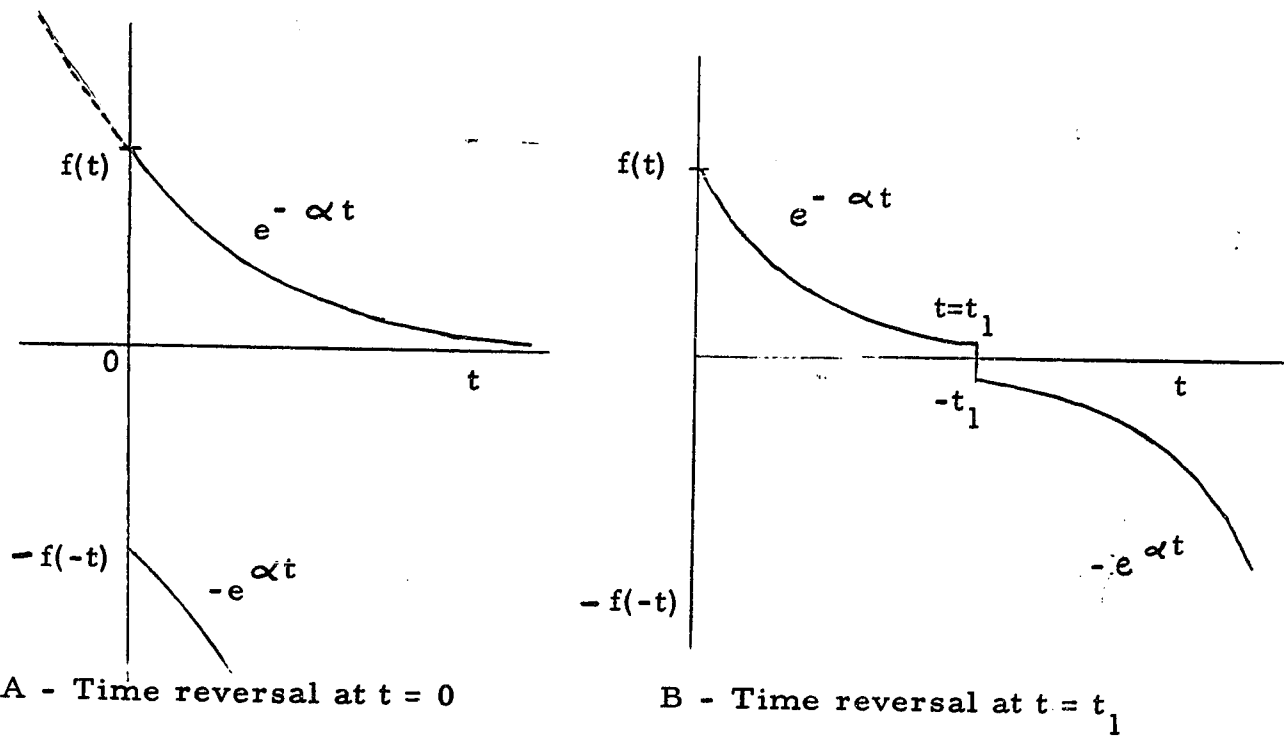


Fig. 3.7.

signal which can be used as a probing signal for identification of the system. The procedure for the evaluation of system coefficient a_k 's was already indicated in previous section.

Computer set up for the generation of first three components from the poles at $s = -\alpha_1, -\alpha_2, -\alpha_3$ is shown in Fig. 3.8. The time reversed signals obtained on computer are shown in Fig. 3.9. The time reversing switches are indicated in Fig. 3.8 and marked S_1, S_2, S_3 . They are to be operated synchronously at the required time t_1 . For negative time signals they assume position marked $-t$. In practice this can be achieved by bringing the computer to "hold" operation and then changing the switches from position $+t$ to $-t$. Bringing the computer back to "operation" will generate growing exponentials which are shown in Fig. 3.9.

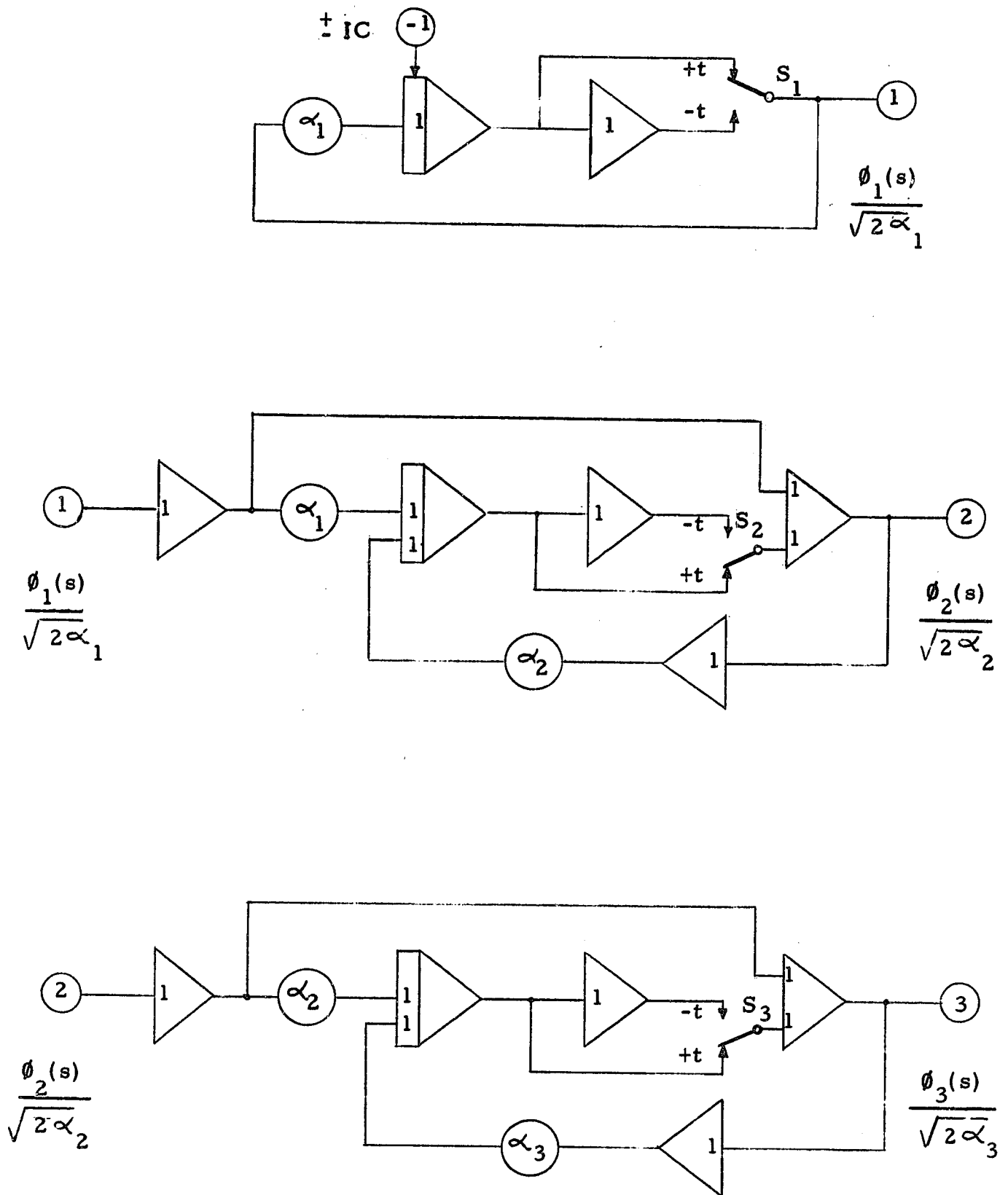
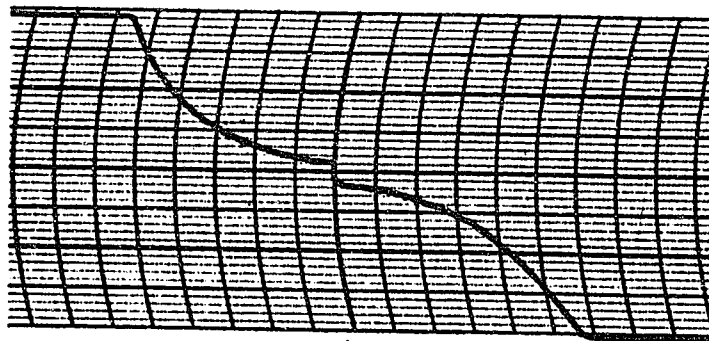
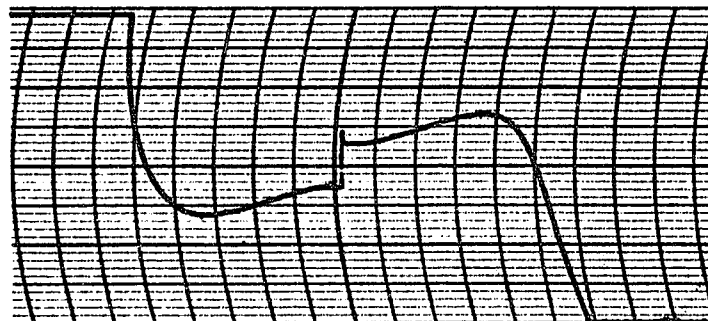


Fig. 3.8. - Computer set-up for generation of time reversed signals



a) Time reversal of $\phi_1(t)$



b) Time reversal of $\phi_2(t)$

Fig. 3.9. Time reversed signals generated on analog computer.

3.3 System Representation in Signal Space

The problem collateral to the representation of the signal on some suitable orthonormal basis, is to express the effect of the transmission system or filter upon the signal. This problem is readily solved by considering the behaviour of the linear system analogous to the linear operator in N-dimensional vector space. Thus the system H can be treated as a linear operator, which transforms the input signal $|r\rangle$ into a new output vector $|c\rangle$ in the same signal space, whose magnitude and direction will be in general different from the input signal vector. Symbolically this can be written as

$$|c\rangle = |H|r\rangle \quad (3.18)$$

If $|r\rangle$ and $|c\rangle$ both are represented on some suitable basis $\{|\phi_j\rangle\}$

then
$$|c\rangle = \sum_{j=1}^N c_j |\phi_j\rangle \quad (3.19)$$

and
$$|r\rangle = \sum_{k=1}^N r_k |\phi_k\rangle \quad (3.20)$$

substituting (3.20) in equation (3.18)

$$\begin{aligned} |c\rangle &= |H| \sum_{k=1}^N r_k |\phi_k\rangle \\ &= \sum_{k=1}^N r_k |H|\phi_k\rangle \end{aligned} \quad (3.21)$$

Thus, basically if the effect of the operator on the base vectors $|\phi_k\rangle$ is known, then the output vector $|c\rangle$ can be computed. If the basis is sufficiently complete, then each of the modified components may itself be expanded in terms of the same basis, viz,

$$|H|\phi_k\rangle = \sum_{j=1}^N h_{jk} |\phi_j\rangle \quad (3.22)$$

so that the equation (3.21) becomes

$$|c\rangle = \sum_{j=1}^N r_k \sum_{k=1}^N h_{jk} |\phi_j\rangle \quad (3.23)$$

which can be rearranged as

$$|c\rangle = \sum_{j=1}^N \sum_{k=1}^N h_{jk} \cdot r_k |\phi_j\rangle \quad (3.24)$$

equating this with (3.19)

$$\sum_{j=1}^N c_j |\phi_j\rangle = \sum_{j=1}^N \sum_{k=1}^N h_{jk} \cdot r_k |\phi_j\rangle \quad (3.25)$$

this gives

$$c_j = \sum_{k=1}^N h_{jk} \cdot r_k \quad j=1, 2, \dots, N \quad (3.26)$$

Now, it is clear from equation (3.26) that the coordinates $\{c_j\}$ of the output vector are related to the coordinates $\{r_k\}$ of the input signal vector by a matrix equation

$$\underline{c} = H \cdot \underline{r} \quad (3.27)$$

The matrix H consists of a collection of numbers which are characteristic of the system and are independent of the signals $\underline{c}$ and $\underline{r}$. This matrix, therefore, constitutes a formal numerical representation of the system.

The matrix representation is amenable to all of the rules for interconnecting system components that are so widely used with Laplace transform representation. Laplace transform is to be simply replaced by its matrix equivalent. Fig. 3.10 illustrates some familiar system interconnections to obtain the overall characterization when the characterizations of sub systems are known.

3.4 Identification of Matrix Elements of System Matrix

In the last section it was shown that the system operator has a matrix form. If the proper instrumentation can be suggested to evaluate these elements, then the system can be identified in signal space. Since these elements represent the transmission characteristics of a system, they are independent of the input and output signals. So, for simplicity we can as well consider the operation of the system on a base vector. Equation (3.22) gives

$$|H|\phi_k\rangle = \sum_{j=1}^N h_{jk} |\phi_j\rangle$$

operating on both sides by a pattern $\langle \tilde{\phi}_m |$ to get the inner product

$$\begin{aligned} \langle \tilde{\phi}_m | H | \phi_k \rangle &= \langle \tilde{\phi}_m | \sum_{j=1}^N h_{jk} |\phi_j\rangle \\ &= \sum_{j=1}^N h_{jk} \langle \tilde{\phi}_m | \phi_j \rangle \\ &= h_{mk} \end{aligned}$$

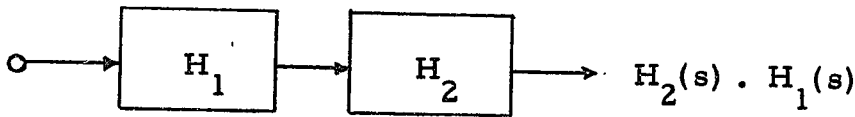
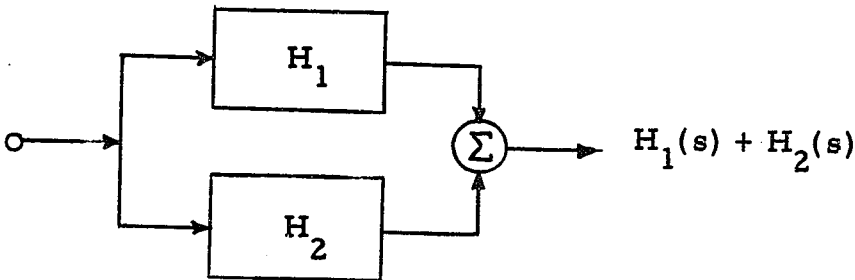
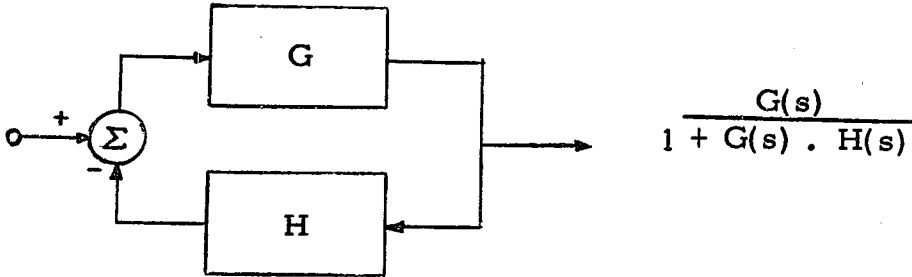
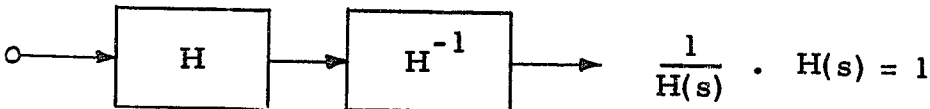
System Configuration	Laplace Transform	Matrix Representation
(a) Cascade Connection		
	$H_2(s) \cdot H_1(s)$	$H_2 \cdot H_1$
(b) Parallel Connection		
	$H_1(s) + H_2(s)$	$H_1 + H_2$
(c) Feedback Connection		
	$\frac{G(s)}{1 + G(s) \cdot H(s)}$	$[I + G \cdot H]^{-1} \cdot G$
(d) Equalizer		
	$\frac{1}{H(s)} \cdot H(s) = 1$	$H^{-1} \cdot H = 1$

Fig. 3.10. - Formal resemblance between the familiar network equations using Laplace transforms and their matrix equivalents.

This gives the matrix element in m^{th} row and k^{th} column of the system matrix H . Thus the system matrix is given by
In time domain

$$|H|\phi_k\rangle = \int_{-\infty}^t h(t-\tau) \phi_k(\tau) d\tau \quad (3.29)$$

Since $\langle \tilde{\phi}_m | H | \phi_k \rangle$ is the inner product of two time functions represented by $|\phi_m\rangle$ and $|H|\phi_k\rangle$.

Thus

$$\langle \tilde{\phi}_m | H | \phi_k \rangle = \int_0^\infty \phi_m(t) \int_{-\infty}^t h(t-\tau) \phi_k(\tau) d\tau \cdot dt \quad (3.30)$$

This expression is not particularly suitable for instrumentation.

In frequency domain

$$|H|\phi_k(s)\rangle = H(s) \cdot \phi_k(s) \quad (3.31)$$

and

$$\langle \tilde{\phi}_m | H | \phi_k \rangle = \int_{-j\infty}^{j\infty} \phi_m(-s) \cdot H(s) \phi_k(s) \frac{ds}{2\pi j} \quad (3.32)$$

This shows that $\langle \tilde{\phi}_m | H | \phi_k \rangle$ will be the value of the output, sampled at $t = 0$, of system with system function $|H|\phi_k\rangle$, the input being the time reversed signal $\phi_m(-t)$. System having system function represented by $|H|\phi_k\rangle$ is simply the cascade connection of the system $H(s)$ and the filter $\phi_k(s)$ as seen from equation (3.32). It has already been shown that the time reversed signals $\{\phi_m(-t)\}$ can be generated on analog computer. The complete set up for evaluating the system coefficients a_m 's and the system matrix elements h_{mk} 's is shown in Fig. 3.11.

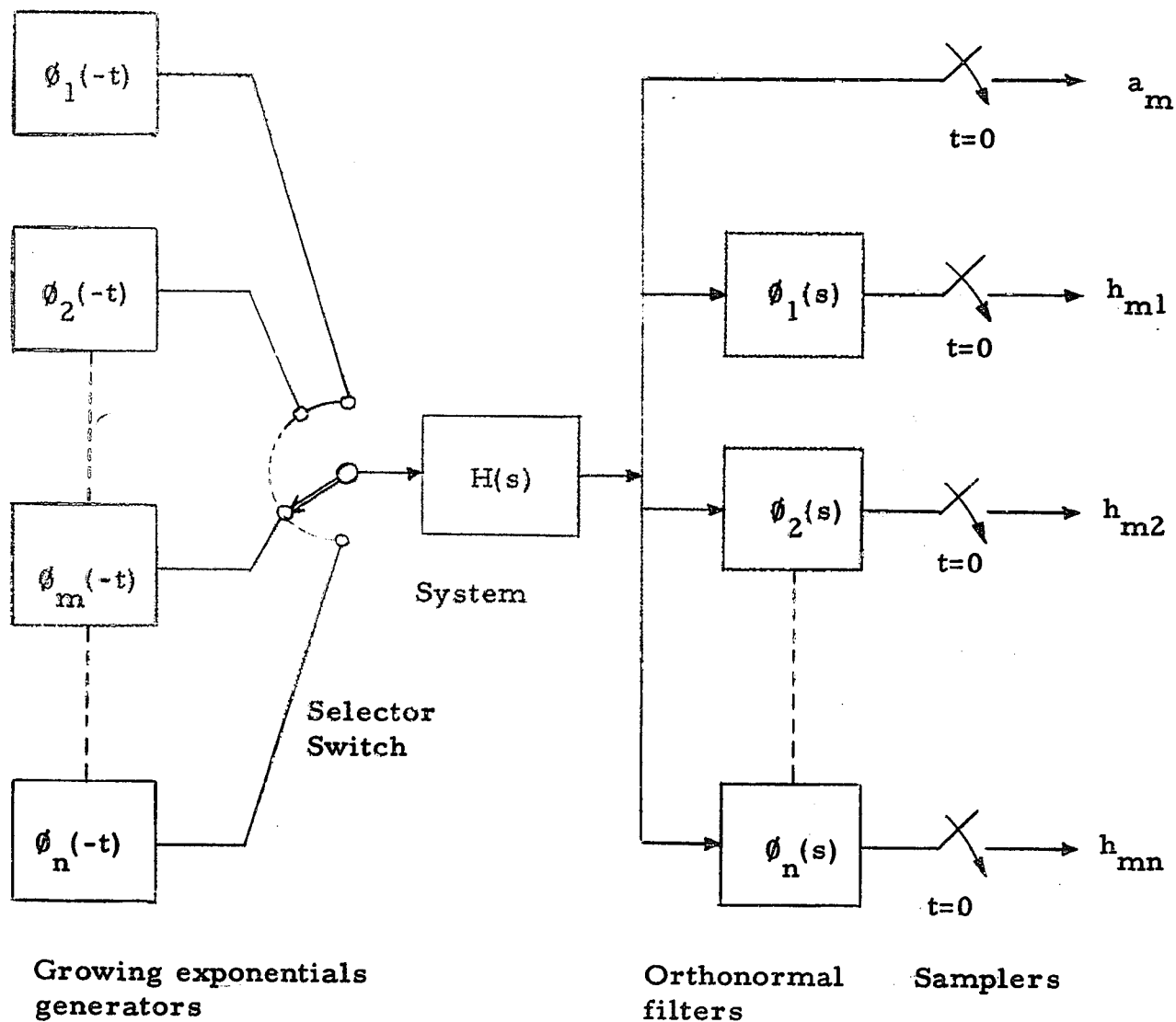


Fig. 3.11. - General scheme showing the instrumentation for evaluation of system coefficients a_m 's and system matrix elements h_{mk} 's.

3.5 Certain Properties of System Matrix

Referring to equation (3.28) the matrix element h_{mk} of the system matrix is given by

$$\begin{aligned} h_{mk} &= \langle \tilde{\phi}_m | H | \phi_k \rangle \\ &= \int_{-j\infty}^{j\infty} \phi_m(-s) \cdot H(s) \cdot \phi_k(s) \frac{ds}{2\pi j} \end{aligned} \quad (3.30)$$

If the system is realizable and has finite energy, then all of its poles are located in the left half of complex frequency s -plane. Similarly all $\{\phi_k(s)\}$ have their poles located in the left half of s -plane. Since $\phi_m(-s)$ is the mirror image of $\phi_m(s)$ across the $j\omega$ -axis, it will have all its poles located in the right half of s -plane. For $k > m$ the zeros of $\phi_k(s)$ are so located, that they cancel all the poles of $\phi_m(-s)$ in the right half plane. Hence the function under the integral sign has all its poles located only in the left half plane, so that the function is analytic in the right half plane. If the function is analytic in any half - either right or left - of s -plane, then the integral vanishes^{19, 23}.

This gives

$$h_{mk} = 0 \quad \text{for all } k > m \quad (3.31)$$

This means that all the elements above the diagonal are zero and the system matrix is a triangular matrix. This offers several remarkable advantages. Since all the elements above the diagonal are zero, the numerical operations such as multiplication and inversion are greatly simplified. Furthermore one can tell

at a glance if the inverse matrix H^{-1} exists by simply noting that none of the diagonal elements are zero. Any element above the diagonal being non zero indicates the instability of the system, and that there exists at least one pole in the right half of s-plane.

Since the product and sum of two or more triangular matrices is again a triangular matrix, it follows that the cascade and parallel combinations of stable, physically realizable systems, are stable and physically realizable also.

3.6 Modifications for Linear Time Varying Systems

The differential equation, given in equation (1.12) representing the operation of the linear time invariant system, can be modified for the case of linear time varying systems and takes the form²⁵

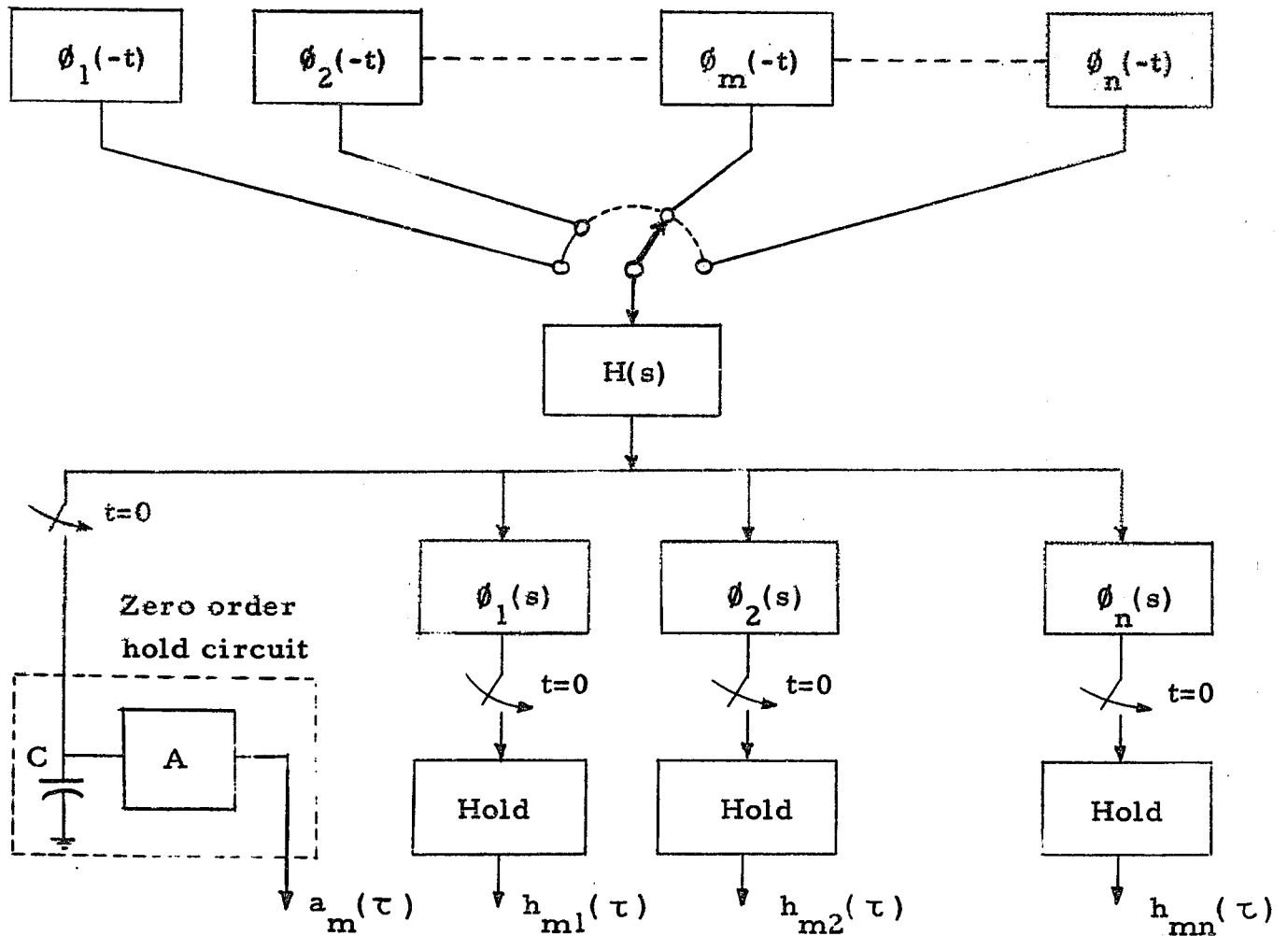
$$\sum_{n=0}^N a_n(t) \frac{d^n}{dt^n} c(t) = \sum_{m=0}^M b_m(t) \frac{d^m}{dt^m} r(t), \quad N \geq M \quad (3.32)$$

where $r(t)$ and $c(t)$ are input and output signals. Some of the coefficients $a_n(t)$ and $b_m(t)$ may be identically zero, except $a_N(t)$. If the input

$$r(t) = \delta(t - \tau) \quad (3.33)$$

where τ is the instant of time at which the unit impulse is applied to the system. The response of the system which is a system function can be written as $h(t, \tau)$. If the variations in the system parameter are slow enough to assume that they are approximately constant during the period required for one set of measurements, then

a suitable modification of the procedure given before is possible. The system coefficients a_m 's and the system matrix elements h_{mk} 's will be the functions of τ . Hence if the zero order hold circuits are connected after the samplers, and if the whole procedure for evaluating the coefficients is put into the repetitive operation, then the outputs of the hold circuits will be the staircase approximation to $a_m(\tau)$ and $h_{mk}(\tau)$. The block schematic is shown in Fig. 3.12.



A - Unity gain buffer amplifier

C - Hold capacitor

Fig. 3.12. - Modified scheme for linear time varying systems.

CHAPTER IV

4.1 Computational Error

A simple calculation shows how the errors in computation of the system coefficients a_n affect the integral squared approximation error.

Let $a_n + \Delta a_n$ be the computed coefficients where Δa_n is the computational error in computing the coefficients a_n . If $\epsilon = \epsilon_{\min} + \Delta \epsilon$ is the integral squared error where $\Delta \epsilon$ is the increase in the minimized error $\epsilon_{\min}$, because of the coefficients errors Δa_n 's, then

$$\begin{aligned} \epsilon &= \epsilon_{\min} + \Delta \epsilon \\ &= \int_0^{\infty} \left[h(t) - \sum_{n=1}^N (a_n + \Delta a_n) \phi_n(t) \right]^2 dt \end{aligned} \quad (4.1)$$

For simplicity of writing here after we shall write only h , ϕ_n etc. instead of $h(t)$, $\phi_n(t)$. Substituting these in equation 4.1 we get

$$\begin{aligned} \epsilon &= \int_0^{\infty} \left[h - \sum_{n=1}^N (a_n + \Delta a_n) \phi_n \right]^2 dt \\ &= \int_0^{\infty} \left[(h - h_a) - \sum_{n=1}^N \Delta a_n \phi_n \right]^2 dt \end{aligned} \quad (4.2)$$

Using the Dirac's notation let the instantaneous error $e(t)$ be denoted by a vector $|e\rangle$ in signal space, so that

$$|e\rangle = |h\rangle - |h_a\rangle \quad (4.3)$$

Then

$$\begin{aligned}
 \epsilon &= (\langle \tilde{e} | - \sum_{n=1}^N \Delta a_n \langle \tilde{\phi}_n |)(|e\rangle - \sum_{n=1}^N \Delta a_n |\phi_n\rangle) \\
 &= \langle \tilde{e} | e \rangle - \langle e | \sum_{n=1}^N \Delta a_n |\phi_n\rangle - \sum_{n=1}^N \Delta a_n \langle \phi_n | e \rangle \\
 &\quad + \sum_{n=1}^N \sum_{m=1}^N \Delta a_n \Delta a_m \langle \tilde{\phi}_n | \phi_m \rangle \\
 &= \epsilon_{\min} - 2 \sum_{n=1}^N \Delta a_n \langle \tilde{\phi}_n | e \rangle + \sum_{n=1}^N \sum_{m=1}^M \Delta a_n \Delta a_m \delta_{mn} \\
 &= \epsilon_{\min} + \sum_{n=1}^N \Delta a_n^2 \tag{4.4}
 \end{aligned}$$

The second term vanishes because the error $|e\rangle$ is orthogonal to all the base vectors $|\phi_n\rangle$ as shown in equation (2.23).

Thus we get

$$\Delta \epsilon = \sum_{n=1}^N \Delta a_n^2 \tag{4.5}$$

When $\Delta \epsilon$ is small compared to $\epsilon_{\min}$ computational errors can be considered negligible. The relative integral squared error is defined by

$$\begin{aligned}
 e_{\min} &= \frac{\epsilon_{\min}}{\int_0^{\infty} [h(t)]^2 dt} \\
 &= \frac{\epsilon_{\min}}{\langle \tilde{h} | h \rangle}
 \end{aligned} \tag{4.6}$$

Hence the relative error due to computational error is given by

$$\Delta e = \frac{\sum_{n=1}^N \Delta a_n^2}{\langle \tilde{h} | h \rangle} \tag{4.7}$$

4.2 Examples

Let us illustrate the use of the theory developed through the previous pages by some simple examples. In all the following examples we shall use the same orthonormal basis formed from the poles on negative real axis of s-plane at $s = -1, -2, -4$. Thus the basis are

$$\begin{aligned}
 \phi_1(s) &= \frac{\sqrt{2}}{s+1} \\
 \phi_2(s) &= \frac{\sqrt{4}}{s+1} \cdot \frac{s-1}{s+2} \\
 \phi_3(s) &= \frac{\sqrt{8}}{s+1} \cdot \frac{s-1}{s+2} \cdot \frac{s-2}{s+4}
 \end{aligned} \tag{4.8}$$

The time domain forms of these three basis are shown in Fig. 4.1.

Example 1

Consider a first order linear time invariant system having a transfer function

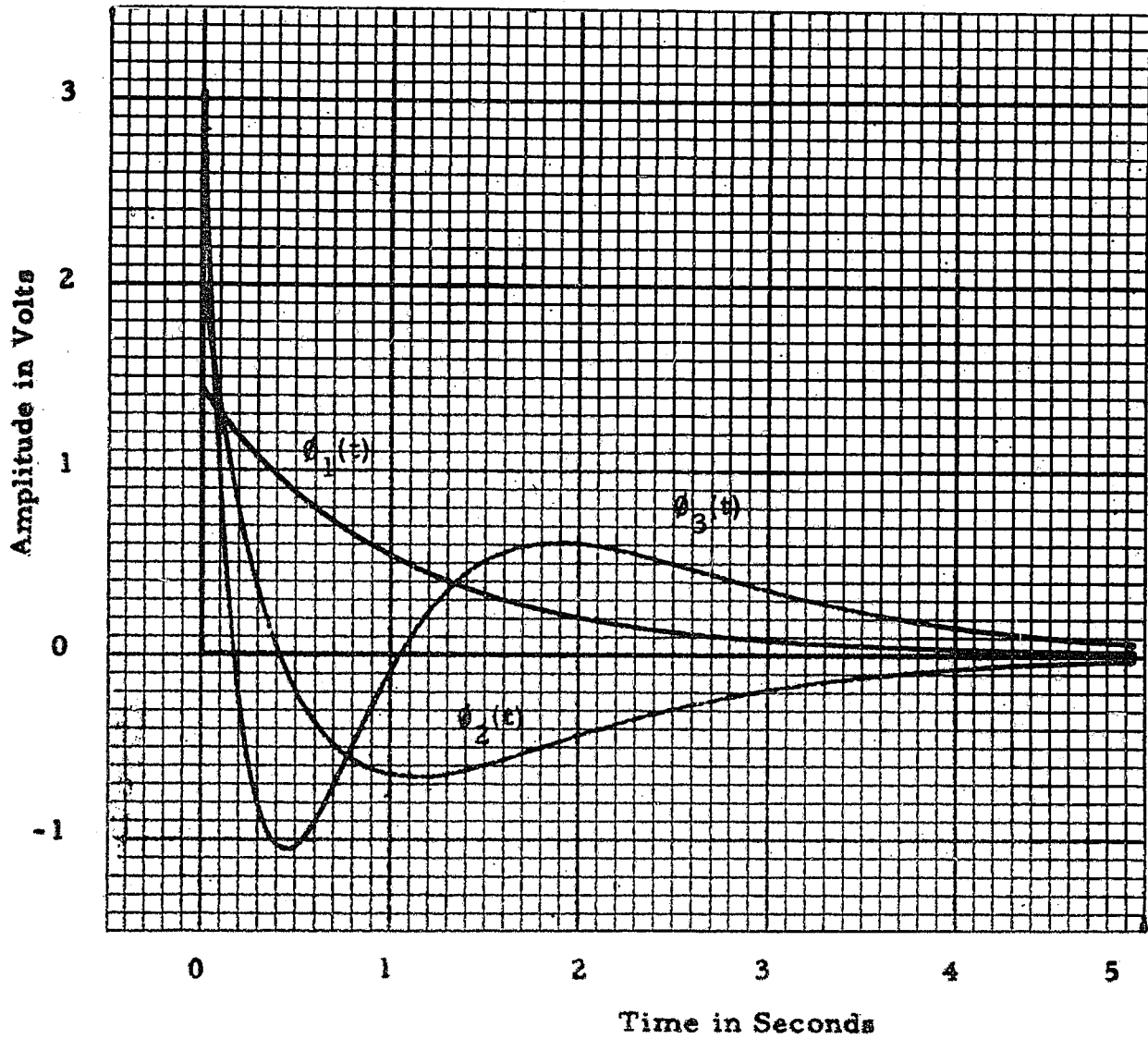


Fig. 4. 1. Time domain form of orthonormal basis, defined in equation (4. 8).

$$H(s) = \frac{1}{s+3} \quad (4.9)$$

Energy in the system function is given by

$$\begin{aligned} \langle \tilde{H} | H \rangle &= \int_{-j\infty}^{j\infty} H(-s) H(s) \frac{ds}{2\pi j} \\ &= \int_{-j\infty}^{j\infty} \frac{1}{-s+3} \cdot \frac{1}{s+3} \cdot \frac{ds}{2\pi j} \end{aligned} \quad (4.10)$$

By residue theorem

$$\langle \tilde{H} | H \rangle = \frac{1}{6} \quad (4.11)$$

The system coefficients a_1, a_2, a_3 which are the projections of the system function $|H\rangle$ onto the basis $|\phi_1\rangle, |\phi_2\rangle, |\phi_3\rangle$ are computed analytically as follows

$$\begin{aligned} a_1 &= \langle \tilde{\phi}_1 | H \rangle \\ &= \int_{-j\infty}^{j\infty} \frac{\sqrt{2}}{-s+1} \cdot \frac{1}{s+3} \cdot \frac{ds}{2\pi j} \\ &= \frac{\sqrt{2}}{4} \\ &= 0.3535 \end{aligned} \quad (4.12)$$

The energy in the first component is a_1^2 . Let us define the normalized energy in individual components as

$$E_n = \frac{a_n^2}{\langle \tilde{H} | H \rangle} \quad n = 1, 2, \dots, N \quad (4.13)$$

Hence

$$\begin{aligned} E_1 &= \frac{\sqrt{2}}{4} \cdot \frac{\sqrt{2}}{4} \cdot \frac{6}{1} \\ &= 0.75 \end{aligned} \quad (4.14)$$

Thus the first component contains 75% of the total energy in the system function and hence one term approximation will have a relative integral squared error equal to 0.25 (i. e. 25%).

By similar procedure we can compute

$$a_2 = \frac{1}{5} = 0.2 \quad (4.15)$$

$$\text{and } a_3 = \frac{\sqrt{8}}{70} = 0.0405 \quad (4.16)$$

while the normalized energy contents in these components are given by

$$E_2 = 0.24 \quad (\text{i. e. } 24\%) \quad (4.17)$$

$$\text{and } E_3 = 0.0098 \quad (\text{i. e. } 0.98\%) \quad (4.18)$$

First two components a_1 and a_2 were computed by using analog computer. The system response to the time reversed input signals $\phi_1(-t)$ and $\phi_2(-t)$ is shown in Fig. 4.2. The values of a_1 and a_2 as obtained from the computer are tabulated below

Coefficients	Analytic Computation	From Computer	Error a_n	Energy
a_1	0.3535	0.3521	-0.0014	0.75
a_2	0.2000	0.2020	0.0020	0.24

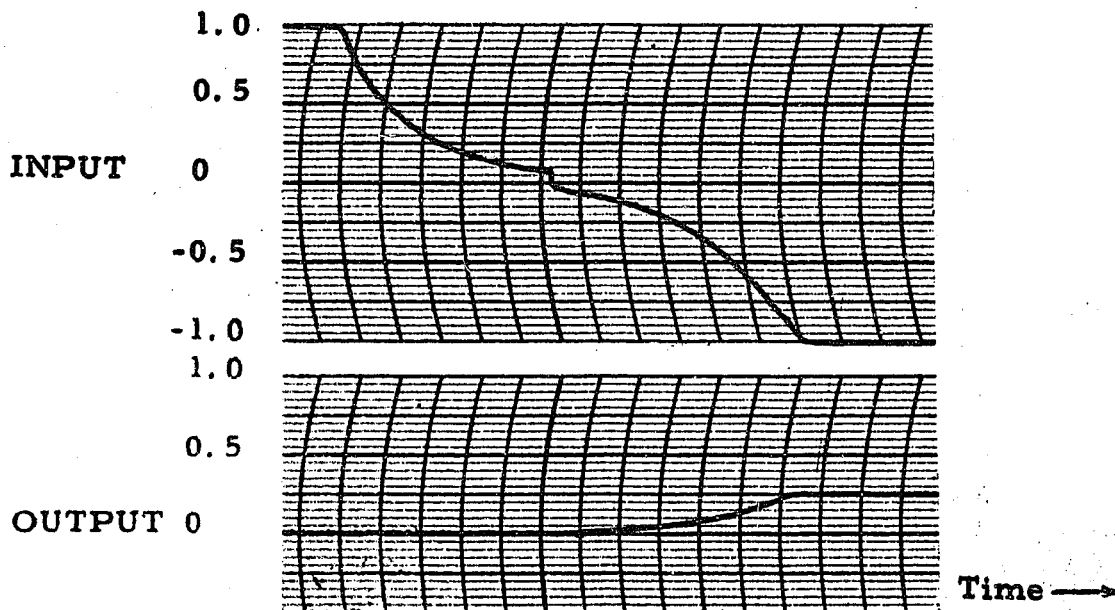


Fig. 4.2A. System response to the time reversed signal $\phi_1(-t)$

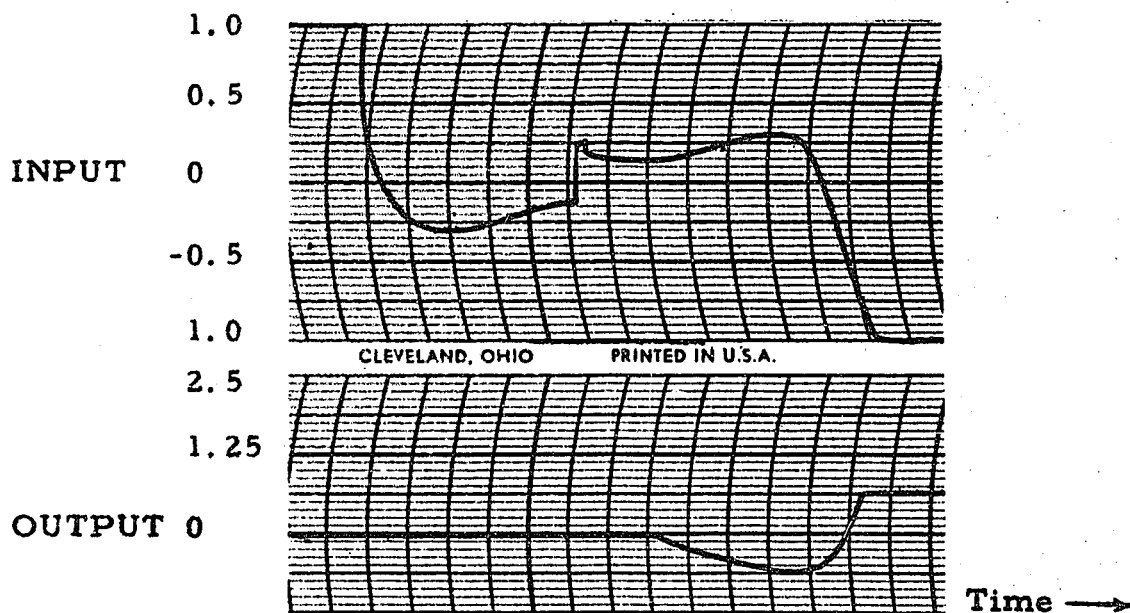


Fig. 4.2B. System response to the time reversed signal $\phi_2(-t)$

The errors Δa_n are due to the analog computation of the coefficients. The relative integral squared error for two term approximation is 0.01 (i.e. 1%) whereas the relative integral squared error due to the errors Δa_n as computed by equation (4.7) is

$$\begin{aligned}
 &= \frac{\Delta a_1^2 + \Delta a_2^2}{\langle \widetilde{H} | H \rangle} \\
 &= 6 \left[(0.0014)^2 + (0.002)^2 \right] \\
 &= 35.76 \times 10^{-6} \quad (\approx 0.0036\%)
 \end{aligned} \tag{4.19}$$

Hence the analog computer errors are completely negligible.

The approximate transfer functions using one, two and three term approximations are computed below:

1) One term approximation

$$\begin{aligned}
 H_a(s) &= a_1 \phi_1(s) \\
 &= \frac{0.5}{s+1}
 \end{aligned} \tag{4.20}$$

2) Two term approximation

$$\begin{aligned}
 H_a(s) &= a_1 \phi_1(s) + a_2 \phi_2(s) \\
 &= \frac{0.9(s + 0.66)}{(s + 1)(s + 2)}
 \end{aligned} \tag{4.21}$$

3) Three term approximation

$$H_a(s) = \frac{71}{70} \cdot \frac{(s + 2.93)(s + 0.87)}{(s + 1)(s + 2)(s + 4)} \tag{4.22}$$

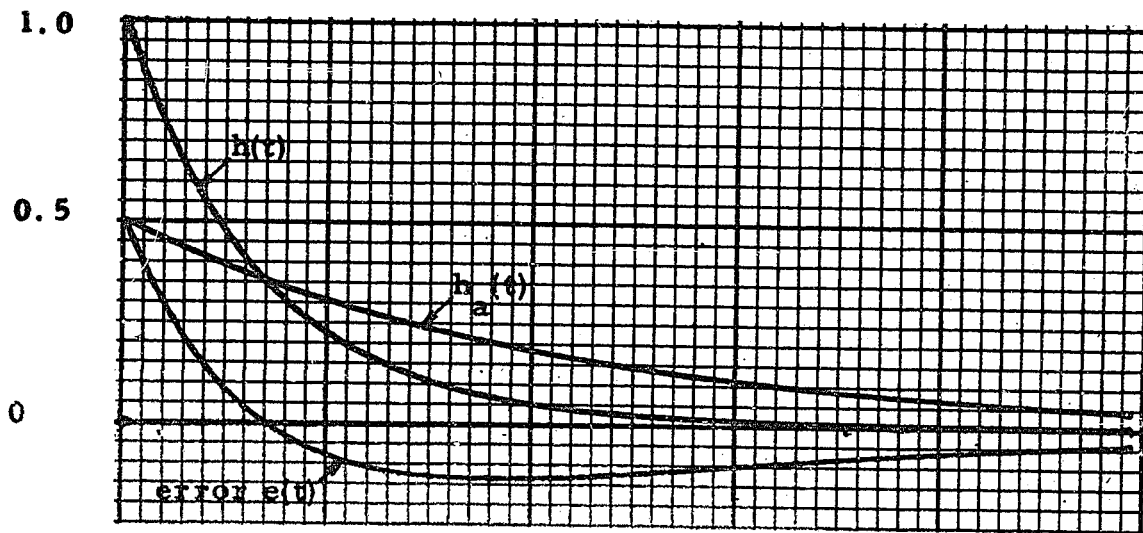


Fig. 4.3A. One term approximation.

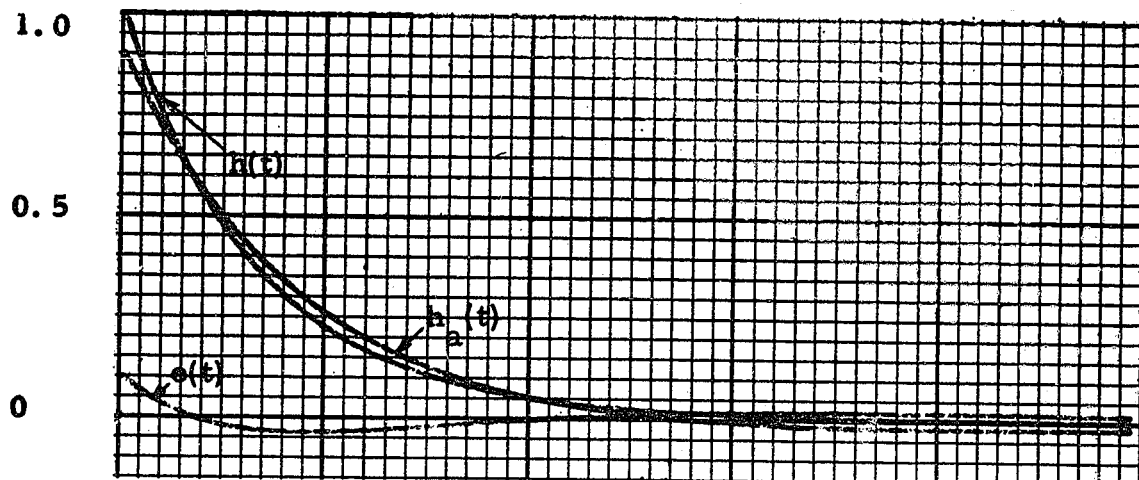


Fig. 4.3B. Two term approximation.

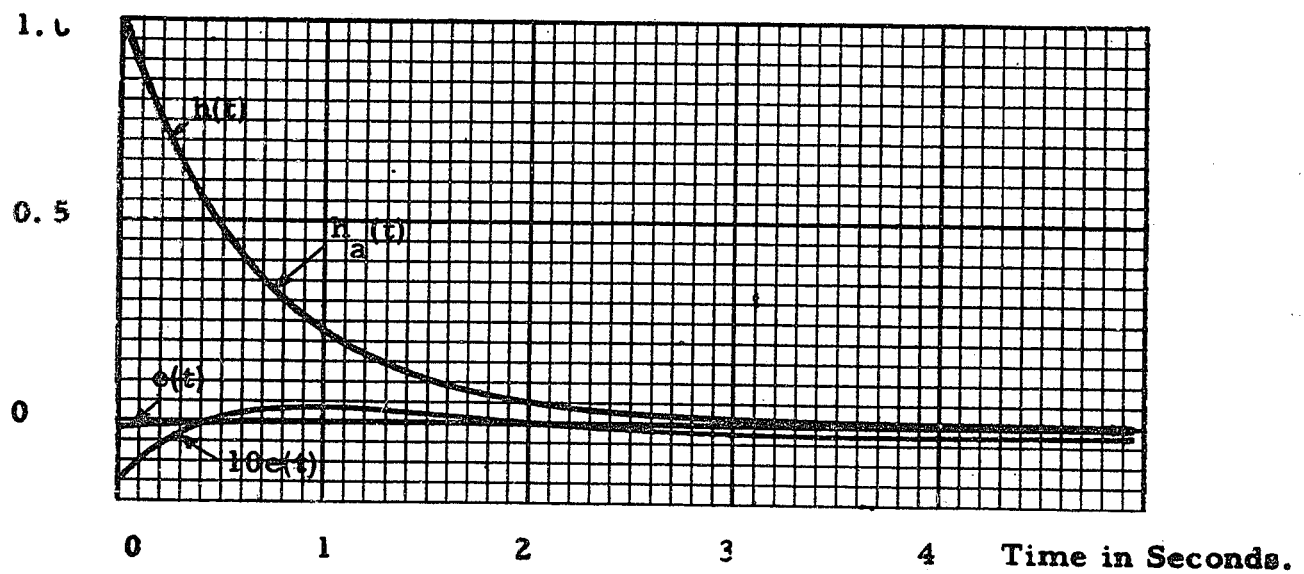


Fig. 4.5C. Three term approximation.

The impulse responses of these three approximate systems are plotted in Fig. 4.3, together with the impulse response of actual system for comparison. As is seen from the graph, the two term approximation is fairly good with 1% integral squared error. Whereas the three term approximation shows almost exact reproduction of the system function with a negligible error of 0.02%.

Example 2

Consider a second order system with transfer function

$$H(s) = \frac{1}{(s+3)(s+5)} \quad (4.23)$$

The energy in the system function

$$\begin{aligned} \langle \widetilde{H} | H \rangle &= \int_{-j\infty}^{j\infty} H(-s) \cdot H(s) \frac{ds}{2\pi j} \\ &= \int_{-j\infty}^{j\infty} \frac{1}{(-s+3)(-s+5)} \cdot \frac{1}{(s+3)(s+5)} \frac{ds}{2\pi j} \\ &= \frac{1}{240} \end{aligned} \quad (4.24)$$

The components of the system function onto the basis $|\phi_1\rangle$, $|\phi_2\rangle$, $|\phi_3\rangle$ defined in equation (4.8) are worked out analytically as well as on computer and they are tabulated in the table shown below. The relative energy contents in different components are also worked out.

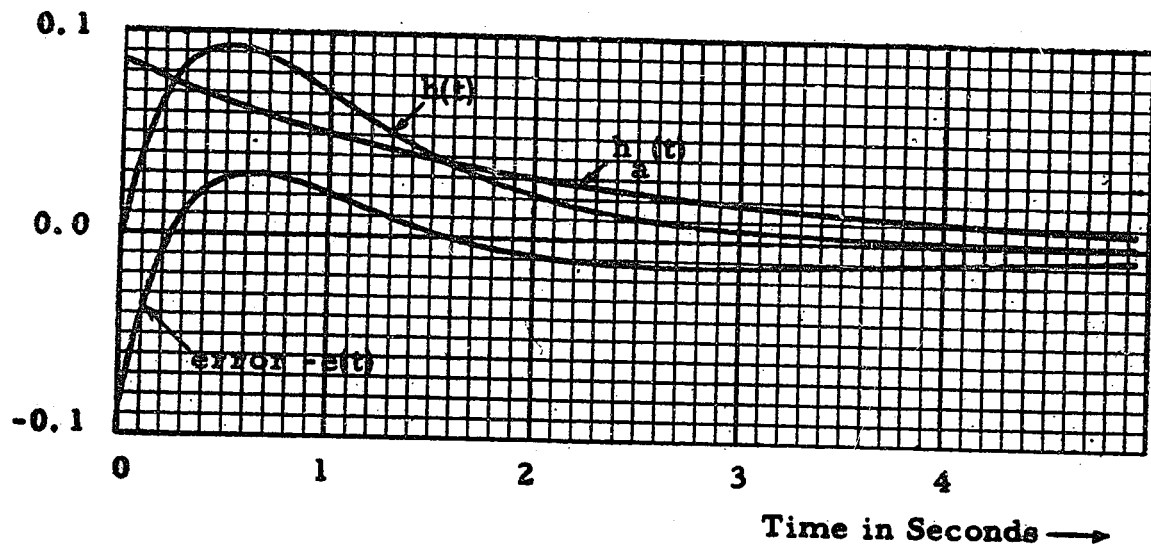


Fig. 4.4A. One term approximation to system function $h(t)$.

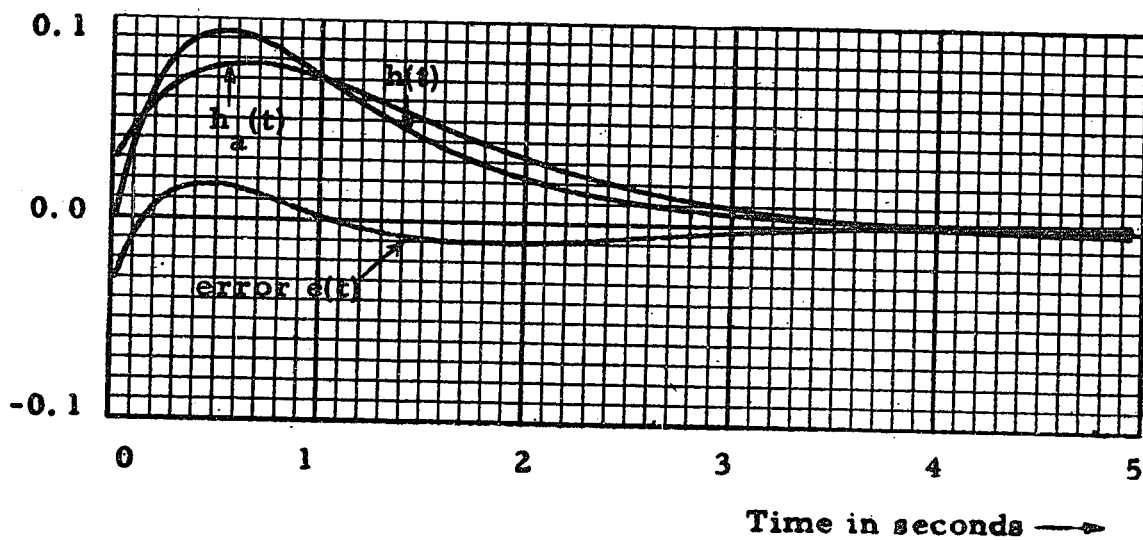


Fig. 4.4B. Two term approximation of system function $h(t)$, using the base vectors $|\phi_1\rangle$ and $|\phi_3\rangle$.

Coefficients	Analytic Computation	From Computer	Error a_n	Energy
a_1	0.059	0.060	0.001	0.8333
a_2	-0.0047	-0.003	0.0017	0.0054
a_3	-0.0247	-	-	0.1462

Considering the first two components the relative integral squared error is 0.1613 whereas the error due to analog computer is 0.093×10^{-2} which is again negligible.

As seen from the table most of energy in system function is directed towards the basis $|\phi_1\rangle$ and $|\phi_3\rangle$. They contain almost 98% of the total energy. The approximate transfer function using only these two components is given by

$$H_a(s) = 0.43 \frac{(s + 1.52)(s + 5.93)}{(s + 1)(s + 2)(s + 4)} \quad (4.25)$$

The impulse response of the original system and the approximate systems with one and two term approximation is shown in Fig. 4.4. It is seen that the two term approximation is fairly good with a small 2% error.

Example 3

Consider a third order system with transfer function

$$H(s) = \frac{1}{(s + 1)(s + 3)(s + 5)} \quad (4.26)$$

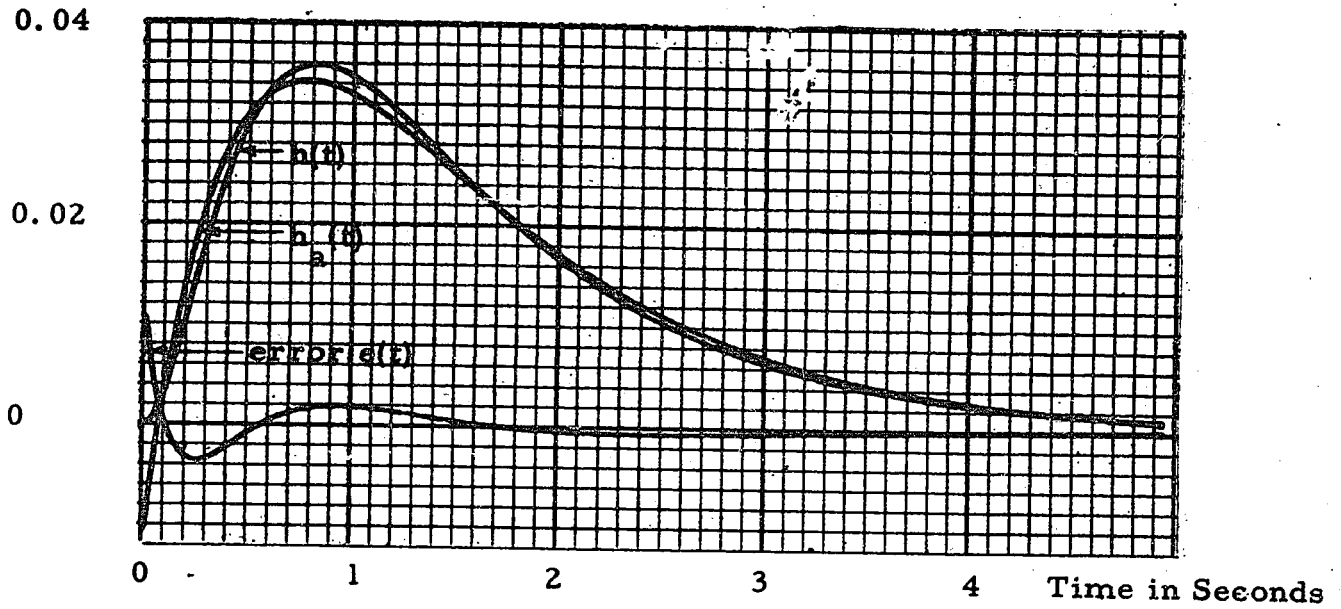


Fig. 4.5. Comparison of impulse response of actual system and approximate system.

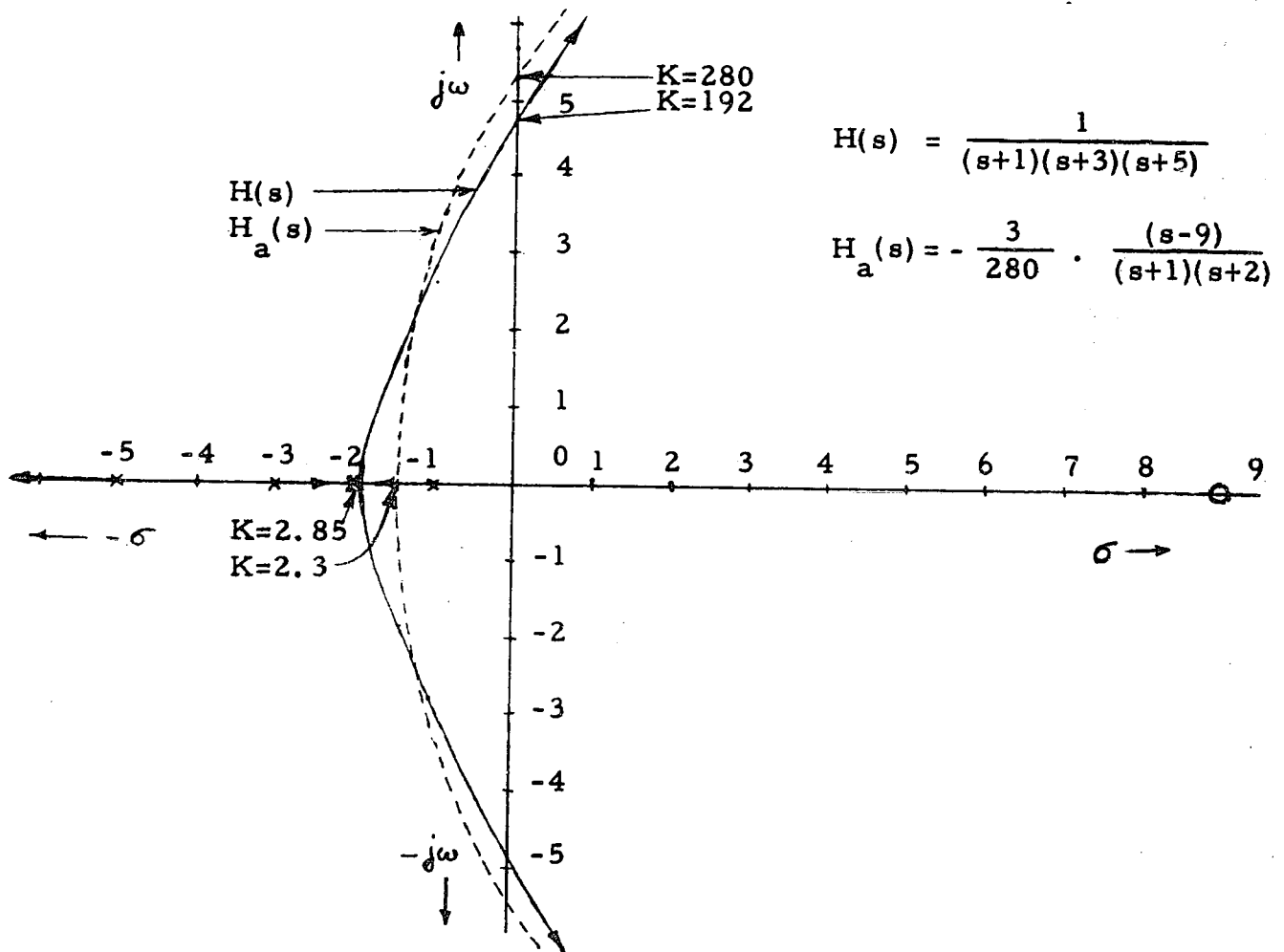


Fig. 4.6. Comparison of root loci of approximate system and actual system.

The energy in the system function is $\frac{1}{640}$. The system coefficients a_1, a_2, a_3 are tabulated below, with their respective relative energy contents.

Coefficients	Analytic Computation	Energy
a_1	$\sqrt{2} / 48$	0.5555
a_2	$-11 / 420$	0.4373
a_3	$\sqrt{8} / 2520$	0.0008

The first two components contain 99.28% of the total energy. The comparison of the impulse response of actual system and the approximate system is shown in Fig. 4.5. Transfer function of the approximate system using first two components is given by

$$H_a(s) = -\frac{3}{280} \cdot \frac{s - 9}{(s + 1)(s + 2)} \quad (4.27)$$

Since original system becomes unstable in feedback configuration for certain value of gain it is of interest to compare the root loci of these two systems. In Fig. 4.6 the root loci are plotted for comparison. It is seen that in the left half of s-plane they match with each other very closely.

Example 4

Let us now consider the system having two complex conjugate poles at $s = -1 \pm j$. The system transfer function is given by

$$H(s) = \frac{1}{s^2 + 2s + 2} \quad (4.28)$$

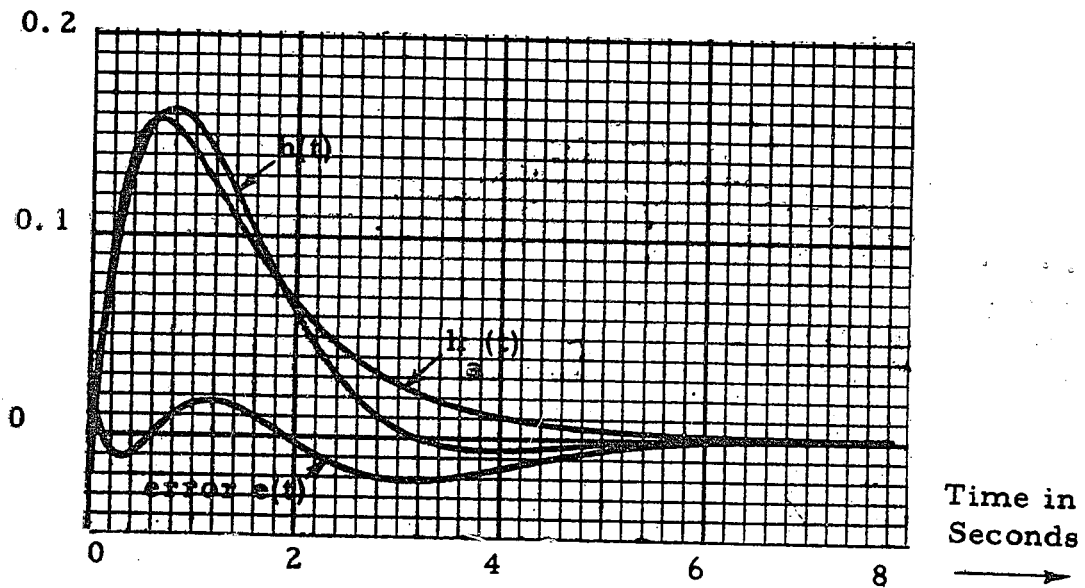


Fig. 4.7. Comparison of impulse response of actual system and approximate system.

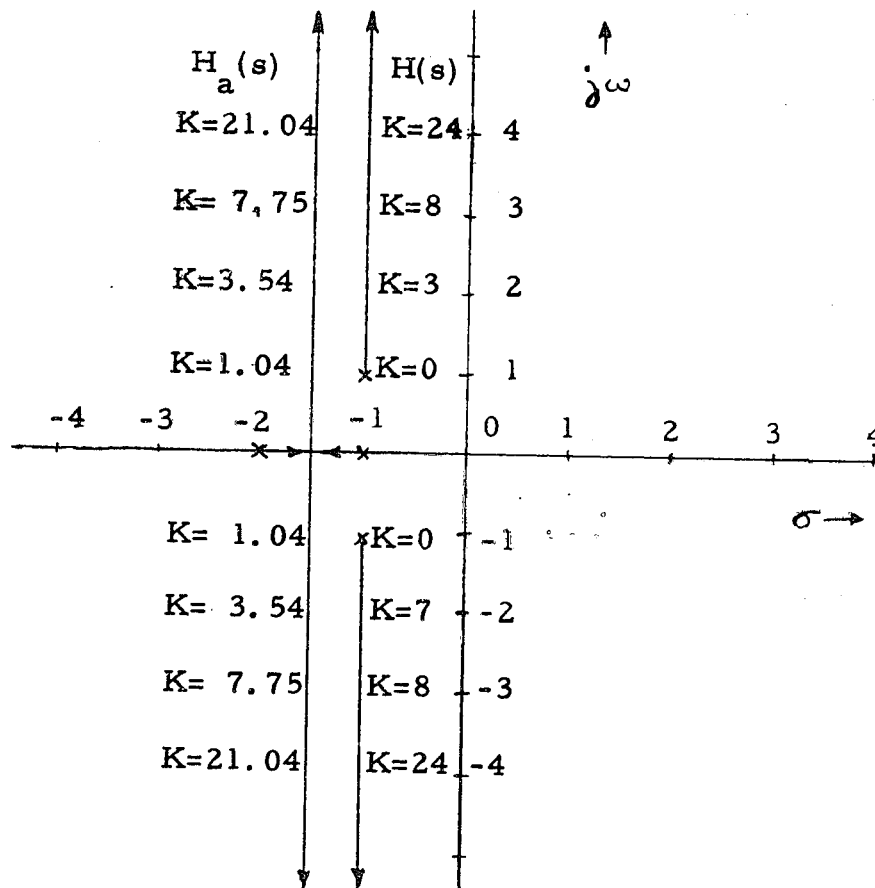


Fig. 4.8. Comparison of root loci of approximate system and actual system.

Energy in the system function is $\frac{1}{8}$. System coefficients and the relative energy in the respective components is analytically computed in the similar manner. Coefficients computed from computer readings are tabulated below.

Coefficients	Analytic Computation	From Computer	Δa_n	Energy
a_1	0.2828	0.2842	0.0014	0.64
a_2	-0.200	-0.204	-0.004	0.32
a_3	-0.021	-	-	0.0038

The first two components contain as much as 96% of the total energy. The approximate transfer function using these two components is

$$H_a(s) = \frac{6}{5} \cdot \frac{1}{(s+1)(s+2)} \quad (4.29)$$

Comparison of the impulse responses of the approximate system with the actual system is shown in Fig. 4.7. Their root loci are compared in Fig. 4.8.

Example 5

In this example we will illustrate the use of matrix representation of the linear system in signal space. Let us consider a simple system same as in example 1.

$$H(s) = \frac{1}{s+3} \quad (4.30)$$

We shall consider only first two basis. The elements of the system matrix are computed as follows using equation (3.28):

$$\begin{aligned}
 h_{11} &= \langle \tilde{\phi}_1 | H | \phi_1 \rangle \\
 &= \int_{-j\infty}^{j\infty} \frac{\sqrt{2}}{-s+1} \cdot \frac{1}{s+3} \cdot \frac{\sqrt{2}}{s+1} \frac{ds}{2\pi j} \\
 &= \frac{1}{4} = 0.25
 \end{aligned} \tag{4.31}$$

$$\begin{aligned}
 h_{12} &= \langle \tilde{\phi}_1 | H | \phi_2 \rangle \\
 &= \int_{-j\infty}^{j\infty} \frac{\sqrt{2}}{-s+1} \cdot \frac{1}{s+3} \cdot \frac{2(s-1)}{(s+1)(s+2)} \frac{ds}{2\pi j} \\
 &= 0
 \end{aligned} \tag{4.32}$$

$$\begin{aligned}
 h_{21} &= \langle \tilde{\phi}_2 | H | \phi_1 \rangle \\
 &= \int_{-j\infty}^{j\infty} \frac{2(-s-1)}{(-s+1)(-s+2)} \cdot \frac{1}{s+3} \cdot \frac{\sqrt{2}}{s+1} \frac{ds}{2\pi j} \\
 &= -\frac{\sqrt{2}}{10} = -0.1414
 \end{aligned} \tag{4.33}$$

$$\begin{aligned}
 h_{22} &= \langle \phi_2 | H | \phi_2 \rangle \\
 &= \int_{-j\infty}^{j\infty} \frac{2(-s-1)}{(-s+1)(-s+2)} \cdot \frac{1}{s+3} \cdot \frac{2(s-1)}{(s+1)(s+2)} \frac{ds}{2\pi j} \\
 &= \frac{1}{5} = 0.2
 \end{aligned} \tag{4.34}$$

The system matrix can be written as

		IN	
OUT		1	2
1	1	0.25	0
	2	-0.1414	0.2

The matrix elements are also evaluated on computer. The graphs shown in Fig. 4.9 are without normalizing coefficients. System matrix elements as computed from the graphs are given here for comparison.

		IN	
OUT		1	2
1	1	0.252	0.0017
	2	-0.1202	0.17

Let the input be

$$r(t) = e^{-2t} \quad (4.35)$$

so that

$$R(s) = \frac{1}{s+2} \quad (4.36)$$

The numerical representative of the input signal vector $|r\rangle$ on the basis $|\phi_1\rangle$ and $|\phi_2\rangle$ is

$$|r\rangle = \begin{bmatrix} 0.471 \\ 0.167 \end{bmatrix} \quad (4.37)$$

The components of the output vector $|c\rangle$ using analytic computation are

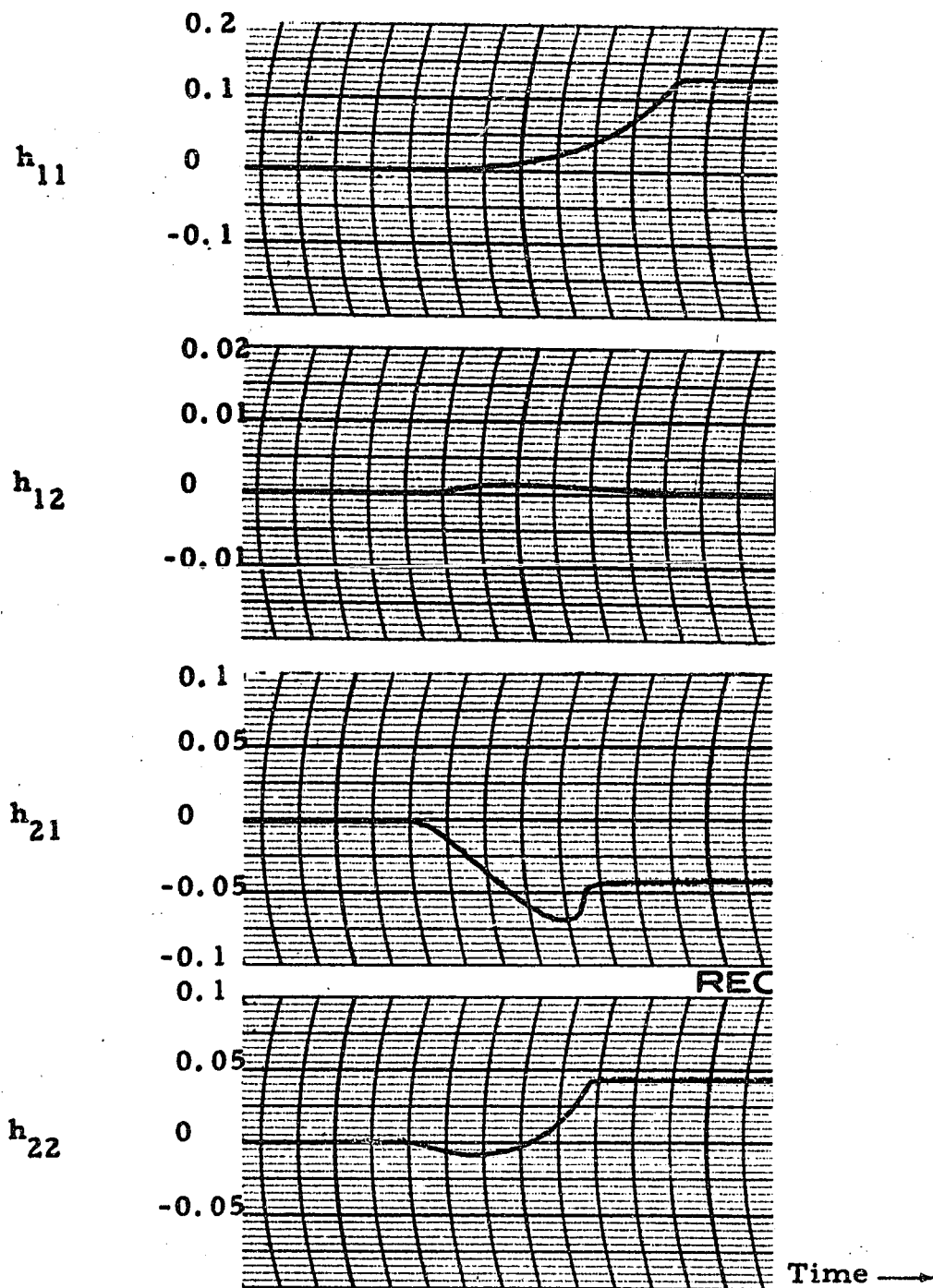


Fig. 4. 9. Evaluation of the matrix element of the system matrix, on analog computer.

$$\begin{aligned} |c\rangle &= \begin{bmatrix} 0.25 & 0 \\ -0.1414 & 0.2 \end{bmatrix} \begin{bmatrix} 0.471 \\ 0.161 \end{bmatrix} \\ &= \begin{bmatrix} 0.1178 \\ -0.0344 \end{bmatrix} \end{aligned} \tag{4.38}$$

Whereas the coefficients evaluated by using computer results are given by

$$|c\rangle = \begin{bmatrix} 0.1188 \\ -0.0392 \end{bmatrix} \tag{4.39}$$

Now the Laplace transform of the output signal is given by

$$C(s) = \frac{1}{(s+1)(s+2)} \tag{4.40}$$

The components of this on the same basis are given by

$$|c\rangle = \begin{bmatrix} 0.1178 \\ -0.0333 \end{bmatrix} \tag{4.41}$$

Comparing equations (4.38) and (4.39) to equation (4.41) it is seen that they compare very well. Thus the computer results are satisfactory.

CONCLUSION

As shown in Chapter II, time reversal of signals is essential for evaluation of system coefficients a_m 's and the system matrix elements h_{mk} 's. Analog computer technique to effect the time reversal of the orthonormal exponentials is described in Chapter III. The use of this technique is illustrated in Chapter IV by considering some simple systems.

Accurate generation of the time reversed signals on analog computer, needs a certain minimum voltage to be present on all the integrators before the time reversal switch is operated. This introduces an error Δa_n in evaluating a system coefficient a_n . Δa_n is given by

$$\Delta a_n = \int_0^{\infty} h(t) \phi_n(t) dt - \int_0^{t_1} h(t) \phi_n(t) dt$$

where t_1 is the finite time, when the time reversal of $\phi_n(t)$ is effected.

Since $h(t)$ is assumed to be a finite energy function

i. e.
$$\int_0^{\infty} [h(t)]^2 dt < \infty$$

which implies that

$$\lim_{t \rightarrow \infty} h(t) \rightarrow 0$$

Thus if t_1 is of the order of 4-5 times the largest time constant of the system then the effect of the computer error is negligible as illustrated by the examples in Chapter IV.

Thus the time reversed orthonormal exponentials are useful probing signals which can reveal the effective order of the system and also help in constructing the system model. It is obvious that these signals can be used more effectively to reveal the equivalent complex frequency structure of the system, wherever the perturbation techniques can be applied.

Although only the systems with single input and single output are considered, there should be no difficulty in applying these techniques to the multivariable systems. Making use of the tensor algebra²⁶ a possibility of using these signals for identification of non linear systems with or without memory is further to be explored.

REFERENCES

1. N. Wiener, "Nonlinear Problems in Random Theory." M.I.T. Press, Cambridge, Massachusetts, 1958.
2. A. G. Bose, "Nonlinear System Characterization and Optimization." IRE Trans., Vol. CT-6, Special Supplement, pp. 30-40, May 1959.
3. Y. W. Lee, "Statistical Theory of Communications." John Wiley & Sons, Inc., N. Y. 1960.
4. L. A. Zadeh, "On the Identification Problem." IRE Trans., Vol. CT-9, pp. 277-281, December 1956.
5. W. H. Huggins, S. Litman, "Growing Exponentials as a Probing Signal for System Identification." Proc. IEEE, pp. 917-923, June 1963.
6. P. Eykhoff, "Some Fundamental Aspects of Process-Parameter Estimation." IEEE Trans. on Automatic Control, pp. 347-57, October 1963.
7. N. Wiener, "Extrapolation, Interpolation and Smoothing of Stationary Time Series." M.I.T. Press, Cambridge, Massachusetts, 1949.
8. E. A. Guillemin, "What is Network Synthesis?" IRE Trans. Vol. CT-1, pp. 4-19, December 1952.
9. L. A. Zadeh, "A General Theory of Linear Signal Transmission Systems." Journal of Franklin Institute, Vol. 253, No. 4, April 1952.
10. L. A. Zadeh, "Generalized Ideal Filters." Journal of Appl. Physics, Vol. 23, No. 2, pp. 223-226, February 1952.

11. D. C. Lai, "Representation and Analysis of Signals: Part VI, Signal Space Concept and Dirac's Notation." Dept. of Electrical Engineering, John Hopkins University, Baltimore, AFCRC Report No. TN-60-167, January 1960.
12. Mishkin and Brown, "Adaptive Control Systmes." McGraw-Hill Book Co. 1961.
13. W. H. Huggins, "Representation and Analysis of Signals: Part I, Use of Orthogonalized Exponentials." Dept. of Electrical Engineering, John Hopkins University, Baltimore, AFCRC Report No. TR-57-357, November 1958.
14. E. G. Gilbert, "Linear System Approximation by Differential Analyser Simulation of Orthonormal Approximation Functions." IRE Trans. on Electronic Computers, pp. 204-209, June 1959.
15. G. S. Glinski and M. N. Jafri, "Modelling and Control Problems of Distillation Columns." Dept. of Electrical Engineering, University of Ottawa, Technical Report No. 64-1, January 1964.
16. R. Bellman, "Introduction to Matrix Analysis." Mc-Graw-Hill Book Co. 1960.
17. F. Ayres, Jr., "Theory and Problems of Matrices." Schaum Publishing Co. New York, 1962.
18. W. H. Kautz, "Transient Synthesis in Time Domain." IRE Trans. on Circuit Theory, Vol. CT-1, No. 3, pp. 29-39, September 1954.

19. W. R. Lepage, "Complex Variables and Laplace Transforms for Engineers." McGraw-Hill Book Co. 1961.
20. Y. W. Lee, "Synthesis of Electrical Networks by Mean of the Fourier Transforms of Laugerre's Functions." Journal of Mathematical Physics, Vol. II, pp. 83-113, 1932.
21. D. C. Ross, "Orthonormal Exponentials." IEEE Trans. on Communications and Electronics, pp. 173-176, March 1964.
22. W. W. Seifert and C. W. Steeg, Jr., "Control Systems Engineering." McGraw-Hill Book Col. 1960.
23. D. Hazony and J. Riley, "Evaluating Residues and Coefficients of High Order Poles." IRE Wescon Convention Record, Part 4, pp. 136-140, August 1959.
24. L. P. Smith, "Mathematical Methods for Scientists and Engineers." Dover Publications Inc., New York, 1961.
25. A. R. Stubberud, "Analysis and Synthesis of Linear Time Variable Systems." University of California Press, Berkeley and Los Angeles, 1964.
26. D. C. Ross, "Representation and Analysis of Signals: Part XVIII, Vector and Tensor Algebra of Signals Applied to Satellite Navigation." Dept. of Electrical Engineering, John Hopkins University, Baltimore, 1964.

VITA

Name: H. B. Kekre

Born: India, April 4th, 1935

Education:

Matriculation: New English High School, Nagpur

Inter Science: College of Science, Nagpur

B. E. Hons: Govt. Engineering College,
Jabalpur

M. Tech: Indian Institute of Technology,
Bombay

Specialization:

B. E. Hons: Telecommunication Engineering
(1958)

M. Tech: Industrial Electronics (1960)

Experience: Member of staff of Dept. of
Electrical Engineering, Indian
Institute of Technology, Bombay,
India, since 1960.