

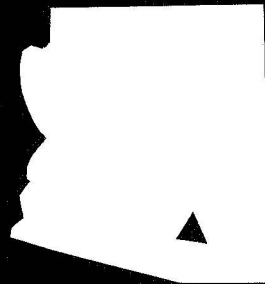
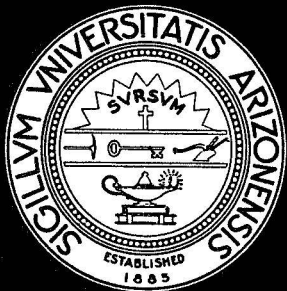
STOCHASTIC ALGORITHMS
FOR SELF-ADAPTIVE FILTERING AND PREDICTION

CASE FILE
COPY

NCL -
NASA 03-002-006

Semi-Annual Report - Part A

January 1971



*ENGINEERING EXPERIMENT STATION
COLLEGE OF ENGINEERING
THE UNIVERSITY OF ARIZONA
TUCSON, ARIZONA*

STOCHASTIC ALGORITHMS FOR SELF-ADAPTIVE FILTERING AND PREDICTION

by

Robert Lee Thomas Hampton

†

NOL -
NASA 03-002-006

Semi-Annual Report - Part A

January 1971

Donald G. Schultz
Principal Investigator

TABLE OF CONTENTS

	Page
LIST OF ILLUSTRATIONS	vii
LIST OF TABLES	viii
ABSTRACT	ix
 CHAPTER	
1 INTRODUCTION	1
1.1 Prologue	1
1.2 Background and Historical Survey	1
1.3 Statement of the Problem and Previous Results	6
1.4 Approach to the Problem	8
1.5 Organization of the Report	8
2 OPTIMAL FILTER THEORY	11
2.1 Introduction and Organization of the Chapter	11
2.2 Filtering, Prediction, and Smoothing: A Definition	12
2.3 Optimal Estimation: Bayesian Approach	12
2.4 Wiener Filter Theory	17
2.5 Kalman Filter Theory	22
2.6 Theoretical Equivalence of Wiener and Kalman Filter Theory	28
2.7 Prediction Theory	33
2.8 Smoothing Theory	35
2.9 Summary	38
3 THE LEARNING CRITERION	39
3.1 Introduction and Organization of the Chapter	39
3.2 The Unsupervised Learning Criterion.	39
3.3 A Linear Regression Model of the Learning Criterion.	52
3.4 Summary	54
4 DERIVATION OF THE ADAPTIVE ALGORITHMS	55
4.1 Introduction and Organization of the Chapter	55
4.2 Motivation for the Adaptive Process	55
4.3 Stochastic Approximation	57

Table of Contents (Continued)

Chapter		Page
4.4	Stochastic Algorithms for Learning the Optimum Filter Matrix	60
4.5	Convergence of the Stochastic Learning Algorithms . .	67
4.6	Acceleration of Convergence	72
4.7	Summary	78
5	ADDITIONS AND EXTENSIONS	79
5.1	Introduction and Organization of the Chapter	79
5.2	Estimation of Σ , Γ , R and Q	79
5.3	Non-Existence of H Inverse	85
	5.3.1 Measurement Stacking With Time Delay	86
	5.3.2 Measurement Stacking With State Augmentation .	87
	5.3.3 Predicted Residual Stacking	89
5.4	Time-Varying Signal and Noise Statistics	91
	5.4.1 Measurement Discounting	92
	5.4.2 Known Time Evolution of $K(n)$	93
5.5	Nonlinear Dynamics	94
5.6	Summary	96
6	NUMERICAL SIMULATION OF THE ADAPTIVE ALGORITHMS	97
6.1	Introduction and Organization of the Chapter	97
6.2	Linear System Simulations	97
	6.2.1 First Order Dynamics	99
	6.2.2 Second Order Dynamics	104
	6.2.3 Third Order Dynamics	108
	6.2.4 Fourth Order Dynamics	124
6.3	Nonlinear Dynamics	135
	6.3.1 Nonlinearities That Satisfy a Global Lipschitz Condition	135
	6.3.2 Nonlinearities That Have Unique Inverses	143
6.4	Application to a Thermionic Reactor Model	145
6.5	Conclusion	151

Table of Contents (Continued)

Chapter	Page
7 EPILOGUE	153
7.1 Summary	153
7.2 Recommendations for Further Study	154
APPENDIX A: LISTING OF THE SIMULATION PROGRAM	155
APPENDIX B: TONELLI'S THEOREM	162
LIST OF REFERENCES	163

LIST OF ILLUSTRATIONS

Figure		Page
2.1	Wiener Filter Model	20
2.2	Kalman Filter Model	27
4.1	Block Diagram of the Feedback System for the Deterministic Algorithm (4.10)	62
4.2	Block Diagram of the Feedback System for the Adaptive Stochastic Algorithm (4.13)	64
6.1	Time Evolution of K_n for Example 6.1: $K_0 = 0.0$	101
6.2	Time Evolution of K_n for Example 6.1: $K_0 = 1.0$	102
6.3	Time Evolution of K_n for Example 6.3: Symmetry Forced . . .	105
6.4	Time Evolution of K_n for Example 6.3: Symmetry Not Forced .	109
6.5	Time Evolution of K_n for Example 6.4	114
6.6	Time Evolution of K_n for Example 6.5	126
6.7	Time Evolution of K_n for Example 6.6	137
6.8	Time Evolution of K_n for Example 6.7	139
6.9	Time Evolution of K_n for Example 6.8	141
6.10	Time Evolution of K_n for Example 6.9	142
6.11	Time Evolution of K_n for Example 6.10	144
6.12	Time Evolution of K_n for Example 6.11	146

LIST OF TABLES

Table		Page
6.1	Comparison of the Estimates and the Optimal Values for Example 6.1	100
6.2	Comparison of the Estimates and the Optimal Values for Example 6.2	104
6.3	Comparison of the Adaptively Filtered and Non-Filtered Results for Example 6.6	136
6.4	Comparison of the Adaptively Filtered and Non-Filtered Results for Example 6.7	140
6.5	Comparison of the Adaptively Filtered and Non-Filtered Results for Example 6.9	143
6.6	Comparison of the Adaptively Filtered and Non-Filtered Results for Example 6.10	143
6.7	Comparison of the Adaptively Filtered and Non-Filtered Results for Example 6.11	145

ABSTRACT

The contribution of this dissertation is the derivation of adaptive stochastic algorithms for learning directly the stationary discrete time optimal Kalman filter in an unknown random environment. To accomplish this derivation, an unsupervised learning criterion is formulated. It is shown that this criterion is a necessary and sufficient condition for optimal filtering. The criterion is then used in conjunction with the theory of Stochastic Approximation to develop the algorithms and to prove that they converge with probability one.

In addition, a method for accelerating the adaptation rate is given. Techniques for estimating other useful quantities such as the error covariance are also presented.

In the latter part of this study the stochastic algorithms are modified to accommodate slowly varying statistics and special classes of nonlinear systems. Numerical simulations are performed to verify experimentally the theoretical expectations and predictions for both linear and nonlinear systems. The adaptive technique is then applied to learning the optimal filter for a thermionic point reactor using the "linearized prompt jump approximation" to the kinetics model of the reactor.

CHAPTER 1

INTRODUCTION

1.1 Prologue

Recently considerable attention has been directed toward self-adaptive (or self-learning) techniques for optimizing system performance. The basic idea is quite simple: one wishes to design a system to operate efficiently in an unknown or changing environment without the necessity of direct human intervention. Such systems are extremely important in the context of control and communication theory where it is often impractical or impossible to obtain the a priori information required to specify the optimum system.

The goal of this report is to provide a self-adaptive solution to the problem of optimal filtering, prediction, and smoothing of stochastic signals imbedded in random noise. However, before discussing the specific problem in more detail, a historical survey of this topic is appropriate.

1.2 Background and Historical Survey

One of the most important topics in control theory is the stochastic control problem. Here one is required to determine the optimum controller for a given plant without precise knowledge of the state $\underline{x}(t)$ of the plant. The stochastic approach to optimum control is motivated by the fact that in general:

- 1) some of the state variables are not available for measurement,
- 2) the measurements contain noise,
- 3) the plant is subject to random input disturbances.

By using the familiar state variable representation, a discrete time linear dynamic system model of the plant can be described by

$$\underline{\dot{x}}(n+1) = \Phi(n+1,n) \underline{\dot{x}}(n) + D(n) \underline{\dot{u}}(n) + \underline{\dot{w}}(n) \quad (1.1)$$

and the measurement of the state $\underline{\dot{x}}(n)$ by

$$\underline{\dot{z}}(n) = H(n) \underline{\dot{x}}(n) + \underline{\dot{v}}(n) \quad (1.2)$$

where

$\underline{\dot{x}}(n)$ is the $m \times 1$ system state vector

$\underline{\dot{u}}(n)$ is the $\ell \times 1$ control vector

$\underline{\dot{w}}(n)$ is the $m \times 1$ white perturbation noise input vector

$\Phi(n+1,n)$ is the one-step $m \times m$ state transition matrix

$D(n)$ is the $m \times \ell$ control matrix

$\underline{\dot{z}}(n)$ is the $p \times 1$ measurement vector

$\underline{\dot{v}}(n)$ is the $p \times 1$ white measurement noise vector

$H(n)$ is the $p \times m$ observation matrix.

The properties of the noise vectors $\underline{\dot{v}}(n)$ and $\underline{\dot{w}}(n)$ are defined by

$$E\{\underline{\dot{w}}(n) \underline{\dot{w}}^T(n)\} = Q(n) \delta(n-s)$$

$$E\{\underline{\dot{v}}(n) \underline{\dot{v}}^T(n)\} = R(n) \delta(n-s)$$

$$E\{\underline{\dot{v}}(n) \underline{\dot{w}}^T(s)\} = [0] \text{ for all } n,s$$

where

$Q(n)$ is the $m \times m$ perturbation noise covariance matrix

$R(n)$ is the $p \times p$ observation noise covariance matrix

and

$\delta(\cdot)$ is the Kronecker delta function.

The approach [Lee, 1964] generally used to attack this problem is to first estimate the state $\underline{x}(n)$. Then this estimate $\hat{\underline{x}}(n)$ is used as if it were the actual state to calculate the optimum control employing deterministic methods, such as the maximum principle. In other words, the stochastic control problem is separated into two phases, referred to as estimation and control. It has been proved that for linear systems with a quadratic performance index and subjected to white Gaussian perturbation and observation noise inputs, the optimal stochastic controller consists of an optimal estimator (filter) in cascade with an optimal deterministic controller [Joseph and Tou, 1961]. This result is known as the Separation Theorem. In this paper, only the estimation phase is considered because the deterministic control solution is well known [Schultz and Melsa, 1967 or Sage, 1968].

In communication theory an equally important topic is the stochastic detection problem. Here it is required to determine the optimal receiver for detecting the presence of a stochastic signal $\underline{x}(t)$ imbedded in additive random Gaussian noise $\underline{v}(t)$. If it is assumed that the signal $\underline{x}(t)$ has a rational spectrum, it is possible to represent it as the state of a linear dynamic system with a white Gaussian noise input [Kalman, 1960]. The linear dynamic system is called the signal generating process,

and in discrete time state variable representation is described by

$$\underline{x}(n+1) = \Phi(n) \underline{x}(n) + \underline{w}(n) \quad (1.3)$$

and the measurement of the signal $\underline{x}(n)$ by

$$\underline{z}(n) = H(n) \underline{x}(n) + \underline{v}(n) \quad (1.4)$$

The notation is the same as that of Eqns. (1.1) and (1.2) except $\Phi(n+1,n)$ for $r = 1$ is simplified to $\Phi(n)$. The signal generating process (1.3) is identical to the control plant process (1.1), without the control input $\underline{u}(n)$. However, the Separation Theorem states that the control term can be disregarded in the estimation phase. Therefore, the estimation problem and its solution are identical for both control and communication theory.

There also exists an analogous Separation Theorem solution to the stochastic detection problem [Kailath, 1963] which states that for a Gaussian signal with rational spectrum observed in white additive Gaussian noise, the optimal stochastic detector consists of an optimal estimator (filter) in cascade with the optimal detector for a completely known signal, i.e., the output of the filter is considered to be the actual signal. Again, only the estimation phase is considered, since the optimum detection solution is well known [Hancock and Wintz, 1966 or Van Trees, 1968].

Because of the identical mathematical framework of estimation in a control or communication context, no distinction between the two areas is made in the text that follows.

Wiener [1949] and Kolmogorov [1941] are credited with the solution for a single input-single output system. Wiener formulated the problem in terms of finding the optimum (in a minimum mean-square error sense) linear filter. He showed that a necessary and sufficient condition for optimality is that the filter satisfy the Wiener-Hopf equation, and he developed a method (spectral-factorization) for solving this equation for signals with a known stationary rational spectrum and for noise with a known stationary white spectrum.

Following Wiener's pioneering work, there developed an extensive literature which interpreted, simplified, modified, and extended his results. Detailed bibliographies may be found in Stumper [1955] and Balakrishnan [1963].

However, the case of a non-stationary multidimensional signal in non-stationary multidimensional noise remained unsolved in an engineering sense until 1960-1961 when Kalman [1960] and Kalman and Bucy [1961] published their fundamental papers. Instead of seeking a solution to the Wiener-Hopf equation in the frequency domain with the attendant problem of spectral factorization, Kalman combined the concept of state variable representation of dynamic systems with the orthogonal projection in a Hilbert space formulation of linear filtering to obtain a direct solution in the time domain. In contrast to Wiener's method, Kalman's results are in recursive form, and therefore ideally suited to real-time sequential digital computation. However, both the Wiener and Kalman theories require complete knowledge of the message generating

and observation noise covariance matrices, denoted by Q and R respectively. Kalman's method also requires the initial value of the error covariance matrix.

In the real world such extensive a priori information is generally not available. The consequence of not knowing R and/or Q is a suboptimal filter, i.e., an increase in the error covariance matrix. In some cases the increase is unbounded [Sorenson, 1966]. Detailed investigations of the suboptimal performance caused by insufficient a priori information have been widely reported in the literature, e.g., Soong [1965], Heffes [1966], and Nishimura [1966, 1967].

In addition, the inverse of R is often required to exist to avoid ill-conditioned matrices in the Kalman filter computation. The presence of either noiseless measurements or correlated observation noise can render R singular. In practice R itself is often ill-conditioned simply because one measurement is an order of magnitude more accurate than the others. Thus the Kalman filter formulation often generates application difficulties. Bryson and Johansen [1965] and Bryson and Mehra [1968] have modified the Kalman framework to handle aspects of this particular problem, but their technique necessitates state augmentation, which increases the dimension of the filter, increases the computation time, and requires taking derivatives.

1.3 Statement of the Problem and Previous Results

The inadequacy or absence of a priori knowledge leads naturally to the consideration of adaptive or learning approaches to optimum

estimation. Specifically, a self-adaptive solution to the discrete time, stationary optimum filtering and prediction problem is sought which does not require a priori specification of R and Q, and which retains the recursive features of Kalman's formulation.

Previous adaptive techniques can be divided into two types. The first, due to Magill [1965], assumes that the parameters of R and Q belong to a finite ensemble of a priori known possibilities. An optimum Bayesian pattern recognition algorithm for Gaussian distributions is used to learn which sampled data process is being observed. With this knowledge, Q and R are uniquely specified. Magill's method is valid only for a scalar observation process and is cumbersome to apply. For example, given N unknown elements of Q and the unknown scalar element R, with M possible values for each variance, there are $(N+1)^M$ combinations. Each combination requires the implementation of the corresponding Kalman filter equations. Hilborn and Lainiotis [1969] extended Magill's technique to a vector observation process and prove mean square and probability one convergence. In their case, there exist $(N+P)^M$ combinations where P is the number of unknown elements in the matrix R.

The second approach is to estimate directly the components of R and Q. Shellenbarger [1966,1967] devised a scheme using the likelihood principle to accomplish this estimation under the assumption of Gaussian distributions and other more restrictive requirements which limit its utility. One of these restrictions is fundamental. It implies that Q cannot be learned if it contains more than $m \times p$ elements. Proof of convergence was not considered. It is important to note that

these two approaches require the estimation of both the R and Q matrices. Then the entire set of Kalman's equations must be solved for the filter matrix each time the estimates of R and Q are updated.

1.4 Approach to the Problem

In this thesis an unsupervised learning criterion is formulated from which self-adaptive algorithms are derived. These algorithms learn the optimum discrete time stationary Kalman filter directly. This eliminates both the necessity of estimating R and Q as an intermediate step and the need to solve the entire set of filtering equations. The number of parameters to be determined and the computation time is thereby reduced. In addition, the problem associated with the occurrence of an ill-conditioned estimate of the R matrix is avoided. Satisfaction of the learning criterion is shown to be a necessary and sufficient condition for optimal filtering. The stochastic algorithms developed for estimating the optimum filter converge in mean-square and with probability one. The results are valid for scalar and vector valued signal and noise processes.

The learning criterion and the associated stochastic algorithms and their application constitute the principal contribution of this research paper.

1.5 Organization of the Dissertation

The second chapter presents and compares Wiener and Kalman filter theory. These results provide the foundation for the adaptive learning criterion.

Chapter 3 formulates the learning criterion and proves that it is necessary and sufficient for optimal filtering. A Markov sequence model of the learning criterion is given.

The unsupervised stochastic algorithms required for performing the adaptation indicated by the learning criterion are derived in Chapter 4. The theory of Stochastic Approximation is invoked to prove convergence of the algorithms to the optimum Kalman filter matrix. In addition, a method for optimizing the adaptive process is developed.

Chapter 5 contains additions and extensions to the previous results. Algorithms are derived for estimating Q , R , the error covariance P and the one-step prediction error covariance Σ . These quantities are required for optimum smoothing. The rather restrictive assumption, needed in Chapter 4, that $(H)^{-1}$ exist is removed. The algorithms of Chapter 4 are also modified to accommodate slowly time-varying signal and noise statistics and altered to include special classes of nonlinear dynamics.

The experimental results obtained by simulating the learning algorithms on the CDC 6400 digital computer are presented in Chapter 6. Many different systems are examined in detail. These investigations serve to confirm experimentally the conclusions of Chapters 4 and 5.

Chapter 7 summarizes the results of this study along with recommendations for further research.

The contribution of this report is the development of an unsupervised learning criterion which is a necessary and sufficient condition for optimal filtering, the derivation from this criterion of adaptive

stochastic algorithms for learning the optimum filter matrix directly, and the application of the algorithm to specific problems of communication and control theory.

CHAPTER 2

OPTIMAL FILTER THEORY

2.1 Introduction and Organization of the Chapter

The subject of this chapter is estimation theory, and its application to the engineering problem of extracting a stochastic signal from noise or estimating the state of a control system from noisy measurements. Section 2.2 defines the terms filtering, prediction, and smoothing which are used throughout the remainder of the chapter. Section 2.3 is devoted to some of the more important aspects of estimation theory. The application of this theory to the field of filtering and prediction leads to the Wiener and Kalman methodologies which are developed in Sections 2.4 and 2.5, respectively, and compared in Section 2.6. This comparison unifies the more general recursive structure of Kalman's solution to the time-varying state estimation problem and the non-recursive structure of Wiener's earlier solution to the stationary signal extraction problem.

The two theories are thus related as a prelude to developing a self-adaptive method of solving the stationary state estimation (or signal extraction) problem which retains the recursive features of Kalman's approach. Specifically, the Wiener-Hopf equation of Section 2.4 provides the theoretical foundation for the adaptive learning criterion, but its proof and application utilizes the recursive filter of Kalman presented

in Section 2.5. Prediction and smoothing are treated briefly in Sections 2.7 and 2.8.

2.2 Filtering, Prediction, and Smoothing: A Definition

At this point it is convenient to specify precisely what is meant by the terms filtering, prediction and smoothing of a stochastic signal $\underline{x}(t)$ observed in additive noise $\underline{v}(t)$.

Definition: Given the measurement time series¹ $\{z(kT) = H\underline{x}(kT) + \underline{v}(kT) : k \in n-r, n-r+1, \dots, n\}$ which is the sum of the random processes $\underline{x}(t)$ and $\underline{v}(t)$, observed over the discrete time interval $n-r, n-r+1, \dots, n$ where n, r are positive integers such that $n > r$,

- (1) Filtering is the estimation of $\underline{x}(\tau)$ at $\tau = nT$.
- (2) Prediction is the estimation of $\underline{x}(\tau)$ for $\tau > nT$.
- (3) Smoothing is the estimation of $\underline{x}(\tau)$ for $\tau < nT$.

All three cases are considered in the succeeding pages, but the emphasis is on filtering because it is the key operation.

Note that even though $\underline{x}(t)$ and $\underline{v}(t)$ may be continuous functions of time, the data $\underline{z}(nT)$ is observed only at discrete times. That is, in this work only sampled data is considered.

2.3 Optimal Estimation: Bayesian Approach

To discuss optimality, a criterion of optimality must be defined. Suppose that a random variable x is to be estimated from the data set $Z = \{z(1), \dots, z(n)\}$. Then $\hat{x}$ is called the optimal estimate of x given

¹The capital T in the argument is generally suppressed.

Z if and only if the average loss

$$E\{\ell(x-\hat{x})\} = E_Z\{E_X[\ell(x-\hat{x})|Z]\} = E_Z\{L(\hat{x}|Z)\} \quad (2.1)$$

is a minimum, where $\ell(x-\hat{x})$ is an appropriately defined loss function.

In Eq. (2.1) the expectation with respect to Z is not dependent upon $\hat{x}$; therefore it suffices to choose $\hat{x}$ such that

$$L(\hat{x}|Z) = E_X\{\ell(x-\hat{x})|Z\} \quad (2.2)$$

is minimized. A solution based on minimizing (2.1) or, equivalently, Eq. (2.2) is called a Bayes estimator. It has been shown [Sherman, 1955, 1958] that for a rather general class of loss function $\ell(\cdot)$ and a posterior densities that the Bayesian estimate is the conditional mean

$$\hat{x} = E\{x|Z\} \quad (2.3)$$

THEOREM 2-1. Let $S = \{\ell(\cdot): \ell \text{ is even and convex}\}$. If the a posteriori density $p(x|Z)$ is even about its conditional mean $E\{x|Z\}$, then the conditional mean $E\{x|Z\}$ is the optimum estimate of x given Z in the sense that it minimizes Eq. (2.2) for all $\ell \in S$.

Proof:

$$\ell(e) = \ell(-e) \quad \text{evenness} \quad (a)$$

$$\ell(ae_1 + be_2) \leq a\ell(e_1) + b\ell(e_2) \quad e_1, e_2 \quad \text{convexity} \quad (b)$$

where $a + b = 1$, $a \in (0,1)$

$$\text{and } p(y|Z) = p(-y|Z) \quad \text{evenness} \quad (c)$$

where $y = x - E\{x|Z\}$

$$\begin{aligned}
L(\hat{x}|Z) &= E_x\{\ell(x-\hat{x})|Z\} \\
&= E_x\{\ell(\hat{x}-x)|Z\} && \text{by (a)} \\
&= E_y\{\ell(\hat{x} - E[x|Z] - y)|Z\} \\
&= E_y\{\ell(\hat{x} - E[x|Z] + y)|Z\} && \text{by (c)} \\
&= E_y\{\ell(E[x|Z] - \hat{x} - y)|Z\} && \text{by (a)} \\
&= E_y\{\ell(E[x|Z] - \hat{x} + y)|Z\} && \text{by (c)} \\
&= E_y\{\ell(y - \{E[x|Z] - \hat{x}\})|Z\} \\
&= \frac{1}{2}E_y\{\ell(y + \{E[x|Z] - \hat{x}\}) \\
&\quad + \ell(y - \{E[x|Z] - \hat{x}\})|Z\} && \text{by (a)} \\
&\geq E_y\{\ell(\frac{1}{2}[y + \{E[x|Z] - \hat{x}\}] \\
&\quad + \frac{1}{2}[y - \{E[x|Z] - \hat{x}\}])|Z\} && \text{by (b)} \\
&= E_y\{\ell(y)|Z\} \text{ with equality iff } \hat{x} = E\{x|Z\} \\
&&& \text{Q.E.D.}
\end{aligned}$$

The class S can be greatly extended if two restrictions are added to the conditional density function.

THEOREM 2-2. Let $S = \{\ell(\cdot): \ell \text{ is even and } \ell(e_2) \geq \ell(e_1) \geq 0 \text{ for } e_2 \geq e_1 \geq 0, \ell(0) = 0\}$. If the a posteriori density $p(x|Z)$ is

- (1) even and monotone non-increasing about its conditional mean and
- (2) decreases rapidly enough so that $\lim_{y \rightarrow \infty} [\ell(y)p(y|Z)] = 0$,
where $y = x - E\{x|Z\}$

then $E\{x|Z\}$ is the optimum Bayes estimate.

Proof: [Viterbi, 1966, p. 308].

Some examples of the $\ell(\cdot) \in S$ are

$$\ell(e) = \begin{cases} C, & |e| > c \\ 0, & |e| \leq c \end{cases}$$

$$\ell(e) = C |e|$$

$$\ell(e) = C [1 - \exp(-e^2)] .$$

Note that under the conditions of Theorem 2 the conditional mean $E\{x|Z\}$ coincides with the maximum a posteriori estimate.

In general this study deals with vector-valued random processes $\underline{x}(n)$. Equation 2.2 then becomes

$$L(\hat{\underline{x}}(n)|Z(n)) = E_{\underline{x}}\{\ell[\|\underline{x}(n) - \hat{\underline{x}}(n)\|] | Z(n)\} \quad (2.4)$$

where $Z(n) = \{z(i), \dots, z(n)\}$ and $\|\cdot\|$ denotes the norm of the vector. Theorems 2-1 and 2-2 extend readily to include this case [Kalman, 1960].

If the error squared is chosen as the loss function, then restrictions (1) and (2) of Theorem 2-2 on the a posteriori density are unnecessary.

THEOREM 2-3. Let $\ell[\|\underline{x}(n) - \hat{\underline{x}}(n)\|] = \|\underline{x}(n) - \hat{\underline{x}}(n)\|^2$

$$= [\underline{x}(n) - \hat{\underline{x}}(n)]^T [\underline{x}(n) - \hat{\underline{x}}(n)],$$

then $\hat{\underline{x}}(n) = E\{\underline{x}(n)|Z(n)\}$ minimizes (2.4) without the restraints of Theorem 2-2 on $p(\underline{x}(n)|Z(n))$.

Proof:

$$\begin{aligned}
 E_{\underline{x}}\{[\underline{x} - \hat{\underline{x}}]^T[\underline{x} - \hat{\underline{x}}]|Z\} &= \hat{\underline{x}}^T \hat{\underline{x}} - 2\hat{\underline{x}}^T E\{\underline{x}|Z\} + E\{\underline{x}^T \underline{x}|Z\} \\
 &= (\hat{\underline{x}} - E\{\underline{x}|Z\})^T (\hat{\underline{x}} - E\{\underline{x}|Z\}) + E\{\underline{x}^T \underline{x}|Z\} - [E\{\underline{x}|Z\}]^T E\{\underline{x}|Z\} \\
 &\geq E\{\underline{x}^T \underline{x}|Z\} - [E\{\underline{x}|Z\}]^T E\{\underline{x}|Z\}
 \end{aligned}$$

with equality iff $\hat{\underline{x}}(n) = E\{\underline{x}(n)|Z(n)\}$ Q.E.D.

The contents of theorems 2-1, 2-2, 2-3 give the "in principle" solution of the Bayesian estimation problem for a wide class of loss functions and probability structures. However, the explicit computation of the optimum filtering estimate $E\{\underline{x}(n)|Z(n)\}$ is formidable except in the important case when $\{\underline{x}(n)\}$ and $\{\underline{z}(n)\}$ are Gaussian. Here the well-known result that $E\{\underline{x}(n)|Z(n)\}$ is a linear function of the data $Z(n)$ may be exploited.

THEOREM 2-4: Let $\{\underline{x}(n)\}$ and $\{\underline{z}(n)\}$ be Gaussian random vector sequences,

$$\text{then} \quad E\{\underline{x}(n)|Z(n)\} = \sum_{j=1}^n A(n,j) \underline{z}(j) = \hat{\underline{x}}(n) \quad (2.5)$$

where the $A(n,j)$ are $m \times p$ matrices and represent the response of the optimum filter matrix at time n to an impulse applied at time j .

Proof: [Deutsch, 1965, p. 32].

As mentioned briefly in Chapter 1, the classical method of specifying the optimal linear transformations $A(n,j)$ is through the Wiener-Hopf equation. This important theory is the subject of the next section and provides the theoretical basis for the adaptive learning

criterion developed in Chapter 3. Kalman's more recent contribution, which utilizes the orthogonal projection theorem to determine the optimal $A(n,j)$, is presented in Section 2.5. His results provide the recursive solution to the optimum filter problem used in proving the learning criterion.

2.4 Wiener Filter Theory

The treatment here of optimal filter theory is in the spirit of Wiener's original work on single input-single output continuous time systems, but the formulation is for multiple input-multiple output sampled data systems.

THEOREM 2-5: The Wiener-Hopf Equation

Let $\{\tilde{x}(n)\}$ and $\{\tilde{z}(n)\}$ be zero mean random sequences.

If either

- (i) the random sequences are Gaussian or
- (ii) the estimator $\hat{\tilde{x}}(n)$ is restricted to be a linear function of the data

$$Z(n) = \{\tilde{z}(1), \dots, \tilde{z}(n)\} \quad \text{and}$$

$$\begin{aligned} E(\|\tilde{x}(n) - \hat{\tilde{x}}(n)\|^2) &= \|\tilde{x}(n) - \hat{\tilde{x}}(n)\|^2 \\ &= [\tilde{x}(n) - \hat{\tilde{x}}(n)]^T [\tilde{x}(n) - \hat{\tilde{x}}(n)] \end{aligned}$$

a necessary and sufficient condition for

$$\hat{\tilde{x}}(n) = \sum_{j=1}^n A(n,j) \tilde{z}(j) \tag{2.5}$$

to be the optimum estimate of $\underline{x}(n)$ is that the filtering transformations $A(n,j)$ be chosen such that $\hat{\underline{x}}(n)$ satisfies the Wiener-Hopf equation

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{z}^T(v)\} = [0] \quad \text{for all } \underline{z}(v) \in Z(n) \quad (2.6)$$

Proof: See Theorem 2-8 where it is shown that the orthogonal projection theorem implies the Wiener-Hopf equation and conversely.

The discrete multidimensional Wiener-Hopf equation may be put in a more familiar form by substituting Eq. (2.5) into Eq. (2.6) yielding

$$\begin{aligned} E\{[\underline{x}(n) - \sum_{j=i}^n A(n,j) \underline{z}(j)] \underline{z}^T(v)\} \\ &= E\{\underline{x}(n) \underline{z}^T(v)\} - \sum_{j=i}^n A(n,j) E\{\underline{z}(j) \underline{z}^T(v)\} \\ &= R_{xz}(n,v) - \sum_{j=i}^n A(n,j) R_{zz}(j,v) \\ &= [0] \end{aligned} \quad (2.7)$$

where

$$R_{xz}(n,v) = E\{\underline{x}(n) \underline{z}^T(v)\}$$

and

$$R_{zz}(j,v) = E\{\underline{z}(j) \underline{z}^T(v)\}$$

If $i = -\infty$, Eq. (2.7) is assumed to satisfy the conditions of Theorem B.1 in the Appendix so that interchanging the order of expectation and summation is valid.

Wiener's approach¹ to explicitly determining the optimum filter consists of transforming Eq. (2.6) into the frequency domain² and using spectral-factorization. The solution is given in terms of the z-transform $A(z)$ of the $m \times p$ filter matrix $A(n-j)$ as illustrated in Fig. 2.1. The problem of synthesizing the filter remains. Employment of the frequency domain requires the following restrictions on the system defined by Eqns. (1.3) and (1.4)

- (1) The state transition matrix Φ and observation matrix H must be time invariant.
- (2) The statistics of $w(n)$ and $v(n)$ must be stationary.
- (3) The data $z(n)$ must be known for all past time, i.e.,
 $i = -\infty$ in Eq. (2.7).

Under these conditions, Eqns. (2.5) and (2.7) become

$$\hat{\tilde{x}}(n) = \sum_{j=-\infty}^n A(n-j) \tilde{z}(j) \quad (2.8)$$

and

$$\begin{aligned} E\{\tilde{x}(n) \tilde{z}^T(v)\} &= \sum_{j=-\infty}^n A(n-j) E\{\tilde{z}(j) \tilde{z}^T(v)\} \\ &= R_{xz}(n-v) - \sum_{j=-\infty}^n A(n-j) R_{zz}(j-v) \\ &= [0], \quad v \in (-\infty, \dots, n) \end{aligned} \quad (2.9)$$

1. The first general technique for determining the optimum multiple input-multiple output discrete filter using Wiener's method was given by Motyka and Cadzow [1967]. The results presented in this section follow their work.

2. Since only sampled-data is considered, frequency domain means the Z-transform domain.

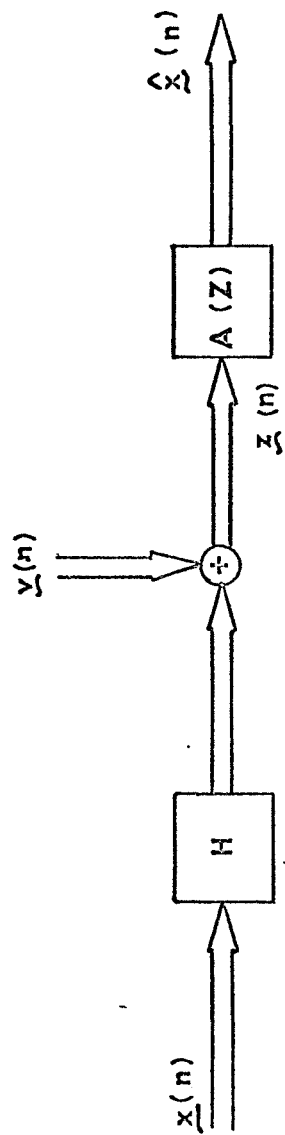


Fig. 2.1 Wiener Filter Model

By a change of variable (2.9) can be rewritten

$$R_{xz}(\alpha) - \sum_{\tau=0}^{\infty} A(\tau) R_{zz}(\alpha-\tau) = [0], \alpha \in (0, \dots, \infty) \quad (2.10)$$

which is just the discrete time stationary Wiener-Hopf equation.

The cross-spectral (generating function) matrix representation of (2.10) is

$$\phi_{xz}(Z) - A(Z) \phi_{zz}(Z) = [0] \quad (2.11)$$

Since each element of (2.10) is zero for each α , each element of (2.11) is a polynomial in Z only. Thus each polynomial element converges for all Z inside the unit circle. Assuming $\phi_{zz}(Z)$ has a spectral factorization of the form

$$\phi_{zz}(Z) = \Delta(Z^{-1}) \Delta^T(Z)$$

where $\Delta(Z)$ is a $p \times p$ matrix whose elements represent the Z -transforms of stable, linear, causal systems (i.e., polynomials in Z) and have no poles outside the unit circle, then a physically realizable Wiener filter exists. The frequency domain expression for this optimum $m \times p$ filter matrix is

$$A(Z) = [\{\Delta^{-1}(Z^{-1}) \phi_{xz}(Z)\}_c + \{\Delta^{-1}(Z^{-1}) \phi_{xz}(Z)\}_+^T [\Delta(Z)]^{-1}] \quad (2.12)$$

where $\{\}_c$ is a matrix whose elements are constant and $\{\}_+$ is a matrix whose elements contain only poles inside the unit circle. The "in principle" solution given by (2.12) is cumbersome, not easy to synthesize, and not suited to machine computation.

Kalman's time domain approach not only eliminates these two difficulties, but also the three restrictions listed previously. He used the orthogonal projection theorem to obtain a recursive set of equations for optimum filtering and prediction.

2.5 Kalman Filter Theory

By applying the orthogonal projection theorem a different specification for the optimum linear transformations $A(n,j)$ is obtained.

THEOREM 2-6: Orthogonal Projection

Let $\{\underline{x}(n)\}$ and $\{\underline{z}(n)\}$ be zero mean random sequences. Let $Y(n)$ be the linear manifold generated by the data:

$$Y(n) = \{\underline{y}(n): \underline{y}(n) = \sum_{j=1}^n B(n,j)\underline{z}(j)\}$$

where the $B(n,j)$ are arbitrary $m \times p$ matrices whose elements are real.

Note that

$$\hat{\underline{x}}(n) = \sum_{j=1}^n A(n,j)\underline{z}(j) \quad (2.5)$$

must belong to $Y(n)$.

If either

- (i) the random sequences $\{\underline{x}(n)\}$ and $\{\underline{z}(n)\}$ are Gaussian or
- (ii) the estimator $\hat{\underline{x}}(n)$ is specified to be a linear function of the data $\{\underline{z}(1), \dots, \underline{z}(n)\}$ and the loss function is the error squared

$$\ell(\underline{e}(n)) = \underline{e}^T(n)\underline{e}(n) = \|\underline{e}(n)\|^2$$

where the error $\underline{e}(n) = \underline{x}(n) - \hat{\underline{x}}(n)$, then the filtering estimate $\hat{\underline{x}}(n)$ of $\underline{x}(n)$ is optimum if and only if the $A(n,j)$ are chosen such that the error

$$\underline{e}(n) = \underline{x}(n) - \hat{\underline{x}}(n) = \underline{x}(n) - \sum_{j=i}^n A(n,j) \underline{z}(j)$$

is orthogonal to $Y(n)$, i.e.,

$$(\underline{x}(n) - \hat{\underline{x}}(n), \underline{y}(n)) = 0 \quad \text{for all } \underline{y}(n) \in Y(n) \quad (2.13)$$

where $(\cdot, \cdot)$ is the inner product induced on $Y(n)$ by $E\{(\cdot)^T(\cdot)\}$.

Proof: Note that from Theorem 2-4 the optimum estimate is the conditional mean and it takes the form of Eq. (2.5) for a large class of loss functions, including the error squared criterion. Thus, it suffices to prove only part (ii).

Let $\Theta_n = [A(n,i), \dots, A(n,n)]$ and $\underline{Z}^T(n) = [\underline{z}(i), \dots, \underline{z}(n)]$, where Θ_n is a $m \times (n-i)p$ matrix and $\underline{Z}(n)$ is a $(n-i)p \times 1$ vector.

Substituting Θ_n and $\underline{Z}(n)$ into (2.5) gives

$$\hat{\underline{x}}(n) = \Theta_n \underline{Z}(n) \quad (2.14)$$

The linear transformations $A(n,j)$ embodied in Θ_n minimize $E\{\underline{e}^T(n)\underline{e}(n)\}$ iff for all $m \times (n-i)p$ matrices Δ_n ,

$$\begin{aligned} & E\{[\underline{x}(n) - (\Theta_n + \Delta_n)\underline{Z}(n)]^T [\underline{x}(n) - (\Theta_n + \Delta_n)\underline{Z}(n)] \\ & - [\underline{x}(n) - \Theta_n \underline{Z}(n)]^T [\underline{x}(n) - \Theta_n \underline{Z}(n)]\} \geq 0 \end{aligned} \quad (2.15)$$

Expanding the left side of Ineq. (2.15) leads to

$$\begin{aligned} & E\{[\underline{x}(n) - \Theta_{\underline{n}} \underline{Z}(n)]^T \Delta_{\underline{n}} \underline{Z}(n) + [\Delta_{\underline{n}} \underline{Z}(n)]^T [\underline{x}(n) - \Theta_{\underline{n}} \underline{Z}(n)] + [\Delta_{\underline{n}} \underline{Z}(n)]^T \Delta_{\underline{n}} \underline{Z}(n)\} \\ & = 2E\{[\underline{x}(n) - \Theta_{\underline{n}} \underline{Z}(n)]^T \Delta_{\underline{n}} \underline{Z}(n)\} + E\{[\Delta_{\underline{n}} \underline{Z}(n)]^T \Delta_{\underline{n}} \underline{Z}(n)\} \geq 0 \end{aligned} \quad (2.16)$$

Now $E\{[\Delta_{\underline{n}} \underline{Z}(n)]^T \Delta_{\underline{n}} \underline{Z}(n)\} \geq 0$, and since it can be made arbitrarily small with respect to the term linear in $\Delta_{\underline{n}}$, the sign of the left side of Ineq. (2.16) is determined by this linear term. It, therefore, follows that

$$E\{[\underline{x}(n) - \Theta_{\underline{n}} \underline{Z}(n)]^T \Delta_{\underline{n}} \underline{Z}(n)\} = 0 \text{ for all } \Delta_{\underline{n}} \quad (2.17)$$

If Eq. (2.17) was not true, the sign of Eq. (2.16) would depend on $\Delta_{\underline{n}}$ which would imply a contradiction because $\Delta_{\underline{n}}$ is arbitrary.

Since $\underline{y}(n) = \Delta_{\underline{n}} \underline{Z}(n)$, Eq. (2.17) may be written

$$\begin{aligned} E\{[\underline{x}(n) - \sum_{j=1}^n A(n,j) \underline{z}(j)]^T \underline{y}(n)\} & = E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \underline{y}(n)\} \\ & = (\underline{x}(n) - \hat{\underline{x}}(n), \underline{y}(n)) \\ & = 0 \text{ for all } \underline{y}(n) \in Y(n) \end{aligned} \quad (2.13)$$

The converse follows easily. If Eq. (2.13) is true, then Eq. (2.16) reduces to

$$E\{[\Delta_{\underline{n}} \underline{Z}(n)]^T \Delta_{\underline{n}} \underline{Z}(n)\} \geq 0 \text{ for all } \Delta_{\underline{n}}$$

Hence satisfaction of Eq. (2.13) minimizes the $E\{\underline{e}_{\underline{n}}^T \underline{e}_{\underline{n}}\}$.

Q.E.D.

Under conditions (i) of the theorem, the orthogonal projection of $\underline{x}(n)$ on $Y(n)$ is identical to the conditional mean $E\{\underline{x}(n)|Z(n)\}$. Thus this theorem implies that the optimal linear estimator cannot be improved upon unless the random phenomenon are non-Gaussian and, even then, only if knowledge of at least third order joint probability distribution functions is available.¹ Kalman [1960] employed Theorem 2-6 to solve the optimum filtering problem, but his approach and proof differ from that given here.

COROLLARY 2-7:

$$(\underline{x}(n) - \hat{\underline{x}}(n), \hat{\underline{x}}(n)) = E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \hat{\underline{x}}(n)\} = 0 \quad (2.18)$$

where

$$\hat{\underline{x}}(n) = \sum_{j=i}^n A(n,j) \underline{z}(j) \quad (2.5)$$

By combining the state variable representation of dynamic systems with the orthogonal projection theorem, Kalman [1960] and Kalman and Bucy [1961] obtained the following recursive solution to the non-stationary multidimensional state estimation problem. Given the dynamic system of Eq. (1.3) and the measurement system of Eq. (1.4), then

$$\hat{\underline{x}}(n) = \Phi(n)\hat{\underline{x}}(n-1) + K(n)[\underline{z}(n) - H(n)\Phi(n)\hat{\underline{x}}(n-1)] \quad (2.19)$$

$$= E\{\underline{x}(n)|Z(n)\} \quad \text{for Gaussian noise}$$

with the initial condition $\hat{\underline{x}}(0) = \underline{0}$, and

1. Given any random sequence, there exists a unique Gaussian sequence with the same mean and covariance kernel.

$$\begin{aligned}
K(n) &= \Gamma(n)H^T(n)R^{-1}(n) \\
&= \Sigma(n|n-1) H^T(n) [H(n)\Sigma(n|n-1) H^T(n) + R(n)]^{-1} \quad (2.20) \\
&= \text{Kalman filter matrix}
\end{aligned}$$

$$\begin{aligned}
\Gamma(n) &= E\{[\tilde{x}(n) - \hat{\tilde{x}}(n)][\tilde{x}(n) - \hat{\tilde{x}}(n)]^T\} \\
&= \Sigma(n|n-1) - \Sigma(n|n-1)H^T(n) [H(n)\Sigma(n|n-1)H^T(n) + R(n)]^{-1}H(n)\Sigma(n|n-1) \\
&= [I - K(n)H(n)]\Sigma(n|n-1) \quad (2.21) \\
&= \text{Error covariance matrix}
\end{aligned}$$

with the initial condition $\Gamma(0) = \Gamma_0$

$$\begin{aligned}
\Sigma(n+1|n) &= E\{[\tilde{x}(n+1) - \Phi(n)\hat{\tilde{x}}(n)][\tilde{x}(n+1) - \Phi(n)\hat{\tilde{x}}(n)]^T\} \\
&= \Phi(n)\Gamma(n)\Phi(n)^T + Q(n+1) \quad (2.22) \\
&= \text{One-step prediction error covariance matrix}
\end{aligned}$$

The block diagram for the Kalman filter is shown in Fig. 2.2. Note that it is in the form of a non-autonomous feedback system. The recursive nature of the Kalman filter equations makes it ideally suited for digital computation and sequential processing of new data in real time. These attributes are just as outstanding in the stationary system, where Φ , H , R , Q and therefore, K , Γ , ϕ , Σ , are time invariant. In this case, once the optimum Kalman filter matrix K_{opt} is determined, Eq. (2.19) is the only real time calculation required. When the covariance matrices Q of the plant perturbation noise and R of the observation noise are known a priori, K_{opt} can be determined off-line. In this report Q and R are not assumed to be known. The absence of this information necessitates

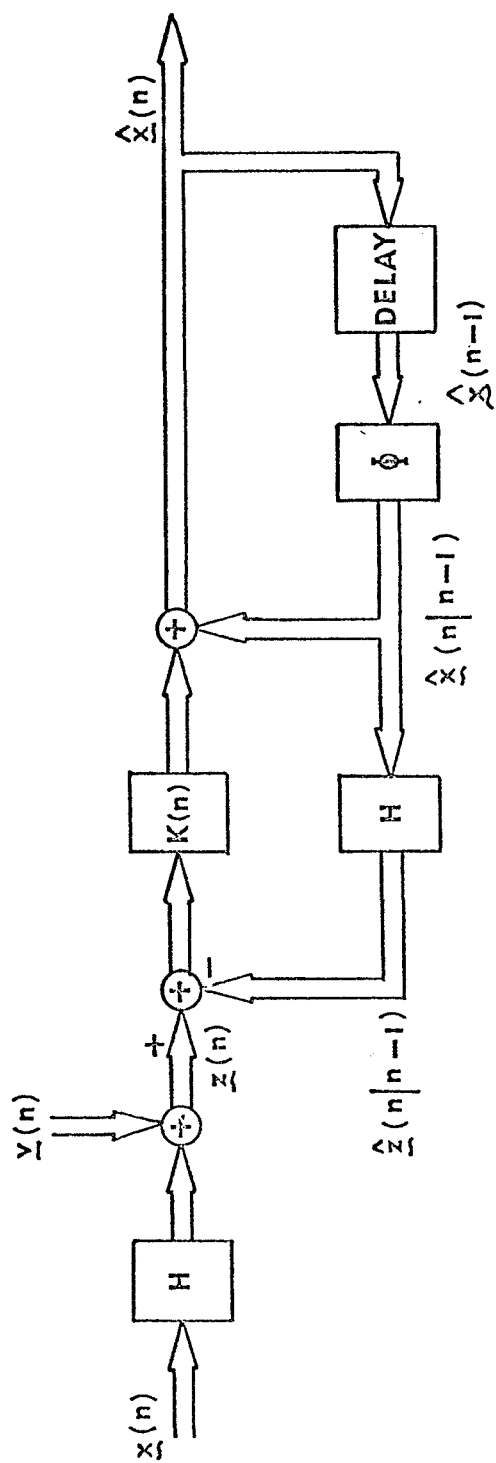


Fig. 2.2 Kalman Filter Model

two on-line recursions: one for estimating K_{opt} and one for estimating $\underline{x}(n)$ with K_{opt} replaced by its estimate.

2.6 Theoretical Equivalence of Wiener and Kalman Filter Theory

In order to provide a common framework for comparing the approaches of Wiener and Kalman, it is necessary to show that the orthogonality of the error $\underline{x}(n) - \hat{\underline{x}}(n)$ and the linear manifold $Y(n)$ implies the Wiener-Hopf equation and conversely. Note that Eq. (2.6) may be viewed as an outer product induced by $E\{(\cdot)(\cdot)^T\}$. With this observation the generalized form of the multidimensional Wiener-Hopf equation is given by the outer product

$$\begin{aligned} [\underline{x}(n) - \hat{\underline{x}}(n)](\underline{z}(v)) &= E\{[\underline{x}(n) - \hat{\underline{x}}(n)]\underline{z}^T(v)\} \\ &= [0] \quad \text{for all } \underline{z}(v) \in Z(n) \end{aligned}$$

The purpose of writing the Wiener-Hopf equation as an outer product is to place it in contradistinction to the principle of orthogonality as given by the inner product Eq. (2.13). However, they are equivalent as proved in the next theorem.

THEOREM 2-8: The Wiener-Hopf Equation

$$\begin{aligned} [\underline{x}(n) - \hat{\underline{x}}(n)](\underline{z}(v)) &= E\{[\underline{x}(n) - \hat{\underline{x}}(n)]\underline{z}^T(v)\} \\ &= [0] \quad \text{for all } \underline{z}(v) \in Z(n) \end{aligned} \tag{2.6}$$

is satisfied if and only if the orthogonal projection theorem

$$\begin{aligned} ([\underline{x}(n) - \hat{\underline{x}}(n)], \underline{y}(n)) &= E\{[\underline{x}(n) - \hat{\underline{x}}(n)]\underline{y}^T(n)\} \\ &= 0 \quad \text{for all } \underline{y}(n) \in Y(n) \end{aligned} \tag{2.13}$$

holds, where

$$\underline{y}(n) = \sum_{k=1}^n B(n,k) \underline{z}(k) \quad (2.16)$$

the $B(n,k)$ are arbitrary $m \times p$ matrices whose elements are real, and

$$\underline{\hat{x}}(n) = \sum_{j=1}^n A(n,j) \underline{z}(j) \quad (2.5)$$

Proof: Assume Eq. (2.13) holds. By substituting Eqns. (2.16) and (2.5) into Eq. (2.13) becomes

$$E\left\{\left[\underline{x}(n) - \sum_{j=1}^n A(n,j) \underline{z}(j)\right]^T \sum_{k=1}^n B(n,k) \underline{z}(k)\right\} = 0 \text{ for all } B(n,k) \quad (2.13)$$

To prove that Eq. (2.13) implies the Wiener-Hopf equation, it is first expanded and rewritten as a sum of its scalar components.

The inner product of Eq. (2.13) is equivalent to the following summations

$$\begin{aligned} &= E\left\{\sum_{a=1}^m [\underline{x}(n) - \sum_j A(n,j) \underline{z}(j)]_a \left[\sum_k B(n,k) \underline{z}(k)\right]_a\right\} \\ &= E\left\{\sum_{a=1}^m \sum_{b=1}^p \left\{[\underline{x}(n) - \sum_j A(n,j) \underline{z}(j)]_a \sum_k B_{ab}(n,k) \underline{z}_b(k)\right\}\right\} \\ &= E\left\{\sum_{a,b} \sum_k [\underline{x}(n) - \sum_j A(n,j) \underline{z}(j)]_a B_{ab}(n,k) \underline{z}_b(k)\right\} \\ &= E\left\{\sum_k \sum_{a,b} B_{ab}(n,k) [\underline{x}(n) - \sum_j A(n,j) \underline{z}(j)]_a \underline{z}_b(k)\right\} \end{aligned}$$

Now by interchanging the linear operations of summation and expectation, the above equation becomes

$$\begin{aligned} \sum_k \sum_{a,b} B_{ab}(n,k) E \left\{ \sum_{c=1}^m [\underline{x}_a(n) - \sum_j A_{ac}(n,j) \underline{z}_c(j)] [\underline{z}_b(k)] \right\} \\ = \sum_k \sum_{a,b} B_{ab}(n,k) \sum_c [E\{\underline{x}_a(n) \underline{z}_b(k)\} - \sum_j A_{ac}(n,j) E\{\underline{z}_c(j) \underline{z}_b(k)\}] \end{aligned}$$

Performing the sum over c gives

$$\sum_k \sum_{a,b} B_{ab}(n,k) [R_{xz}(n,k) - \sum_j A(n,j) R_{zz}(j,k)]_{ab} = 0 \quad (2.23)$$

Since (2.23) is zero for arbitrary $B(n,k)$, the term in brackets must be identically zero for all a , b , and k . Otherwise, a contradiction exists. Thus, it has been shown that each element of Eq. (2.6) is zero.

$$\begin{aligned} \therefore R_{xz}(n,k) - \sum_{j=1}^n A(n,j) R_{zz}(j,k) \\ = E\{[\underline{x}(n) - \sum_{j=1}^n A(n,j) \underline{z}(j)] \underline{z}^T(k)\} \\ = E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{z}^T(k)\} \\ = [0] \text{ for all } \underline{z}(k) \in Z(n) \end{aligned}$$

By reversing the argument, the converse follows readily. Q.E.D.

Corollary 2-9:

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \hat{\underline{x}}^T(n)\} = [0]$$

$$\text{iff } E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \hat{\underline{x}}(n)\} = 0$$

The previous theorem and the accompanying corollary provide a common foundation for the filter theory of Wiener and Kalman.

Even though the theoretical equivalence of Wiener and Kalman filter theory has been demonstrated, the explicit computational relations for the optimum filters of Wiener and Kalman represented by Eqns. (2.12) and (2.20) respectively bear no external resemblance. The solution for the Wiener filter $A(n-v)$ is given in the frequency domain as the Z-transform of the impulse response of the optimum filter matrix. The Kalman filter matrix $K(n)$ is given in the time domain as a function of the error covariance matrix, the observation matrix, and the measurement noise covariance matrix. In the filtering Eq. (2.19), $K(n)$ acts simply as a weighting for the correction term. However, when the restrictions required for the Wiener approach are satisfied, the two filters must be identical. For this case Φ , H , R , and Q are time invariant. Therefore, the optimum Kalman filter, denoted by K_{opt} , is also constant. With these observations the relationship of $A(n-v)$ to K_{opt} is derived. Since the estimate $\hat{\underline{x}}(n)$ is the same regardless of the filter used, it is seen that

$$\hat{\underline{x}}(n) = \Phi \hat{\underline{x}}(n-1) + K_{\text{opt}} [z(n) - H \Phi \hat{\underline{x}}(n-1)] = \sum_{\alpha=-\infty}^n A(n-\alpha) z(\alpha)$$

From Theorem 2-5

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)] z^T(n)\} = E\{[\underline{x}(n) - \hat{\underline{x}}(n)] [\underline{x}^T(n) H^T + \underline{v}^T(n)]\} = [0]$$

or

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{x}^T(n)\} H^T = E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{v}^T(n)\} \quad (2.24)$$

The left hand side of (2.24) is equal to $\Gamma(n)H^T$ from Corollary 2-9, and the right hand side is $E\{\hat{\underline{x}}(n) \underline{v}^T(n)\}$ since $\underline{x}(n)$ is independent of $\underline{v}(n)$. Therefore, (2.24) becomes

$$\begin{aligned} \Gamma(n) H^T &= E\{\hat{\underline{x}}(n) \underline{v}^T(n)\} \\ &= E\left\{ \sum_{\alpha=-\infty}^n A(n-\alpha) \underline{z}(\alpha) \underline{v}^T(n) \right\} \\ &= E\left\{ \sum_{\alpha=-\infty}^n A(n-\alpha) [H\underline{x}(\alpha) + \underline{v}(\alpha)] \underline{v}^T(n) \right\} \\ &= \sum_{\alpha=-\infty}^n A(n-\alpha) E\{\underline{v}(\alpha) \underline{v}^T(n)\} \\ &= A(0)R \end{aligned}$$

which implies

$$A(0) = \Gamma(n) H^T R^{-1}$$

The optimum Kalman filter is by definition (2.15) equal to $\Gamma(n) H^T R^{-1}$. Therefore,

$$A(0) = K_{\text{opt}}$$

Thus the impulse response of the optimum Wiener Filter matrix evaluated at time equal to zero gives the optimum Kalman filter matrix. This means that if a method can be devised to learn $A(0)$ when Q and R are unknown,

then the adaptive filter solution does not require estimation of the entire impulse response of the optimum filter matrix in the Wiener formulation or estimation of the elements of Q and R , and the subsequent necessity of calculating Σ in the Kalman formulation. Such a method represents a potentially great reduction in computation time, memory space, and algorithmic complexity. An iterative scheme for realizing this potential is developed in the next two chapters. First, however, optimal prediction and smoothing are introduced.

2.7 Prediction Theory

The previous sections contain a complete theory for optimal filtering as required in this report. The necessary theorems and their proofs for optimal prediction are directly analogous to those for optimal filtering.

THEOREM 2-10: Let $\{\underline{x}(n)\}$ and $\{\underline{z}(n)\}$ be Gaussian random vector sequences, then

$$E\{\underline{x}(n+r) | Z(n)\} = \sum_{j=i}^n A(r,j) \underline{z}(j) \quad (2.25)$$

$$= \hat{\underline{x}}(n+r|n), \quad r > 0$$

where the $A(r,j)$ are $m \times p$ matrices. $\hat{\underline{x}}(n+r|n)$ is the estimate of $\underline{x}(n)$ given the data up to time r . The optimal linear transformations $A(r,j)$ can be determined from the orthogonal projection theorem.

THEOREM 2-11: Let $\{\underline{x}(n)\}$ and $\{\underline{z}(n)\}$ be zero mean random sequences. Let $Y(r)$ be the linear manifold generated by the data,

$$Y(r) = \{\underline{y}(n) : \underline{y}(n) = \sum_{j=1}^n B(r,j) \underline{z}(j)\}$$

where $\underline{y}(n)$ is any $m \times 1$ vector belonging to $Y(n)$ and the $B(n,j)$ are $m \times p$ real matrices. If either

- (i) the random sequences $\{\underline{x}(n)\}$ and $\{\underline{z}(n)\}$ are Gaussian or
- (ii) the predictor $\hat{\underline{x}}(n+r|n)$ is specified to be a linear function of the data $Z(n) = \{\underline{z}(1), \dots, \underline{z}(n)\}$ and

$$E\{[\underline{x}(n+r) - \hat{\underline{x}}(n+r|n)]\} = \|\underline{x}(n+r) - \hat{\underline{x}}(n+r|n)\|^2$$

Then the prediction $\hat{\underline{x}}(n+r|n)$ is optimum if and only if the $A(r,j)$ are chosen such that the error $\underline{x}(n+r) - \hat{\underline{x}}(n+r|n)$ is orthogonal to $Y(n)$,

$$E\{[\underline{x}(n+r) - \hat{\underline{x}}(n+r|n)]^T \underline{y}(n)\} = 0 \text{ for all } \underline{y}(n) \in Y(n) \quad (2.26)$$

The equivalence of the orthogonal projection theorem and the Wiener-Hopf equation for prediction is given by Theorem 2-12.

THEOREM 2-12: A necessary and sufficient condition for

$$E\{[\underline{x}(n+r) - \hat{\underline{x}}(n+r|n)] \underline{z}^T(k)\} = [0] \text{ for all } \underline{z}(k) \in Z(n)$$

is that

$$E\{[\underline{x}(n+r) - \hat{\underline{x}}(n+r|n)]^T \underline{y}(n)\} = 0 \text{ for all } \underline{y}(n) \in Y(n)$$

Kalman's solution for predicting $\underline{x}(n+r)$ is directly dependent on the optimum filtering equations.

$$\hat{\underline{x}}(n+r|n) = \Phi(n+r, n) \hat{\underline{x}}(n) \quad (2.27)$$

= the prediction of $\underline{x}(n+r)$ given the data

$$\{\underline{z}(1), \dots, \underline{z}(n)\} = \underline{Z}(n)$$

= $E\{\underline{x}(n+r) | \underline{Z}(n)\}$ for Gaussian noise

where $\hat{\underline{x}}(n)$ is the optimum filtering estimate of $\underline{x}(n)$.

In the stationary case Eq. (2.27) reduces to

$$\hat{\underline{x}}(n+r|n) = \Phi^r \hat{\underline{x}}(n) \quad (2.28)$$

Regardless, once the optimum filter is determined, so is the optimum predictor.

2.8 Smoothing Theory

Recall from Section 2.1 that smoothing is the estimate $\hat{\underline{x}}(m|n)$ of $\underline{x}(m)$ at time n , when $m < n$. The optimum smoothing estimate for the class of loss functions of Section 2.2 is given by Theorem 2-13.

THEOREM 2-13: Let $\{\underline{x}(n)\}$ and $\{\underline{z}(n)\}$ be Gaussian random vector sequences, then

$$E\{\underline{x}(m) | \underline{Z}(n)\} = \sum_{j=i}^n A(n, j) \underline{z}(j) = \hat{\underline{x}}(m|n) \quad (2.29)$$

Proof: [Meditch 1967, p. 162]

In practice it is convenient to dichotomize smoothing into three classes

(1) Fixed-Interval Smoothing: $\hat{\underline{x}}(m|N)$ where $m \in [1, \dots, N-1]$

and $n = N$, a fixed integer.

- (2) Fixed-Point Smoothing: $\hat{x}(M|n)$ where $m = M$, a fixed integer and $n \in [M+1, M+2, \dots]$
- (3) Fixed-Lag Smoothing: $\hat{x}(m|m+N)$ where $m \in [0, 1, 2, \dots]$ and N is a fixed integer.

The results presented here for these three classifications essentially follow Meditch [1969]. The interested reader is referred to his work for details and proofs.

THEOREM 2-14: The stationary optimal fixed-interval smoothed estimate $\hat{x}(m|N)$ of $x(m)$ is determined from

$$\hat{x}(m|N) = \hat{x}(m) + F_{\text{opt}} [\hat{x}(m+1|N) - \Phi \hat{x}(m)] \quad (2.30)$$

and the smoothing gain matrix by

$$\begin{aligned} F_{\text{opt}} &= \Gamma \Phi^T \Sigma^{-1} \\ &= [I - K_{\text{opt}} H] \Sigma \Phi^T \Sigma^{-1} \end{aligned} \quad (2.31)$$

As in prediction the fixed-interval smoothing estimate is directly dependent on the optimum filter K_{opt} , but the one-step prediction error covariance Σ must also be known. Alternately, knowledge of the error covariance matrix Γ can be substituted for K_{opt} .

THEOREM 2-15: The stationary optimum fixed-point smoothing estimate $\hat{x}(M|n)$ or $\hat{x}(M)$ is determined from

$$\hat{x}(M|n) = \hat{x}(M|n-1) + S(n)H^T R^{-1} [z(n) - H\Phi \hat{x}(n-1)] \quad (2.32)$$

and the smoothing gain matrix by

$$\begin{aligned}
S(n) &= S(n-1) \Phi^T [I - H^T R^{-1} H \Gamma] \\
&= S(n-1) \Phi^T [I - H^T R^{-1} H (I - K_{opt} H) \Sigma] \quad (2.33)
\end{aligned}$$

with the initial condition

$$S(M) = \Gamma = [I - K_{opt} H] \Sigma$$

Again the fixed-point optimal smoothing matrix is a function of K_{opt} . However, the observation noise covariance matrix R is also required to compute $S(n)$.

THEOREM 2-16: The stationary optimum fixed-lag smoothing estimate $\hat{\tilde{x}}(m|m+N)$ of $\hat{\tilde{x}}(m)$ is determined from

$$\begin{aligned}
\hat{\tilde{x}}(m|m+N) &= \Phi \hat{\tilde{x}}(m-1|m-1+N) + U_{opt} [\hat{\tilde{x}}(m-1|m-1+n) - \hat{\tilde{x}}(m-1)] \\
&\quad + J(m+N) K_{opt} [z(m+N) - H \Phi \hat{\tilde{x}}(m-1+N)] \quad (2.34)
\end{aligned}$$

and the smoothing gain matrices by

$$\begin{aligned}
U_{opt} &= Q[\Phi^{-1}]^T \Gamma^{-1} \\
&= Q[\Phi^{-1}]^T [\Sigma - K_{opt} H \Sigma]^{-1} \quad (2.35)
\end{aligned}$$

$$\begin{aligned}
J(m+N) &= F_{opt}^{-1} J(m-1+N) F_{opt} \\
&= \Sigma [\Sigma \Phi^T - K_{opt} H \Sigma \Phi^T]^{-1} J(m-1+N) [I - K_{opt} H] \Sigma \Phi^T \Sigma^{-1} \quad (2.36)
\end{aligned}$$

with the initial condition

$$J(N) = [F_{opt}]^N$$

This is the most complex smoothing algorithm from both a computational and informational standpoint. It requires knowledge of K_{opt} , Σ , and the plant perturbation noise covariance matrix Q . Numerous matrix inversions must also be calculated.

Sequential methods of learning the Q , R , and Σ are derived in Chapter 6. This permits an adaptive solution to the problem of optimal smoothing without a priori knowledge of Q and R . Since Σ is also being estimated directly, the necessity of solving the discrete matrix Riccati equation each time the estimates of Q and R are updated is eliminated.

2.9 Summary

This chapter provides the theoretical foundation for the rest of this study. It has reviewed the fundamental concepts of estimation theory and their application. It was shown that for Gaussian random sequences, the optimum estimator is linear. For a given system this important fact yields the filter theory of Wiener and Kalman, which are proved to be theoretically equivalent. However, the assumptions required to obtain an explicit solution using the Kalman method are much less restrictive. The Kalman filter equations are also in recursive form which provides a great computational advantage, even in the stationary case. The theory developed here is used extensively in the next chapter to derive the learning criterion for adaptive learning of the optimal Kalman filter.

CHAPTER 3

THE LEARNING CRITERION

3.1 Introduction and Organization of the Chapter

In Chapter 2 some of the important concepts of estimation theory were reviewed, and the results of Wiener and Kalman filter theory were presented and compared. There it was shown that for optimum filtering the estimator must satisfy the Wiener-Hopf equation. This equation is also the fundament of the learning criterion developed in Section 3.2. The learning criterion (Theorem 3-3) provides the theoretical foundation for learning the optimum Kalman filter matrix. Section 3.3 reduces the learning criterion to a Markov sequence model. This Markov model is used in Chapter 4 to motivate the adaptive technique.

3.2 The Unsupervised Learning Criterion

The purpose of the unsupervised learning criterion is to provide a necessary and sufficient condition for an adaptive solution to the stationary optimum filtering and prediction problem, where the signal and noise covariance matrices Q and R are unknown. In addition, this criterion must have two other characteristics.

1. It must be a function of measurable and/or calculable quantities only.
2. It must provide information from which convergent algorithms can be derived.

Otherwise, the criterion is meaningless from an engineering viewpoint.

As observed in Chapter 2 optimal filtering is a necessary condition for optimal smoothing. As a consequence, the unsupervised learning criterion is a necessary condition for adaptively solving the optimal smoothing problem.

The learning criterion developed in this chapter is said to be unsupervised because it operates on-line using only the noisy information generated by the actual filtering process. Specifically, it utilizes the residual time series $\{\delta_n\}$ which is the difference between the measurement process $\{z(n)\}$ and the predicted measurement process $\{H\hat{x}(n-1)\}$. By contrast a supervised learning criterion operates off-line using noise free externally supplied information. For example, in the context of this paper the learning process would be classified as supervised if, in addition to the noisy measurements $z(n)$, knowledge of the actual states $x(n)$ were available over a required training period.

In the search for an appropriate learning criterion the following elementary lemma is useful. It is included for completeness.

LEMMA 3-1:

Let the vector γ be orthogonal to the data set $\{z(1), \dots, z(n)\}$. Then γ is orthogonal to any linear combination of the data.

Since the only information available is the residual time series $\{\delta_n\}$, it is natural to determine what properties it has when the filter is optimum and suboptimum. The independence and zero mean characteristics of $v(n)$ and $w(n)$ are exploited freely in the following theorem proofs without comment.

THEOREM 3-2:

Given the dynamic system

$$\underline{\hat{x}}(n+1) = \Phi \underline{\hat{x}}(n) + \underline{w}(n) \quad (1.3)$$

the observation process

$$\underline{z}(n) = H \underline{\hat{x}}(n) + \underline{v}(n) \quad (1.4)$$

where the inverse of H exists, and the filtering equations

$$\begin{aligned} \underline{\hat{x}}(i) &= \underline{0}, \text{ where } i \text{ is the initial time} \\ \underline{\hat{x}}(n) &= \Phi \underline{\hat{x}}(n-1) + K[\underline{z}(n) - H\Phi \underline{\hat{x}}(n-1)] \end{aligned} \quad (2.19)$$

If $K = K_{\text{opt}}$

$$= \Sigma H^T [H \Sigma H^T + R]^{-1}$$

i.e., K_{opt} is the optimum filter matrix chosen such that Theorem 2-6 is satisfied, then

$$E\{[\underline{\Delta}_{j+1} - H\Phi K_{\text{opt}} \underline{\delta}_j]^T \underline{\delta}_j\} = 0 \quad \text{for all } j \leq n \quad (3.1)$$

where

$$\underline{\delta}_j = \underline{z}(j) - H\Phi \underline{\hat{x}}(j-1)$$

$$\underline{\Delta}_{j+1} = \underline{z}(j+1) - H\Phi \underline{\hat{x}}(j-1)$$

and conversely.

Proof: First examine the bracketed term in the above expectation

$$\begin{aligned}
 \Delta_{j+1} - H\Phi K_{\text{opt}} \delta_j &= z(j+1) - H\Phi \hat{x}(j-1) - H\Phi K_{\text{opt}} \delta_j \\
 &= z(j+1) - H\Phi \hat{x}(j-1) - H\Phi K_{\text{opt}} [z(j) - H\Phi \hat{x}(j-1)] \\
 &= z(j+1) - H\Phi \hat{x}(j) \\
 &= \delta_{j+1}
 \end{aligned}$$

Therefore, Eq. (3.1) may be rewritten as

$$E\{\delta_{j+1}^T \delta_j\} = 0 \quad \text{for all } j \leq n \quad (3.2)$$

Assume $K = K_{\text{opt}}$, which from Theorem (2-6) implies that

$$E\{[x(n) - \hat{x}(n)]^T B(n, j) z(j)\} = 0 \quad \text{for all } n, j$$

where $j \leq n$, and $B(n, j)$ is an arbitrary $m \times p$ matrix whose elements are real. Now expanding Eq. (3.2) results in

$$\begin{aligned}
 &E\{[z(j+1) - H\Phi \hat{x}(j)]^T [z(j) - H\Phi \hat{x}(j-1)]\} \\
 &= E\{[H\Phi x(j+1) + v(j+1) - H\Phi \hat{x}(j)]^T [z(j) - H\Phi \hat{x}(j-1)]\} \\
 &= E\{[x(j+1) - \hat{x}(j)]^T H^T [z(j) - H\Phi \hat{x}(j-1)]\} + \\
 &\quad E\{v^T(j+1) [z(j) - H\Phi \hat{x}(j-1)]\} \\
 &= E\{[\phi x(j) + w(j) - \phi \hat{x}(j)]^T H^T [z(j) - H\Phi \hat{x}(j-1)]\} + \\
 &\quad E\{v^T(j+1)\} E\{z(j) - H\Phi \hat{x}(j-1)\}
 \end{aligned}$$

$$\begin{aligned}
&= E \{ [\underline{x}(j) - \hat{\underline{x}}(j)]^T \Phi^T H^T [\underline{z}(j) - H\hat{\underline{x}}(j-1)] \} + \\
&\quad E \{ \underline{w}^T(j) H^T [\underline{z}(j) - H\hat{\underline{x}}(j-1)] \} + 0 \\
&= E \{ [\underline{x}(j) - \hat{\underline{x}}(j)]^T \Phi^T H^T \underline{z}(j) \} - \\
&\quad E \{ [\underline{x}(j) - \hat{\underline{x}}(j)]^T \Phi^T H^T H\hat{\underline{x}}(j-1) \} + 0
\end{aligned}$$

The first term in the last equation is zero from the projection theorem. The second term is zero since $\hat{\underline{x}}(j-1)$ is a linear combination of the $\underline{z}(k)$, $k \leq j-1$ and, therefore, $\Phi^T H^T H\hat{\underline{x}}(j-1) \in Y(n)$.

Consequently,

$$E\{\delta_{j+1}^T \delta_j\} = 0 \text{ for all } j \leq n.$$

Converse: Assume $E\{\delta_{j+1}^T \delta_j\} = 0$ for all $j \leq n$.

As a consequence of this assumption, the following preliminary result is obtained

$$\begin{aligned}
E\{\delta_{j+1}^T \delta_j\} &= E\{[\underline{z}(j+1) - H\hat{\underline{x}}(j)]^T [\underline{z}(j) - H\hat{\underline{x}}(j-1)]\} \\
&= E\{[H\underline{x}(j+1) + \underline{v}(j+1) - H\hat{\underline{x}}(j)]^T [\underline{z}(j) - H\hat{\underline{x}}(j-1)]\} \\
&= E\{[\underline{x}(j+1) - \hat{\underline{x}}(j)]^T H^T \delta_j\} = 0 \quad \text{for all } j \leq n \quad (3.3)
\end{aligned}$$

By hypothesis it is known that the lag-one residuals are orthogonal.

It must now be shown that

$$E\{\delta_{n+1}^T \delta_j\} = 0 \quad \text{for all } j \leq n$$

This is accomplished by iterating backward in time, starting with $j = n-1$. The state, observation and filtering equations are substituted freely in the following text.

$$\begin{aligned}
E\{\delta_{n+1}^T \delta_{n-1}\} &= E\{[z(n+1) - H\hat{x}(n)]^T [z(n-1) - H\hat{x}(n-1)]\} \\
&= E\{[H \underline{x}(n+1) + y(n+1) - H\hat{x}(n)]^T \delta_{n-1}\} \\
&= E\{[\underline{x}(n+1) - \hat{x}(n)]^T H^T \delta_{n-1}\} \\
&= E\{[\phi \hat{x}(n) + \underline{w}(n) - \phi \hat{x}(n-1) - K[z(n) - H\hat{x}(n-1)]]^T H^T \delta_{n-1}\} \\
&= E\{[\underline{x}(n) - \hat{x}(n-1)]^T \phi^T H^T \delta_{n-1}\} - \\
&\quad E\{[z(n) - H\hat{x}(n-1)]^T K^T \phi^T H^T \delta_{n-1}\}
\end{aligned}$$

The first term in the above equation is zero from Lemma 3-1 and Eq. (3.3).

The second term is zero from Lemma 3-1 and Eq. (3.2). Thus

$$E\{\delta_{n+1}^T \delta_{n-1}\} = 0$$

Continuing the iteration, it is concluded that

$$E\{\delta_{n+1}^T \delta_j\} = 0 \quad \text{for all } j \leq n \quad (3.4)$$

A more illustrated form of Eq. (3.4) is

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \Phi^T H^T \delta_{\underline{j}}\} = 0 \quad \text{for all } j \leq n \quad (3.5)$$

In particular for $j = i+1$, Eq. (3.4) becomes

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \Phi^T H^T [\underline{z}(i+1) - H\hat{\underline{x}}(i)]\} = 0$$

Since $\hat{\underline{x}}(i) = \underline{0}$,

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \Phi^T H^T \underline{z}(i+1)\} = 0$$

or

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T B(n, i+1) \underline{z}(i+1)\} = 0$$

Proceeding by induction, assume

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T B(n, k) \underline{z}(k)\} = 0 \quad \text{for all } k \leq n-1$$

Now for $k = n$, Eq. (3.5) gives

$$\begin{aligned} E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \Phi^T H^T \delta_{\underline{n}}\} &= 0 \\ &= E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \Phi^T H^T \underline{z}(n)\} \\ &\quad - E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \Phi^T H^T H\hat{\underline{x}}(n-1)\} \end{aligned}$$

Note that $\hat{\underline{x}}(n-1)$ is a linear combination of $\{\underline{z}(i), \underline{z}(i+1), \dots, \underline{z}(n-1)\}$.

Thus, from Lemma 3-1, the second term in the above equation is zero.

Consequently,

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T B(n,k) \underline{z}(k)\} = 0 \quad \text{for all } k \leq n$$

or equivalently

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)]^T \underline{y}(n)\} = 0 \quad \text{for all } \underline{y}(n) \in Y(n)$$

where $\underline{y}(n) = \sum_{k=i}^n B(n,k) \underline{z}(k)$ Q.E.D.

This theorem is important because it states that the filter is optimal if and only if the residual time series $\{\delta_{\underline{n}}\}$ is an orthogonal series. Since the residual process is observable, this theorem provides a method for determining whether or not the filter matrix is optimum. But if it is suboptimum, Theorem 3-2 does not indicate how it should be altered to approach K_{opt} . Consequently, it fails to meet the second characteristic required of a learning criterion. This deficiency motivates Theorem 3-3 which is used in Chapter 4 to develop algorithms that learn K_{opt} .

THEOREM 3-3: The Learning Criterion

Given the system of Theorem 3-2, a necessary and sufficient condition for the satisfaction of the Wiener-Hopf equation

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{z}^T(j)\} = [0] \quad \text{for all } n, j. \quad (2.6)$$

where $j \leq n$

is

$$E\{[\Delta_{\underline{j}+1} - H\Phi K_{\text{opt}} \delta_{\underline{j}}] \delta_{\underline{j}}^T\} = [0] \quad \text{for all } j \leq n \quad (3.6)$$

Proof: From the development of Eq. (3.2) from Eq. (3.1), Eq. (3.6) may be written

$$\begin{aligned} E\{[\Delta_{j+1} - H\Phi K_{\text{opt}}\delta_j] \delta_j^T\} &= E\{\delta_{j+1}\delta_j^T\} \\ &= [0] \quad \text{for all } j \leq n \end{aligned} \quad (3.6)$$

Now assume the Wiener-Hopf equation (2.6) holds, then

$$\begin{aligned} E\{\delta_{j+1}\delta_j^T\} &= E\{[z(j+1) - H\Phi\hat{x}(j)][z(j) - H\Phi\hat{x}(j-1)]^T\} \\ &= E\{[Hx(j+1) + v(j+1) - H\Phi\hat{x}(j)] \delta_n^T\} \\ &= E\{[H\Phi x(j) + w(j) - H\Phi\hat{x}(j)] \delta_j^T\} \\ &= H\Phi E\{[x(j) - \hat{x}(j)] z^T(j)\} \\ &\quad - H\Phi E\{[x(j) - \hat{x}(j)] \hat{x}^T(j-1)\} \Phi^T H^T \end{aligned}$$

The first term is zero directly from (2.6). The second term is also zero from (2.6) because $\hat{x}(j-1)$ is a linear combination of $z(k)$, $k \leq j-1$. Thus,

$$E\{\delta_{j+1}\delta_j^T\} = [0] \quad \text{for all } j \leq n$$

Converse: Assume (3.6) is true.

As a consequence of this assumption, the following preliminary result is obtained

$$\begin{aligned}
 E\{\delta_{\underline{z}(j+1)} \delta_{\underline{z}(j)}^T\} &= E\{[z(j+1) - H\Phi \hat{x}(j)][z(j) - H\Phi \hat{x}(j-1)]^T\} \\
 &= E\{[Hx(j+1) + v(j+1) - H\Phi \hat{x}(j)] \delta_{\underline{z}(j)}^T\} \\
 &= HE\{[x(j+1) - \Phi \hat{x}(j)] \delta_{\underline{z}(j)}^T\} \\
 &= H\Phi E\{[x(j) - \hat{x}(j)] \delta_{\underline{z}(j)}^T\} = [0] \quad \text{for all } j \leq n
 \end{aligned}$$

Since the inverse of H exists,

$$E\{[x(j) - \hat{x}(j)] \delta_{\underline{z}(j)}^T\} = [0] \quad \text{for all } j \leq n \quad (3.7)$$

It must now be shown that

$$E\{\delta_{\underline{z}(n+1)} \delta_{\underline{z}(j)}^T\} = [0] \quad \text{for all } j \leq n$$

This is accomplished by iterating backward in time, starting with $j = n-1$.

Proceeding as in Theorem 3-2,

$$\begin{aligned}
 E\{\delta_{\underline{z}(n+1)} \delta_{\underline{z}(n-1)}^T\} &= E\{[z(n+1) - H\Phi \hat{x}(n)] \delta_{\underline{z}(n-1)}^T\} \\
 &= E\{[Hx(n+1) + v(n+1) - H\Phi \hat{x}(n)] \delta_{\underline{z}(n-1)}^T\} \\
 &= HE\{[\Phi x(n) + w(n) - \Phi \hat{x}(n)] \delta_{\underline{z}(n-1)}^T\}
 \end{aligned}$$

$$\begin{aligned}
&= H\Phi E\{[\Phi \tilde{x}(n-1) + \tilde{w}(n-1) - \hat{\tilde{x}}(n)] \delta_{n-1}^T\} \\
&= H\Phi E\{[\Phi \tilde{x}(n-1) - \Phi \hat{\tilde{x}}(n-1) - K[\tilde{z}(n) - H\Phi \hat{\tilde{x}}(n-1)]] \delta_{n-1}^T\} \\
&= H\Phi E\{[\tilde{x}(n-1) - \hat{\tilde{x}}(n-1)] \delta_{n-1}^T\} - \\
&\quad H\Phi K E\{[\tilde{z}(n) - H\Phi \hat{\tilde{x}}(n-1)] \delta_{n-1}^T\}
\end{aligned}$$

The first term is zero from Eq. (3.7) and the second term is zero from Eq. (3.6). Therefore,

$$E\{\delta_{n+1} \delta_{n-1}^T\} = [0] \quad .$$

Using an identical argument and continuing the iteration, it can be shown that

$$E\{\delta_{n+1} \delta_j^T\} = [0] \quad \text{for all } j \leq n \quad (3.8)$$

A more useful form of the above equation is obtained by substituting the observation equation for $\tilde{z}(n+1)$ and the state equation for $\tilde{x}(n+1)$. It is

$$E\{[\tilde{x}(n) - \hat{\tilde{x}}(n)] \delta_j^T\} = [0] \quad \text{for all } j \leq n \quad (3.9)$$

Proceeding inductively, it can now be shown that the Wiener-Hopf equation follows from Eq. (3.9). In particular, for $j = i+1$, it becomes

$$E\{[\tilde{x}(n) - \hat{\tilde{x}}(n)] \delta_{i+1}^T\} = [0]$$

and since

$$\hat{\tilde{x}}(i) = 0$$

$$\begin{aligned}
& E\{[\underline{x}(n) - \hat{\underline{x}}(n)][\underline{z}(i+1) - H\Phi\hat{\underline{x}}(i)]^T\} \\
& = E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{z}^T(i+1)\} = [0]
\end{aligned}$$

Again proceeding by induction, assume that

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{\delta}_k^T\} = [0] \quad \text{for all } k \leq n-1$$

or equivalently

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{z}^T(k)\} = [0] \quad \text{for all } k \leq n-1$$

Now for $k = n$, Eq. (3.9) becomes

$$\begin{aligned}
& E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{\delta}_n^T\} = [0] \\
& = E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{z}^T(n)\} - E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \hat{\underline{x}}(n-1)\} \Phi^T H^T
\end{aligned}$$

Because $\hat{\underline{x}}(n-1)$ is just a linear combination of $\{\underline{z}(i), \dots, \underline{z}(n-1)\}$, the second term in the last equation is null. Therefore, it follows that

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)] \underline{z}^T(k)\} = [0] \quad \text{for all } k \leq n$$

Q.E.D.

This theorem states that the lack of orthogonality when K is not equal to the optimal filter K_{opt} is reflected in a non-null correlation matrix and conversely. That is, $K = K_{\text{opt}}$ if and only if

$$\begin{aligned}
E\{\underline{\delta}_{n+1} \underline{\delta}_n^T\} & \stackrel{\Delta}{=} \Omega(n+1, n) \\
& = [0]
\end{aligned} \tag{3.10}$$

$\Omega(n+1, n)$ can be used as a basis for learning K_{opt} , as shown in Chapter 4, by devising a scheme to utilize this correlation to adjust K_n^{-1} such that it converges to K_{opt} as n approaches infinity.

The following corollary is a direct consequence of Theorem 3-3.

Corollary 3-4

Equation (3.6), $E\{\delta_{j+1} \delta_j^T\} = [0]$ for all j implies $K = \Sigma H^T [H \Sigma H^T + R]^{-1}$ which is the optimum Kalman filter matrix.

Proof:

$$E\{\delta_{j+1} \delta_j^T\} = E\{[z(j+1) - H\hat{x}(j)][z(j) - H\hat{x}(j-1)]^T\}$$

$$= H E\{[x(j+1) - \hat{x}(j)] \delta_j^T\}$$

Substituting Eq. (2.19) for $\hat{x}(j)$ gives

$$= H\Phi E\{[x(j) - \hat{x}(j-1) - K[Hx(j) + v(j) - H\hat{x}(j-1)]] \delta_j^T\}$$

$$= H\Phi E\{[x(j) - \hat{x}(j-1)][x(j) - \hat{x}(j-1)]^T\} H^T$$

$$- H\Phi K E\{[Hx(j) + v(j) - H\hat{x}(j-1)][Hx(j) + v(j) - H\hat{x}(j-1)]^T\}$$

$$= H\Phi \Sigma H^T - H\Phi K H E\{[x(j) - \hat{x}(j-1)][x(j) - \hat{x}(j-1)]^T\} H^T$$

$$- H\Phi K E\{v(j) v^T(j)\}$$

1. Since in the adaptive process the filter matrix is being changed as n increases, its value at time n is denoted K_n .

$$= H\Phi\Sigma H^T - H\Phi K [H\Sigma H^T + R] = [0]$$

Solving for K yields

$$K = \Sigma H^T [H\Sigma H^T + R]^{-1}$$

which is the optimal Kalman filter.

Q.E.D.

3.3 A Linear Regression Model of the Learning Criterion

Since Theorem 3-3 indicates through the correlation between δ_{n+1} and δ_n whether or not the filter matrix is optimal, it satisfies the first requirement for a learning criterion. To show that it permits the development of a method for satisfying the second requirement and, therefore, qualifies as an unsupervised learning criterion, it is convenient to represent the correlation between the residuals as a linear regression

$$\delta_{n+1} = P_n \delta_n + e_n \quad (3.11)$$

where P_n is the regression matrix and e_n is the error term. Post-multiplying both sides of Eq. (3.11) by δ_n^T and taking the expected value gives

$$E\{\delta_{n+1} \delta_n^T\} = P_n E\{\delta_n \delta_n^T\} + E\{e_n \delta_n^T\}$$

or

$$\Omega(n+1,n) = P_n \Omega(n,n) + E\{e_{\sim n} \delta_n^T\} \quad (3.12)$$

Selecting the regression matrix to be

$$P_n = \Omega(n+1,n) \Omega^{-1}(n,n) \quad (3.13)$$

forces

$$E\{e_{\sim n} \delta_n^T\} = [0]$$

P_n then represents the correlation between $\delta_{\sim n+1}$ and $\delta_{\sim n}$. It is easy to see that when $P_n = [0]$ for all n , the learning criterion is satisfied and conversely. Thus, an equivalent statement of Theorem 3-3 can be made in terms of P_n .

THEOREM 3-4:

Given the system of Theorem 3-2, the filter matrix K is optimal if and only if

$$P_n = [0] \quad \text{for all } n.$$

Thus, an algorithm is required which utilizes P_n to adjust K_n in a way that forces P_n to converge to the null matrix as n approaches infinity. Equivalently, the adjustment must force $\delta_{\sim n+1}$ to approach $e_{\sim n}$ as n increases without bound. The method for achieving this goal is deferred to Chapter 4, where it is shown that P_n is proportional to the difference between the optimum filter matrix K_{opt} and the actual filter matrix K_n .

3.4 Summary

This chapter has developed a learning criterion which is a necessary and sufficient condition for optimality of the filter matrix. Furthermore, it is a function of measurable quantities only and it specifies what must be accomplished to learn K_{opt} . The learning process was reduced to a regression equation such that when the correlation matrix P_n is null, Theorem 3-3 is satisfied. In the next chapter, the regression equation is used to indicate how the adaptation should be performed. Stochastic algorithms are derived from a mean-square error performance index to estimate the optimum filter matrix. The theory of stochastic approximation is invoked to prove that the adaptation algorithms converge.

CHAPTER 4

DERIVATION OF THE ADAPTIVE ALGORITHMS

4.1 Introduction and Organization of the Chapter

This chapter is concerned with the derivation of stochastic algorithms which converge to the optimum Kalman filter matrix. This development is based on the unsupervised learning criterion of Theorem 3-3. Section 4.2 utilizes the regression model of the learning procedure to motivate the adaptation process. Section 4.3 introduces the theory of stochastic approximation which provides the theoretical foundation for this chapter. In Section 4.4 a performance index is introduced from which the adaptive learning algorithms are derived. Section 4.5 uses the theory of stochastic approximation and the learning criterion (Theorem 3-3) to prove probability one convergence of the algorithms. A method for accelerating the rate of convergence is given in Section 4.6. Section 4.7 summarizes the chapter.

4.2 Motivation for the Adaptive Process

At the end of Chapter 3 the learning criterion was reduced to a linear regression

$$\delta_{n+1} = \Delta_{n+1} - H\Phi K_{n \sim n} \delta_{n \sim n} = P_{n \sim n} \delta_{n \sim n} + e_n \quad (3.11)$$

such that the learning criterion is satisfied if and only if $P_n = [0]$.

From Theorem 3-4 this condition implies that the filter matrix is optimum and conversely.

Consider now the correlation P_n that exists when the filter matrix is suboptimum. The goal is to use this correlation matrix to force K_n to approach K_{opt} as n approaches infinity. As a consequence, P_n must converge to the null matrix or equivalently δ_{n+1} must approach e_n in the limit. The role of P_n in accomplishing the above goal is made evident by first rewriting Eq. (3.11) as

$$e_n = \delta_{n+1} - P_n \delta_n \quad (4.1)$$

and setting

$$P_n = H\Phi\Delta K_n \quad (4.2)$$

where ΔK_n at this point is an arbitrary matrix chosen to satisfy Eq. (4.2). The following convention is observed in the ensuing development: a suboptimal estimate obtained using the structure of the filtering Eq. (2.19), i.e., the filter used is not optimal, is denoted $\hat{x}_n$ to distinguish it from the optimal estimate denoted $\hat{x}(n)$. Assume that optimum estimation has occurred up to, but not including, time n . Now substituting for δ_{n+1} , δ_n and P_n in Eq. (4.1) gives

$$e_n = z(n+1) - H\Phi\hat{x}_n - H\Phi\Delta K_n [z(n) - H\Phi\hat{x}(n-1)]$$

By replacing $\hat{x}_n$ with the filtering equation, the above expression becomes

$$\begin{aligned} e_n = z(n+1) - H\Phi \{ \phi\hat{x}(n-1) + K_n [z(n) - H\Phi\hat{x}(n-1)] \} \\ - H\Phi\Delta K_n [z(n) - H\Phi\hat{x}(n-1)] \end{aligned}$$

$$= \underline{z}(n+1) - H\Phi \{ \Phi \hat{\underline{x}}(n-1) + (K_n + \Delta K_n)[\underline{z}(n) - H\Phi \hat{\underline{x}}(n-1)] \}$$

Observe that ΔK_n appears as a correction to the filtering matrix K_n employed at time n . If $K_n + \Delta K_n$ had been used instead, the correlation P_n would have been zero. This means that $K_n + \Delta K_n = K(n)$, thus the equation for $\underline{e}_n$ reduces to

$$\begin{aligned} \underline{e}_n &= \underline{z}(n+1) - H\Phi \{ \Phi \hat{\underline{x}}(n-1) + K(n)[\underline{z}(n) - H\Phi \hat{\underline{x}}(n-1)] \} \\ &= \underline{z}(n+1) - H\Phi \hat{\underline{x}}(n) \\ &= \delta_{n+1} \end{aligned} \tag{4.3}$$

Equation (4.3) implies that the filter matrix has been adjusted so that the new correlation matrix is null. Thus from Theorem 3-4, the new filter matrix $K(n) = K_n + \Delta K_n$ is the optimum Kalman filter. The value of ΔK_n as determined from (4.2) is

$$\Delta K_n = (H\Phi)^{-1} P_n \tag{4.4}$$

under the assumption that H inverse exists. However, since the distribution function necessary to determine P_n is unknown, ΔK_n cannot be calculated. But P_n and, therefore ΔK_n , can be estimated by using the theory of stochastic approximation. In the next section a brief introduction to the subject is presented.

4.3 Stochastic Approximation

Basically, stochastic approximation is an algorithmic search technique for optimization in the presence of uncertainty. This

uncertainty or noise may arise from an ignorance of the underlying phenomena, experimental errors or inherent random fluctuations. The nomenclature stochastic approximation emphasizes the stochastic nature of the errors in, say, the process measurements and the use of these measurements to approximate the location of the goal.

Such stochastic search algorithms are naturally more complex than their deterministic counterparts. However, the fundamental considerations are the same:

- (1) selecting a direction in which to search, then
- (2) selecting the distance to move (choosing the step size).

The effect of random error on an algorithm may cause it to converge to some non-optimum value or even to diverge. Therefore, correct convergence takes priority over speed of convergence. In stochastic approximation this consideration is reflected in the choice of step sizes. The direction to move is selected as if the process was deterministic. This is reminiscent of the Separation Theorem of control theory. The neglect of the stochastic nature of the problem means that some step directions may be incorrect, but such setbacks are swamped-out in the long run by additional data if the step sizes are properly selected. This is the crux of stochastic approximation and is directly related to the law of large numbers.

A powerful theorem due to Gladyshev [1965] serves to lay the mathematical groundwork for stochastic approximation and is used later in conjunction with the Learning Criterion (Theorem 3.3) to prove convergence of the algorithm for learning the optimum Kalman filter matrix.

THEOREM 4.1:

Let $\{y_n: n=1, 2, \dots\}$ be an observable random sequence which depends on a real parameter α such that

$$E\{y_n(\alpha)\} = f_n(\alpha) \quad (4.5)$$

and for the real constant C ,

$$f_n(\alpha_{\text{opt}}) = C \quad (4.6)$$

$$f_n(\alpha) < C \quad \text{for } \alpha < \alpha_{\text{opt}}$$

$$f_n(\alpha) > C \quad \text{for } \alpha > \alpha_{\text{opt}}$$

The value of α which satisfies the regression equation (4.6) is estimated by observing y_n and using the recursive relation

$$\alpha_{n+1} = \alpha_n - a_n [y_n(\alpha_n) - C] \quad (4.7)$$

If $\{a_n\}$ is a sequence of positive numbers such that

$$(i) \quad \sum_{n=1}^{\infty} a_n = \infty$$

$$(ii) \quad \sum_{n=1}^{\infty} a_n^2 < \infty$$

and

$$(iii) \quad \text{g.l.b.}_{|\alpha_n - \alpha_{\text{opt}}| < \varepsilon} (\alpha_n - \alpha_{\text{opt}}) [f_n(\alpha_n) - C] \leq 0 \quad \text{for all } \varepsilon > 0$$

(iv) There exists a positive number r such that for all α

$$E\{y_n^2(\alpha)\} < r(1 + \alpha^2)$$

then the sequence $\{\alpha_n\}$ defined by (4.7) converges with probability one to α_{opt} , i.e., $P\{\lim_{n \rightarrow \infty} \alpha_n = \alpha_{\text{opt}}\} = 1$.

If $C = 0$, then the stochastic approximation sequence (4.7) becomes an algorithm for finding the unique zero of $f_n(\alpha)$. Historically, this application was the motivation for the original work performed by the statisticians Robbins and Monro [1951]. Dvoretzky [1956] extended and unified their work and that of other investigators. However, the viewpoint is strictly that of mathematical statistics. A comprehensive survey of stochastic approximation and its engineering applications is contained in Hampton [1967]. For an in-depth study of the subject, the reader is referred to the interesting monographs by Albert and Gardner [1967] and Fu [1968].

4.4 Stochastic Algorithms for Learning the Optimum Filter Matrix

To facilitate the application of stochastic approximation, the problem of estimating P_n is rewritten in the structure of a performance index. Let $L(P_n)$ be the expected value of the convex performance index to be minimized,

$$L(P_n) = E\{\ell(\delta_{n+1} - \hat{\delta}_{n+1})\} \quad (4.8)$$

where $\ell(\cdot)$ is the loss function and $\hat{\delta}_{n+1} = P_n \delta_n$. When $L(P_n)$ is known (the deterministic case), Eq. (4.8) can be minimized by solving Eq. (4.9) iteratively.

$$\nabla_p L(P_{\text{opt}}) = E\{\nabla_p \ell(\delta_{n+1} - \hat{\delta}_{n+1})\} = [0] \quad (4.9)$$

Using standard computational techniques such as the gradient method, the recursive relation for finding P_{opt} is

$$\hat{P}_{n+1} = \hat{P}_n + \nabla_p L(\hat{P}_n) W_{n+1} \quad (4.10)$$

where $\{W_{n+1}\}$ is a weighting sequence which determines the magnitude of the correction term. Under appropriate convergence conditions, $\lim_{n \rightarrow \infty} \hat{P}_n = P_{\text{opt}}$. A block diagram, Fig. 4.1, of the process generated by the deterministic algorithm (4.10) reveals it to be a nonlinear feedback system. It has no input which implies that all the information needed to determine P_{opt} is included in the system a priori.

In an adaptive stochastic setting the probability distribution function required to perform the expectation indicated by Eq. (4.8) is not available. Thus, the performance index $L(P_n)$ and its gradient are unknown. This condition is precisely the motivation for stochastic approximation. What is needed is a method to utilize the available information, specifically the gradient of the loss function, to solve Eq. (4.9). By identifying $\nabla_p \ell(\delta_{n+1} - \hat{\delta}_{n+1})$ with the random sequence $y_n(\alpha)$ of Eq. (4.5) and using the stochastic approximation algorithm (4.7) with $C = 0$, the following recursive relation is obtained

$$\hat{P}_{n+1} = \hat{P}_n - a_{n+1} \nabla_p \ell(\delta_{n+1} - \hat{P}_{n-n} \delta_{n-n}) W_{n+1} \quad (4.11)$$

where $\{a_n\}$ satisfies conditions (i) and (ii) of Theorem 4.1 and $\{W_{n+1}\}$

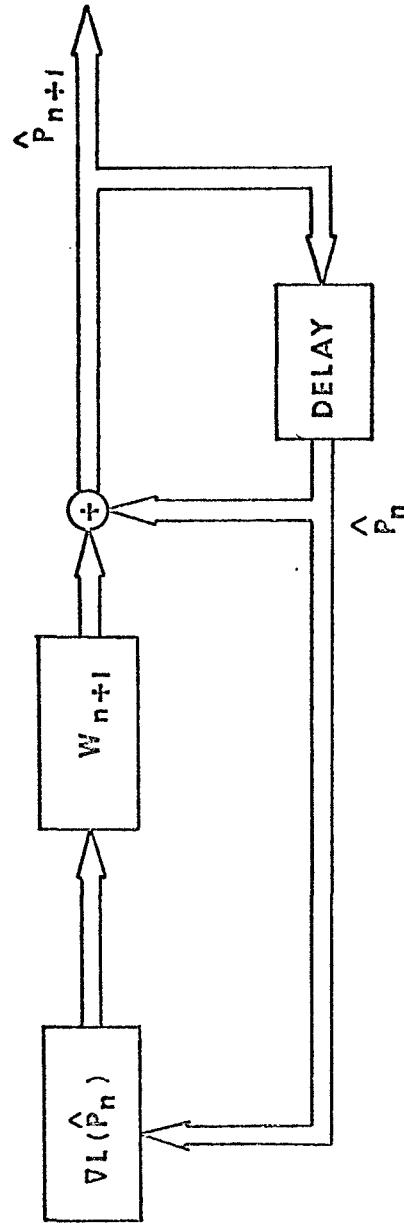


Fig. 4.1 Block Diagram of the Feedback System for the Deterministic Algorithm (4.10)

is an appropriate set of weighting terms. If W_{n+1} is chosen to be the identity matrix for all n , then Eq. (4.11) reduces to a standard stochastic approximation recursion. In general, W_{n+1} is a matrix chosen to accelerate the rate of convergence. By defining

$$G_n \triangleq a_n W_n \quad (4.12)$$

Eq. (4.11) becomes

$$\hat{P}_{n+1} = \hat{P}_n - \nabla_p \ell(\delta_{n+1} - \hat{P}_n \delta_n) G_{n+1} \quad (4.13)$$

This gradient descent algorithm is the direct stochastic analogue of the deterministic algorithm (4.10). Since Eq. (4.13) is a random sequence, its convergence must be established in a probabilistic sense. A block diagram realization of this process is illustrated in Fig. 4.2. It represents a nonlinear feedback system similar to that of Fig. 4.1, but it is forced. This important difference implies that a continuous input of external information is necessary to solve Eq. (4.9) iteratively in an unknown random environment. The concept of viewing the algorithm as a feedback system means that algorithmic convergence may be interpreted in terms of system stability.

To proceed with the development of Eq. (4.11), a differentiable and convex loss function must be specified. A quadratic loss function

$$\ell(\delta_{n+1} - \hat{\delta}_{n+1}) = \frac{1}{2} (\delta_{n+1} - P_n \delta_n)^T (\delta_{n+1} - P_n \delta_n) \quad (4.14)$$

is selected because it permits an optimum selection of the weighting matrices W_{n+1} and a systematic proof that the associated algorithm

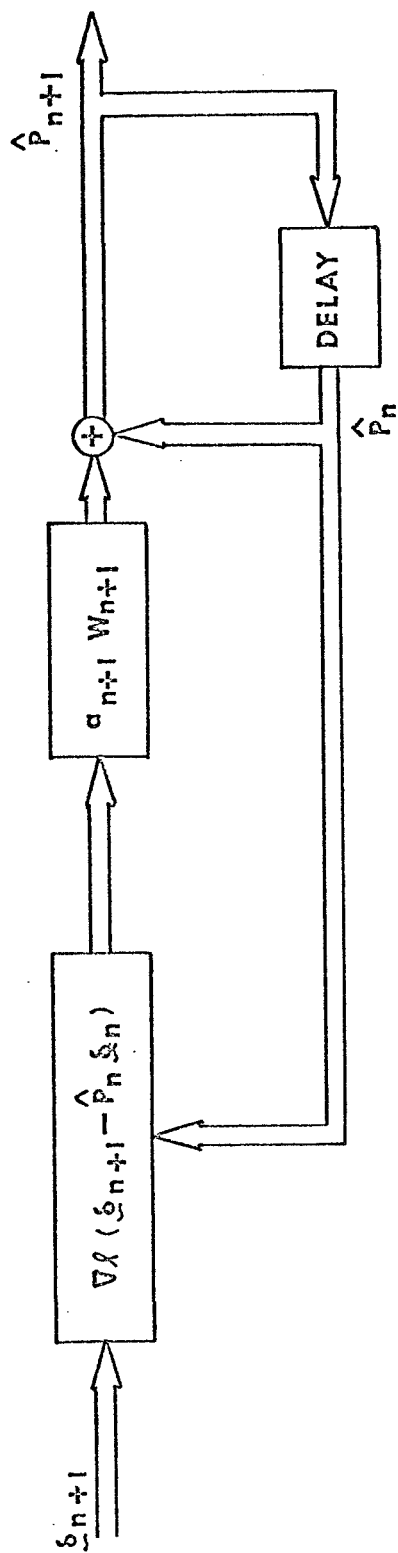


Fig. 4.2 Block Diagram of the Feedback System for the Adaptive Stochastic Algorithm (4.13)

converges to the Kalman filter matrix. This choice of loss function results in the familiar mean square performance index. Substituting the gradient of $\ell(\cdot)$ with respect to P_n into Eq. (4.11) gives

$$\hat{P}_{n+1} = \hat{P}_n + a_{n+1} [(\delta_{n+1} - \hat{P}_n \delta_n) \delta_n^T] W_{n+1} \quad (4.15)$$

It is readily verified that conditions (i) and (ii) of Theorem 4.1 are satisfied by any sequence of the form $(n)^{-b}$, where $1/2 < b \leq 1$. Therefore, letting $a_n = n^{-1}$, the harmonic series, gives the fastest rate of asymptotic convergence. Using this sequence, Eq. (4.15) becomes

$$\begin{aligned} \hat{P}_{n+1} &= \hat{P}_n + [(\delta_{n+1} - \hat{P}_n \delta_n) \delta_n^T] \frac{W_{n+1}}{n+1} \\ &= \hat{P}_n + \{[z(n+1) - H\phi \hat{x}_n] - \hat{P}_n [z(n) - H\phi \hat{x}_{n-1}]\} \delta_n^T \frac{W_{n+1}}{n+1} \end{aligned}$$

Substituting $\hat{P}_n = H\phi \Delta \hat{K}_n$ from (4.2) into the above equation, results in

$$\begin{aligned} \hat{P}_{n+1} &= \hat{P}_n + \{z(n+1) - \\ &\quad H\phi [\hat{x}_n + \Delta \hat{K}_n (z(n) - H\phi \hat{x}_{n-1})]\} \delta_n^T \frac{W_{n+1}}{n+1} \end{aligned}$$

Using the filtering equation for $\hat{x}_n$, it is seen that

$$\begin{aligned} \hat{P}_{n+1} &= \hat{P}_n + \{z(n+1) - \\ &\quad H\phi [\phi \hat{x}_{n-1} + (\hat{K}_n + \Delta \hat{K}_n)(z(n) - H\phi \hat{x}_{n-1})]\} \delta_n^T \frac{W_{n+1}}{n+1} \end{aligned}$$

The expression in brackets is just an updated estimate of $\underline{x}(n)$, therefore,

$$\begin{aligned}\hat{\underline{P}}_{n+1} &= \hat{\underline{P}}_n + \{ \underline{z}(n+1) - H\Phi \hat{\underline{x}}_n \} \delta_n^T \frac{W_{n+1}}{n+1} \\ &= \hat{\underline{P}}_n + \delta_{n+1} \delta_n^T \frac{W_{n+1}}{n+1}\end{aligned}\tag{4.16}$$

Assuming the inverse of H exists,

$$\Delta \hat{\underline{K}}_{n+1} = (H\Phi)^{-1} \hat{\underline{P}}_{n+1}\tag{4.17}$$

Instead of estimating $\underline{P}_{opt}$ and then calculating $\Delta \hat{\underline{K}}_{n+1}$ as required when using the stochastic algorithm (4.16), it is preferable to estimate $\underline{K}_{opt}$ directly. This is accomplished by re-defining the loss function $\ell(\cdot)$. Let

$$\Delta_{n+1} = (H\Phi)^{-1} [\underline{z}(n+1) - H\Phi \hat{\underline{x}}_n]$$

with $\delta_n = \underline{z}(n) - H\Phi \hat{\underline{x}}_n$ as before, and

$$\ell(\Delta_{n+1} - \hat{\Delta}_{n+1}) = \frac{1}{2} (\Delta_{n+1} - \underline{K}_{n+1} \delta_n)^T (\Delta_{n+1} - \underline{K}_{n+1} \delta_n)\tag{4.18}$$

Then, the iterative solution to

$$\begin{aligned}\nabla_{\underline{K}} L(\underline{K}_{opt}) &= \nabla_{\underline{K}} E\{\ell(\Delta_{n+1} - \underline{K}_{opt} \delta_n)\} \\ &= E\{\nabla_{\underline{K}} \ell(\Delta_{n+1} - \underline{K}_{opt} \delta_n)\} \\ &= [0]\end{aligned}\tag{4.19}$$

is given by the stochastic approximation algorithm

$$\hat{K}_{n+1} = \hat{K}_n + a_{n+1} \nabla_K \ell(\Delta_{n+1} - \hat{K}_n \delta_n) W_{n+1} \quad (4.20)$$

$$= \hat{K}_n + a_{n+1} \{ [\Delta_{n+1} - \hat{K}_n \delta_n] \} \delta_n^T W_{n+1}$$

$$= \hat{K}_n + a_{n+1} (H\Phi)^{-1} [z(n+1) - H \Phi \hat{x}(n-1)] -$$

$$\hat{K}_n [z(n) - H \Phi \hat{x}(n-1)] \} \delta_n^T W_{n+1}$$

$$= \hat{K}_n + a_{n+1} (H\Phi)^{-1} \delta_{n+1} \delta_n^T W_{n+1} \quad (4.21)$$

where $\{a_n = \frac{1}{n}\}$.

The recursive relation of Eq. (4.21) is the desired stochastic algorithm for learning the Kalman filter matrix K_{opt} . In the following section it is shown that Eq. (4.20) converges with probability one.

4.5 Convergence of the Stochastic Learning Algorithms

Proving that Eq. (4.20) converges to the solution of Eq. (4.19) only indicates that the performance index

$$\frac{1}{2} E\{(\Delta_{n+1} - K_n \delta_n)^T (\Delta_{n+1} - K_n \delta_n)\}$$

is minimized. It does not demonstrate that this minimum is also the Kalman filter. This is the first task at hand.

THEOREM 4-2:

The value of the filter that minimizes

$$L(K_n) = \frac{1}{2} E\{(\Delta_{n+1} - K_n \delta_n)^T (\Delta_{n+1} - K_n \delta_n)\} \quad (4.22)$$

coincides with the optimum Kalman filter K_{opt} .

Proof: Consider Eq. (4.19)

$$E\{\nabla_{K_n} L(\Delta_{n+1} - K_n \delta_n)\} = - E\{(\Delta_{n+1} - K_n \delta_n) \delta_n^T\}$$

Using the definition of Δ_{n+1} and δ_n in the above relation gives

$$\begin{aligned} E\{\nabla_{K_n} L(\Delta_{n+1} - K_n \delta_n)\} \\ &= - E\{(H\Phi)^{-1} \{[z(n+1) - H\Phi \hat{x}(n-1)] - K_n [z(n) - H\Phi \hat{x}(n-1)]\} \delta_n^T\} \\ &= - (H\Phi)^{-1} E\{\delta_{n+1} \delta_n^T\} \end{aligned} \quad (4.23)$$

Now taking the limit of Eq. (4.19) as K_n approaches K_{opt} and invoking the results of Theorem 3-3 (the Learning Criterion) yields

$$\begin{aligned} \lim_{K_n \rightarrow K_{opt}} E\{\nabla_{K_n} L(\Delta_{n+1} - K_n \delta_n)\} &= - (H\Phi)^{-1} \lim_{K_n \rightarrow K_{opt}} E\{\delta_{n+1} \delta_n^T\} \\ &= [0] \end{aligned} \quad (4.24)$$

Since the solution to Eq. (4.24) minimizes the performance index $L(K_n)$, the proof of the theorem is complete.

Q.E.D.

Now since $L(K_n)$ is a quadratic function of K_n , proving that algorithm (4.20) converges with probability one to the minimum of $L(K_n)$ implies that it converges to K_{opt} with probability one. The scalar case is considered first.

THEOREM 4-3:

Given the performance index

$$L(K_n) = E\{\ell(\Delta_{n+1} - K_n \delta_n)\}$$

where

$$\ell(\Delta_{n+1} - \hat{\Delta}_{n+1}) = \frac{1}{2} (\Delta_{n+1} - K_n \delta_n) (\Delta_{n+1} - K_n \delta_n) \quad (4.18)$$

and the stochastic algorithm

$$\hat{K}_{n+1} = \hat{K}_n - a_{n+1} \nabla_K \ell(\Delta_{n+1} - \hat{K}_n \delta_n) W_{n+1} \quad (4.20)$$

$$= \hat{K}_n + a_{n+1} (H\Phi)^{-1} \delta_{n+1} \delta_n W_{n+1} \quad (4.21)$$

where $\{a_n\}$ satisfies conditions (i) and (ii) and $\{W_{n+1}\}$ is a uniformly bounded sequence, then Eq. (4.21) converges with probability one to the minimum of the performance index which from Theorem 4-2 is equal to the optimum Kalman filter.

Proof:

By comparing Eqns. (4.20) and (4.22) with Eqns. (4.5) and (4.7), one observes that the gradient of the loss function corresponds to the random sequence $y_n(\alpha_n)$ and the expected value of the gradient of the loss function corresponds to $f_n(\alpha_n)$. Therefore, conditions (i) - (iv) of Theorem 4-1 become

$$(i) \quad \sum_{n=1}^{\infty} a_n W_n = \infty$$

$$(ii) \quad \sum_{n=1}^{\infty} (a_n W_n)^2 < \infty$$

$$(iii) \quad \text{g.l.b. } (K_n - K_{\text{opt}}) E\{\nabla_K \ell(\Delta_{n+1} - K_n \delta_n)\} \geq 0 \text{ for all } \varepsilon > 0 \\ |K_n - K_{\text{opt}}| < \varepsilon$$

$$(iv) \quad E\{[\nabla_K \ell(\Delta_{n+1} - K_n \delta_n)]^2\} \leq r(1 + K_n)^2$$

Since W_n is uniformly bounded for all n , conditions (i) and (ii) are immediately satisfied. Condition (iv) simply requires that the loss function not increase more rapidly than a quadratic in K_n . The loss function of Eq. (4.18) meets this restriction.

The only condition that remains to be verified is (iii). It is now considered.

The left hand side of inequality (iii) may be written

$$\begin{aligned} & E\{(K_n - K_{\text{opt}}) [\nabla_K \ell(\Delta_{n+1} - K_n \delta_n)]\} \\ &= - E\{(K_n - K_{\text{opt}}) \delta_n [\ell'(\Delta_{n+1} - K_n \delta_n)]\}^1 \\ &= - E\{(K_n - K_{\text{opt}}) \delta_n [\ell'(\Delta_{n+1} - (K_n - K_{\text{opt}}) \delta_n - K_{\text{opt}} \delta_n)]\} \\ &= - E\{\Lambda_n \ell'(\psi_{n+1} - \Lambda_n)\} \geq 0 \end{aligned}$$

1. The notation $\ell'(\cdot)$ denotes the derivative of ℓ with respect to its argument.

or equivalently

$$E\{\Lambda_n \ell'(\psi_{n+1} - \Lambda_n)\} \leq 0$$

where

$$\Lambda_n \triangleq (K_n - K_{\text{opt}}) \delta_n$$

and

$$\psi_{n+1} \triangleq \Delta_{n+1} - K_{\text{opt}} \delta_n$$

For the optimum filter matrix, Eq. (4.23) gives

$$\begin{aligned} E\{\nabla_K \ell(\Delta_{n+1} - K_{\text{opt}} \delta_n)\} &= - E\{\ell'(\Delta_{n+1} - K_{\text{opt}} \delta_n) \delta_n\} \\ &= - E\{\delta_n \ell'(\psi_{n+1})\} = 0 \end{aligned} \quad (4.25)$$

Therefore,

$$\begin{aligned} (K_n - K_{\text{opt}}) E\{\delta_n \ell'(\psi_{n+1})\} &= E\{(K_n - K_{\text{opt}}) \delta_n \ell'(\psi_{n+1})\} \\ &= E\{\Lambda_n \ell'(\psi_{n+1})\} \\ &= 0 \end{aligned} \quad (4.26)$$

By virtue of the convexity of $\ell(\cdot)$, the following property holds

$$\ell'(\psi_m) - \ell'(\psi_p) \leq 0 \text{ for all } \psi_m \leq \psi_p$$

Consequently,

$$\ell'(\psi_{n+1} - \Lambda_n) - \ell'(\psi_{n+1}) \leq 0 \text{ for all } \Lambda_n > 0$$

Multiplying both sides of the above inequality by Λ_n gives

$$\Lambda_n \ell'(\psi_{n+1} - \Lambda_n) - \Lambda_n \ell'(\psi_{n+1}) \leq 0 \quad \text{for all } \Lambda_n > 0 \quad (4.27)$$

Taking the expectation of (4.27) yields

$$E\{\Lambda_n \ell'(\psi_{n+1} - \Lambda_n)\} - E\{\Lambda_n \ell'(\psi_{n+1})\} \leq 0$$

The second term in this inequality is zero from Eq. (4.26). Thus,

$$E\{\Lambda_n \ell'(\psi_{n+1} - \Lambda_n)\} \leq 0$$

which confirms that condition (iii) is satisfied.

Q.E.D.

In general, algorithm (4.21) is a random matrix sequence. To use Theorem 4-1 in this case, the argument of Theorem 4-3 must be applied to each element of the matrix $E\{\nabla_K \ell(\Delta_{n+1} - K_{n \sim n} \delta)\}$ separately.

4.6 Acceleration of Convergence

In Section (4.4) it was pointed out that if W_{n+1} is a constant, algorithms Eqns. (4.11) and (4.20) are basically stochastic gradient schemes. Letting W_{n+1} equal the identity matrix and focusing attention on Eq. (4.20) results in

$$\begin{aligned} \hat{K}_{n+1} &= \hat{K}_n - \nabla_K \ell(\Delta_{n+1} - \hat{K}_{n \sim n} \delta) \frac{I}{n+1} \\ &= \hat{K}_n + \frac{1}{n+1} [\Delta_{n+1} - \hat{K}_{n \sim n} \delta] \delta_{n \sim n}^T \end{aligned} \quad (4.28)$$

For this choice of W_{n+1} the correction term in Eq. (4.28) simply proceeds in the direction of the local gradient of the loss function, with its magnitude weighted by the factor $\frac{1}{n+1}$. This indicates that even though large corrections are possible at the beginning of the estimation process when $\hat{K}_n$ is far from K_{opt} , no information from past estimation stages is utilized. Note that the $\frac{1}{n+1}$ factor forces the correction term to approach zero as the estimation progresses, since it becomes more probable that the observed gradient is due to noise instead of estimation error.

To accelerate convergence a second order stochastic algorithm is developed. This is accomplished by applying Newton's method to the loss function. For the scalar case this procedure gives

$$\hat{K}_{n+1} = \hat{K}_n - \frac{1}{n+1} \nabla_K \ell(\Delta_{n+1} - \hat{K}_n \delta_n) \{ \nabla_K [\nabla_K \ell(\Delta_{n+1} - \hat{K}_n \delta_n)] \}^{-1} \quad (4.29)$$

$$= \hat{K}_n + \frac{1}{n+1} (\Delta_{n+1} - \hat{K}_n \delta_n) \delta_n \{ \delta_n \delta_n \}^{-1} \quad (4.30)$$

$$= \hat{K}_n + \frac{1}{n+1} (\Delta_{n+1} - \hat{K}_n \delta_n) \delta_n W_{n+1}$$

where

$$W_{n+1} = \{ \delta_n \delta_n \}^{-1}$$

Even though Eq. (4.30) should converge more rapidly than Eq. (4.28), it does not utilize the information from past stages to select the step-size. One solution to this problem is to define an intermediate loss function

$$\ell_n(K_n) = \frac{1}{2n} \sum_{j=0}^{n-1} (\Delta_{j+1} - K_n \delta_j)^T (\Delta_{j+1} - K_n \delta_j) \quad (4.31)$$

The filter matrix $\hat{K}_n$ which minimizes Eq. (4.31) is a least squares solution to

$$\nabla_{K_n} \ell_n(\hat{K}_n) = \frac{1}{n} \sum_{j=0}^{n-1} (\Delta_{j+1} - \hat{K}_n \delta_j) \delta_j^T = [0] \quad (4.32)$$

Solving for $\hat{K}_n$ yields

$$\hat{K}_n = \sum_{j=0}^{n-1} \Delta_{j+1} \delta_j^T \left[\sum_{j=0}^{n-1} \delta_j \delta_j^T \right]^{-1} \quad (4.33)$$

and

$$\hat{K}_{n+1} = \sum_{j=0}^n \Delta_{j+1} \delta_j^T \left[\sum_{j=0}^n \delta_j \delta_j^T \right]^{-1} \quad (4.34)$$

It is now necessary to put Eq. (4.34) in recursive form, which also eliminates the matrix inversion. Using Eq. (4.33), Eq. (4.34) becomes

$$\hat{K}_{n+1} = \left[\hat{K}_n \sum_{j=0}^{n-1} \delta_j \delta_j^T + \Delta_{n+1} \delta_n^T \right] \left[\sum_{j=0}^{n-1} \delta_j \delta_j^T + \delta_n \delta_n^T \right]^{-1} \quad (4.35)$$

By defining W_n^{-1} as

$$W_n^{-1} \triangleq \frac{1}{n} \sum_{j=0}^{n-1} \delta_j \delta_j^T \quad (4.36)$$

then W_{n+1}^{-1} is determined by

$$\begin{aligned}
W_{n+1}^{-1} &= (1 - \frac{1}{n+1}) W_n^{-1} + \frac{1}{n+1} \delta_n \delta_n^T \\
&= \frac{n}{n+1} W_n^{-1} + \frac{1}{n+1} \delta_n \delta_n^T \\
&= \frac{1}{n+1} (nW_n^{-1} + \delta_n \delta_n^T)
\end{aligned} \tag{4.37}$$

or equivalently

$$W_{n+1} = (n+1)(nW_n^{-1} + \delta_n \delta_n^T)^{-1} \tag{4.38}$$

Using the well-known "inside-out" lemma (see e.g., Deutsch, 1965, p. 119) of numerical analysis, Eq. (4.38) may be rewritten

$$W_{n+1} = \frac{n+1}{n} \{W_n - W_n \delta_n [n + \delta_n^T W_n \delta_n]^{-1} \delta_n^T W_n\} \tag{4.39}$$

Now substitute Eq. (4.36) into Eq. (4.35). The result is

$$\hat{K}_{n+1} = [\hat{K}_n nW_n^{-1} + \Delta_{n+1} \delta_n^T] (nW_n^{-1} + \delta_n \delta_n^T)^{-1}$$

which can be rewritten using the "inside-out" lemma as

$$\hat{K}_{n+1} = [\hat{K}_n nW_n^{-1} + \Delta_{n+1} \delta_n^T] W_n - \frac{\begin{bmatrix} W_n \delta_n \delta_n^T W_n \\ (n + \delta_n^T W_n \delta_n) \end{bmatrix}}{\frac{1}{n}} \tag{4.40}$$

Performing the indicated multiplication and observing that $\delta_n^T W_n \delta_n$ is a scalar reduces Eq. (4.40) to

$$\begin{aligned}
\hat{K}_{n+1} &= \hat{K}_n + \frac{\Delta_{n+1} \delta_n^T W_n - \hat{K}_n \delta_n \delta_n^T W_n}{(n + \delta_n^T W_n \delta_n)} \\
&= \hat{K}_n + [\Delta_{n+1} - \hat{K}_n \delta_n] \frac{\delta_n^T W_n}{(n + \delta_n^T W_n \delta_n)} \quad (4.41)
\end{aligned}$$

Equation (4.41) is recursive and the matrix inversion has been avoided. To put this algorithm in the form of Eq. (4.20) it is necessary to prove that

$$\frac{\delta_n^T W_{n+1}}{n+1} = \frac{\delta_n^T W_n}{(n + \delta_n^T W_n \delta_n)}$$

Pre-multiplying Eq. (4.37) by W_n and post-multiplying by W_{n+1} gives

$$W_n = \frac{1}{n+1} (nI + W_n \delta_n \delta_n^T) W_{n+1}$$

Again pre-multiplying by δ_n^T

$$\begin{aligned}
\delta_n^T W_n &= \frac{1}{n+1} (n \delta_n^T + \delta_n^T W_n \delta_n \delta_n^T) W_{n+1} \\
&= \frac{1}{n+1} (n + \delta_n^T W_n \delta_n) \delta_n^T W_{n+1}
\end{aligned}$$

or

$$\frac{\delta_n^T W_{n+1}}{n+1} = \frac{\delta_n^T W_n}{(n + \delta_n^T W_n \delta_n)} \quad (4.42)$$

Substituting this equation into the estimation scheme of Eq. (4.41) gives the desired stochastic relation

$$\hat{K}_{n+1} = \hat{K}_n + [\Delta_{n+1} - \hat{K}_n \delta_n] \delta_n^T \frac{W_{n+1}}{n+1} \quad (4.43)$$

where

$$W_{n+1} = \frac{n+1}{n} [W_n - W_n \delta_n (n + \delta_n^T W_n \delta_n)^{-1} \delta_n^T W_n] \quad (4.39)$$

This stochastic approximation algorithm for learning the Kalman filter matrix is an optimal one in the sense that

- (a) it utilizes the information from all past estimation stages rather than one stage,
- (b) it minimizes the intermediate loss function $\ell_n(K_n)$ at each stage of iteration instead of just following the local gradient at each stage, and
- (c) it does not affect the asymptotic convergence of $\hat{K}_n$ to K_{opt} .

In addition, by invoking the "strong law of large numbers" it can be shown that $\ell_n(K)$ converges to $L(K)$ with probability one [Ho and Blaydon, 1966], where $L(K)$ is given by Eq. (4.22).

The increase in statistical efficiency attained when using the weight matrix defined in Eq. (4.39) is paid for in computational complexity. If W_{n+1} is selected to be a constant for all n , only one iteration is required per estimation stage, as opposed to two when W_{n+1} is specified by Eq. (4.39). It is worth noting that the asymptotic behavior of both algorithms is governed by the "law of large numbers",

i.e., the convergence rate for large n is proportional to $1/n$. However, from an engineering viewpoint the adaptive properties of the algorithm for small and intermediate values of n is crucial. In this region the accelerated version is far superior to the non-accelerated one, as is demonstrated experimentally in Chapter 6.

4.7 Summary

In this Chapter adaptive stochastic algorithms were derived for learning the optimal Kalman filter. The theory of stochastic approximation was briefly reviewed. It was employed in conjunction with the Learning Criterion of Chapter 3 to prove probability one convergence of the algorithms. A scheme for accelerating algorithmic convergence was presented. The final results of this acceleration technique are given by Eqns. (4.39) and (4.43). These are the two algorithms used to obtain the principal simulation results of Chapter 6.

However, before the experimental data are presented, the results obtained in this chapter are used in Chapter 5 to estimate other useful quantities such as the error covariance Γ , the predicted error covariance Σ , as well as Q and R . These estimates are needed to adaptively solve the optimal smoothing problem.

CHAPTER 5

ADDITIONS AND EXTENSIONS

5.1 Introduction and Organization of the Chapter

In Chapter 4 adaptive algorithms are derived for solving the optimal filtering. Adaptive prediction is included in this solution as a necessary condition. For optimal smoothing, however, knowledge of the optimum filter must be augmented by additional information. Specifically, the fixed-interval smoothing also requires the one-step error covariance matrix Σ (see Eq. (2.31)); fixed point smoothing requires R^{-1} as well as K_{opt} and Σ (see Eq. (2.33)); fixed-lag smoothing requires K_{opt} , Σ , and Q (see Eqns. (2.35) and (2.36)). In addition, if the error covariance matrix Γ is known, the formulas for optimal smoothing may be significantly simplified. Thus, Section 5.2 uses the results of Chapter 4 to estimate Σ , Γ , R , and Q . Techniques for lifting the rather restrictive assumption that the inverse of H exists are introduced in Section 5.3. Section 5.4 suggests simple modification to the algorithm of Chapter 4 to accommodate slowly time-varying signal and noise statistics. These same algorithms are heuristically altered in Section 5.5 a certain class of nonlinear systems. A summary of the chapter is contained in Section 5.6.

5.2 Estimation of Σ , Γ , R and Q

After the optimal filter has been learned, estimation schemes for Σ , Γ , R , and Q can be derived. To accomplish this, additional use

must be made of the information contained in the residual time series $\{\delta_{\sim n}\}$. The algorithm for estimating Σ is considered first. By performing the expectation indicated in Eq. (3.10), as in Corollary 3-4, it is seen that¹

$$\Omega(1) = E\{\delta_{\sim n+1} \delta_{\sim n}^T\} \quad (3.10)$$

$$= H\Phi\{\Sigma H^T - K[H\Sigma H^T + R]\} \quad (5.1)$$

Similarly, $\Omega(0)$ can be shown to be

$$\begin{aligned} \Omega(0) &= E\{\delta_{\sim n} \delta_{\sim n}^T\} \\ &= H\Sigma H^T + R \end{aligned} \quad (5.2)$$

Solving Eq. (5.1) for Σ produces

$$\Sigma = (H\Phi)^{-1} \Omega(1) (H^T)^{-1} + K[H\Sigma H^T + R] (H^T)^{-1} \quad (5.3)$$

Substituting for $[H\Sigma H^T + R]$ by using Eq. (5.2) gives

$$\Sigma = [(H\Phi)^{-1} \Omega(1) + K\Omega(0)] (H^T)^{-1} \quad (5.4)$$

As noted earlier $\Omega(1)$ and $\Omega(0)$ cannot actually be calculated since the statistics of the distribution functions are unknown. However, they may be estimated recursively using standard stochastic approximation methods, which yield

1. The notation $\Omega(n+1, n)$ and $\Omega(n, n)$ has been simplified to $\Omega(1)$ and $\Omega(0)$ respectively denoting the difference between the time arguments.

$$\hat{\Omega}_n(1) = \hat{\Omega}_{n-1}(1) + a_n [\delta_{n+1} \delta_n^T - \hat{\Omega}_{n-1}(1)] \quad (5.5)$$

$$\hat{\Omega}_n(0) = \hat{\Omega}_{n-1}(0) + b_n [\delta_n \delta_n^T - \hat{\Omega}_{n-1}(0)] \quad (5.6)$$

where $\{a_n\}$ and $\{b_n\}$ satisfy the conditions of Theorem 4.1. Using this theorem, Eqns. (5.5) and (5.6) can be shown to converge to (5.1) and (5.2) respectively. Substituting these estimates into (5.4) along with the estimate of K_{opt} provides the desired relation for estimating Σ .

$$\hat{\Sigma}_n = [(H\Phi)^{-1} \hat{\Omega}_n(1) + \hat{K}_n \hat{\Omega}_n(0)] (H^T)^{-1} \quad (5.7)$$

To derive an expression for estimating the error covariance matrix Γ , a general expression for Γ must be developed. The relation given by Eq. (2.21) is deficient because it is valid only for optimal filtering. By definition

$$\Gamma = E\{[\underline{x}(n) - \hat{\underline{x}}(n)][\underline{x}(n) - \hat{\underline{x}}(n)]^T\}$$

Substituting the filtering equation for $\hat{\underline{x}}(n)$ yields

$$\Gamma = E\{[\underline{x}(n) - \Phi \hat{\underline{x}}(n-1) + K(z(n) - H\Phi \hat{\underline{x}}(n-1))][\underline{x}(n) - \Phi \hat{\underline{x}}(n-1) + K(z(n) - H\Phi \hat{\underline{x}}(n-1))]^T\}$$

By exploiting the independence of $\underline{v}(n)$ and $w(n)$, the above equation may be written

$$\begin{aligned}
\Gamma &= E\{[\underline{x}(n) - \Phi \hat{\underline{x}}(n-1)][\underline{x}(n) - \Phi \hat{\underline{x}}(n-1)]^T\} \\
&\quad - E\{[\underline{x}(n) - \Phi \hat{\underline{x}}(n-1)][\underline{x}(n) - \Phi \hat{\underline{x}}(n-1)]^T\} H^T K^T \\
&\quad - KH E\{[\underline{x}(n) - \Phi \hat{\underline{x}}(n-1)][\underline{x}(n) - \Phi \hat{\underline{x}}(n-1)]^T\} \\
&\quad + K E\{[\underline{z}(n) - \Phi \hat{\underline{x}}(n-1)][\underline{z}(n) - \Phi \hat{\underline{x}}(n-1)]^T\} K^T \\
&= \Sigma - \Sigma H^T K^T - KH \Sigma + K[H \Sigma H^T + R] K^T \\
&= [I - KH] \Sigma + \{K[H \Sigma H^T + R] - \Sigma H^T\} K^T \tag{5.8}
\end{aligned}$$

When K is equal to the optimal filter as given by Eq. (2.20), the term in braces is zero, and Eq. (5.8) reduces to the optimum expression for the error covariance matrix. Observe from Eq. (5.1) that Eq. (5.8) may be rewritten

$$\Gamma = [I - KH] \Sigma - (H\Phi)^{-1} \Omega(1) K^T \tag{5.9}$$

Replacing the unknown quantities with their respective estimates, the algorithm for $\hat{\Gamma}_n$ is obtained

$$\hat{\Gamma}_n = [I - \hat{K}_n H] \hat{\Sigma}_n - (H\Phi)^{-1} \hat{\Omega}_n(1) \hat{K}_n^T \tag{5.10}$$

To estimate R , the estimator for Σ and $\Omega(0)$ is substituted into Eq. (5.2), giving

$$\hat{R}_n = \hat{\Omega}_n(0) - H \hat{\Sigma}_n H^T \tag{5.11}$$

$$= \hat{K}_n^{-1} \hat{\Gamma}_n H^T \tag{5.12}$$

Equation (5.11) is generally preferable to Eq. (5.12) for estimating R because it does not require the existence or calculation of the inverse of $\hat{K}_n$. If R^{-1} is required as it is in the smoothing Eq. (2.32), it may be obtained directly from Eq. (2.20)

$$R^{-1} = (H^T)^{-1} \Gamma^{-1} K$$

from which the estimate of R inverse is

$$\hat{R}_n^{-1} = (H^T)^{-1} \hat{\Gamma}_n^{-1} \hat{K}_n \quad (5.13)$$

This result does require the inverse of $\hat{\Gamma}_n$, but it is usually better conditioned than $\hat{R}_n$.

The last quantity to be estimated is Q . From Eq. (2.22)

$$Q = \Sigma - \Phi \Gamma \Phi^T \quad (5.14)$$

$$= \Sigma - \Phi [I - KH] \Sigma \Phi^T \quad (5.15)$$

Substituting the estimates of the unknown quantities into Eqns. (5.14) and (5.15) respectively results in

$$\hat{Q}_n = \hat{\Sigma}_n - \Phi \hat{\Gamma}_n \Phi^T \quad (5.16)$$

$$= \hat{\Sigma}_n - \Phi [I - \hat{K}_n H] \hat{\Sigma}_n \Phi^T \quad (5.17)$$

Equations (5.5), (5.6), (5.7), (5.10), (5.11), (5.12), (5.13), (5.16) and (5.17) provide a complete set of algorithms for estimating Σ , Γ , R^{-1} , R , and Q . Where more than one method is available for estimating the same quantity, the algorithm used is dictated by computational considerations. However, these algorithms may be significantly

simplified by noting from the Learning Criterion Theorem 3-3 that

$$\Omega(1) = E\{\delta_{n+1}\delta_n^T\}$$

approaches the null matrix as $\hat{K}_n$ approaches K_{opt} . Thus for large n , algorithm (5.7) reduces to

$$\hat{\Sigma}_n = \hat{K}_n \Omega_n(0)(H^T)^{-1} \quad (5.18)$$

This remains a close approximation for small and intermediate values of n if the accelerated algorithm defined by Eqns. (4.43) and (4.37) for estimating K_{opt} is employed. Another benefit accrued by using Eq. (4.43) and (4.37) is made obvious by rewriting it as a stochastic approximation

$$W_{n+1}^{-1} = W_n^{-1} + \frac{1}{n+1} (\delta_{n+1}\delta_n^T - W_n^{-1}) \quad (5.19)$$

Comparing this equation with Eq. (5.6) reveals that W_n^{-1} is just $\hat{\Omega}_n(0)$. Therefore, algorithm (5.18) is simply

$$\hat{\Sigma}_n = \hat{K}_n W_n^{-1}(H^T)^{-1} \quad (5.20)$$

Thus, the algorithm has not only been greatly simplified, but the computation of the two recursions Eqns. (5.5) and (5.6) have been eliminated. The algorithm (5.10) for estimating Γ becomes

$$\hat{\Gamma}_n = [I - \hat{K}_n H] \hat{\Sigma}_n \quad (5.21)$$

The expression for estimating R and Q remain unchanged because they were written in terms of $\hat{\Gamma}_n$ and $\hat{\Sigma}_n$.

5.3 Non-Existence of H Inverse

When the inverse of the observation matrix H does not exist, the stochastic algorithms derived in Chapter 4 for learning the optimum filter must be modified or a new invertible observation matrix must be created. Two methods for forming a new observation matrix and one method for modifying the algorithms are presented. They are

1. measurement stacking with time delay,
2. measurement stacking with state augmentation, and
3. residual stacking.

For purposes of discussion the following typical third order, single-input, single-output system is used. In phase variable representation, the system is described by Eqns. (1.3) and (1.4)

$$\underline{\tilde{x}}(n+1) = \Phi \underline{\tilde{x}}(n) + \underline{\tilde{w}}(n) \quad (1.3)$$

$$\underline{\tilde{z}}(n) = H \underline{\tilde{x}}(n) + \underline{\tilde{v}}(n) \quad (1.4)$$

where

$$\Phi = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ a & b & c \end{bmatrix}$$

$$H = [1 \quad 0 \quad 0]$$

and $\underline{\tilde{w}}^T(n) = [0 \quad 0 \quad w(n)]$

The results to be presented are valid for general Φ and H matrices.

5.3.1 Measurement Stacking with Time Delay

This is a straight forward, brute-force technique. It consists simply of accumulating m observations where m is the order of the system, and creating an equivalent $m \times m$ observation matrix, say H_N , whose inverse does exist.

For illustration, consider the example third order system. Since H is 1×3 , three measurements must be stacked at every iteration. Making use of the deterministic relationships between the state variables and starting at time n , the measurements are

$$\begin{aligned} z(n) &= x_1(n) + v(n) \\ z(n+1) &= x_1(n+1) + v(n+1) \\ &= x_2(n) + v(n+1) \\ z(n+2) &= x_1(n+2) + v(n+2) \\ &= x_3(n) + v(n+2) \end{aligned}$$

Therefore,

$$\underline{z}(n) = H_N \underline{x}(n) + \underline{v}(n)$$

where $H_N = I$

and $\underline{v}^T(n) = [v(n), v(n+1), v(n+2)]$

Note that to avoid violating the assumption that the measurement noise $v(n)$ be uncorrelated from stage to stage, three new observations must be made at each iteration. This means that the filtering estimate of $\underline{x}(n)$ is being performed at time $n+2$. To obtain the estimate of $\underline{x}(n+2)$, $\hat{\underline{x}}(n)$ must be predicted ahead two stages. This is the disadvantage of the technique. However, it is not as deleterious as it appears, because

the message generating noise directly affects only the prediction of $\underline{x}_3(n)$. Still this method is not recommended unless the number of measurements required to form H_N is small compared to m .

5.3.2 Measurement Stacking With State Augmentation

The prediction problem associated with the previous technique can be avoided by using the standard technique of augmenting the state vector $\underline{x}(n)$ with the measurement noise $\underline{v}(n)$ [Meditch, 1969, p. 195]. This forms a new $m \times p$ state vector $\underline{y}(n)$ such that

$$\begin{aligned}\underline{y}(n) &= \begin{bmatrix} \underline{x}(n) \\ \underline{v}(n) \end{bmatrix} \\ \underline{y}(n+1) &= \begin{bmatrix} \Phi & 0 \\ 0 & 0 \end{bmatrix} \underline{y}(n) + \begin{bmatrix} \underline{w}(n) \\ \underline{v}(n+1) \end{bmatrix} \\ &= \Phi_N \underline{y}(n) + \begin{bmatrix} \underline{w}(n) \\ \underline{v}(n+1) \end{bmatrix}\end{aligned}\tag{5.22}$$

where the 0's are null matrices of appropriate dimension, and

$$\begin{aligned}\underline{z}(n) &= [H, I] \underline{y}(n) \\ &= H\underline{x}(n) + \underline{v}(n)\end{aligned}\tag{5.23}$$

This augmented structure requires only one new measurement per estimation stage regardless of the H matrix. The new $(m+p) \times (m+p)$ invertible observation matrix H_N is formed by stacking the new measurement with predicted measurements from the previous iteration.

This technique is made clearer by considering the example system given by Eqns. (5.22) and (5.23). The four required measurements are

$$\begin{aligned}
 z(n) &= [H, I] \tilde{y}(n) \\
 &= [1, 0, 0, 1] \tilde{y}(n) \\
 &= x_1(n) + v(n)
 \end{aligned} \tag{5.24}$$

$$\begin{aligned}
 z(n)_1 &= [H\Phi_N, 0] \hat{\tilde{y}}(n-1) \\
 &= [1, 0, 0, 0] \hat{\tilde{y}}(n-1) \\
 &= x_1(n) + e_1(n)
 \end{aligned} \tag{5.25}$$

$$\begin{aligned}
 z(n)_2 &= [H\Phi_N\Phi_N, 0] \hat{\tilde{y}}(n-1) \\
 &= [0, 1, 0, 0] \hat{\tilde{y}}(n-1) \\
 &= x_2(n) + e_2(n)
 \end{aligned} \tag{5.26}$$

$$\begin{aligned}
 z(n)_3 &= [H\Phi_N\Phi_N\Phi_N, 0] \hat{\tilde{y}}(n-1) \\
 &= [0, 0, 1, 0] \hat{\tilde{y}}(n-1) \\
 &= x_3(n) + e_3(n)
 \end{aligned} \tag{5.27}$$

where $e_i(n)$, $i=1,2,3$ is the difference between the value of $x_i(n)$ and $\hat{x}_i(n|n-1)$. From Eqns. (5.24) - (5.27), it is obvious that H_N is

$$H_N = \begin{bmatrix} 1 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \end{bmatrix}$$

This technique is recommended when the dimension of $z(n)$ is small compared to the dimension of $x(n)$. Note that both techniques for

forming an invertible observation matrix inherently change the original Kalman filter structure. The standard Kalman filter is an $m \times p$ matrix, where p is the dimension of the $\tilde{z}(n)$ and satisfies the inequality

$$1 \leq p \leq m$$

The new observation matrix H_N produces an m dimensional measurement. Thus, the new Kalman filter is always an $m \times m$ matrix. A technique which preserves the original structure of the problem is presented in Section 5.3.3.

5.3.3 Predicted Residual Stacking

The development of this section is based on the common case where H is an $1 \times m$ matrix, i.e., only the output state variable can be observed. This implies that the Kalman filter is an $m \times 1$ vector. A similar development can be performed for any given H matrix. Here the predicted residual is defined by

$$z(n+r) - H\Phi^{r+1} \hat{\tilde{x}}(n-1) \triangleq \Delta_{n+r}$$

Now let

$$\tilde{\Delta}(n+m, n) \triangleq \begin{bmatrix} H\Phi^m \\ \vdots \\ H\Phi^2 \\ H\Phi \end{bmatrix}^{-1} \begin{bmatrix} \Delta_{n+m} \\ \vdots \\ \Delta_{n+2} \\ \Delta_{n+1} \end{bmatrix}$$

The loss function of Eq. (4.18) is now redefined as

$$\begin{aligned}
& \ell(\hat{\Delta}(n+m, n) - \Delta(n+m, n)) \\
& = [\hat{\Delta}(n+m, n) - K_n \delta_n]^T [\hat{\Delta}(n+m, n) - K_n \delta_n]
\end{aligned} \tag{5.29}$$

where δ_n is as defined in Chapter 3. To see that Eq. (5.29) is a meaningful definition for the loss function, form the linear regression equation

$$\hat{\Delta}(n+m, n) = K_n \delta_n + e_n \tag{5.30}$$

Now the prediction residual transition vector is

$$K_n = E\{\hat{\Delta}(n+m, n) \delta_n^T\} (E\{\delta_n \delta_n^T\})^{-1} \tag{5.31}$$

After some computation, it is found that

$$E\{\hat{\Delta}(n+m, n) \delta_n^T\} = \Sigma H^T$$

and

$$E\{\delta_n \delta_n^T\} = (H \Sigma H^T + R)$$

Therefore,

$$\begin{aligned}
K_n &= \Sigma H^T (H \Sigma H^T + R)^{-1} \\
&= K_{\text{opt}}
\end{aligned}$$

Consequently, K_n represents the optimal Kalman filter matrix.

Also, $K_n \delta_n$ may be considered to represent the predicted value of $\Delta(m+m, n)$. Since K_n cannot be calculated, K_{opt} is estimated as in Chapter 4 using the theory of stochastic approximation. Applying this theory to the loss function of Eq. (5.29) gives

$$\begin{aligned}
\hat{K}_{n+1} &= \hat{K}_n - a_{n+1} \nabla_{\hat{K}} \ell(\Delta(n+m, m) - \hat{K}_n \delta_n) W_{n+1} \\
&= \hat{K}_n + a_{n+1} [\Delta(n+m, n) - \hat{K}_n \delta_n] \delta_n^T W_{n+1}
\end{aligned} \tag{5.32}$$

where $a_n = 1/n$ and W_{n+1} is a uniformly bounded sequence. The algorithm of Eq. (5.32) is analogous to Eq. (4.20) and can be shown to converge to K_{opt} with probability one. In general the method of this section is recommended over those of Section 5.3.1 and 5.3.2. However, the author conjectures that the mean-square error associated with the state estimation procedure of 5.3.2 is at least as small as that of this section.

5.4 Time-Varying Signal and Noise Statistics

In practice the stationary assumption of Chapters 3 and 4 is valid only for a finite time interval. If this time interval is less than the essential adaptation time of the learning algorithms, then the results of Chapter 4 may still be employed. Even in this case, it is important to determine when the time variations have changed the optimal filter significantly.

A statistical test for making such a determination can be devised from either Theorem 3-2 or Theorem 3-3 (the Learning Criterion). Theorem 3-2 states that $E\{\delta_{n+1}^T \delta_n\} = 0$ if and only if the filter is optimum, and Theorem 3-3 states that $E\{\delta_{n+1} \delta_n^T\} = [0]$ if and only if the filter is optimum. Either of these two facts may be used to develop a confidence interval test for the time average of $\delta_{n+1}^T \delta_n$ or $\delta_{n+1} \delta_n^T$. Alternately, the chi-square method can be used to test whether or not the times average of $\delta_{n+1}^T \delta_n$ or $\delta_{n+1} \delta_n^T$ is significantly different from

zero or the null matrix respectively. Since these techniques are standard in the statistical literature, the details are not presented here.

Returning now to the problem of a time-varying environment, two different approaches are considered. The first is useful when the time variation is slow and its exact nature is unknown. The second is valid only when the time variation of the filter is known.

5.4.1 Measurement Discounting

The algorithms of Chapter 4 weight all measurements equally over the adaptation interval. Thus as the interval increases, or equivalently as the number of iterations n increases, the effect of each measurement becomes vanishingly small. This effect may be eliminated by attenuating past observations. Such a procedure essentially shapes the memory of the algorithm. One simple way of limiting the memory is to weight the data equally over a finite number of observations under the assumption that the process is stationary over this time interval.

A different approach is to exponentially attenuate past measurements [Korn, 1966, p. 45]. To modify the accelerated algorithm of Eqns. (4.39) and (4.43) for an exponentially weighted memory requires that the intermediate loss function of Eq. (4.31) be redefined as

$$L_n(K_n) = \frac{1}{2n} \sum_{j=0}^{n-1} (\Delta_{j+1} - K_n \delta_j)^T \lambda^{n-1-j} (\Delta_{j+1} - K_n \delta_j) \quad (5.33)$$

where $\lambda \in (0,1)$.

Now following the development which led to Eqns. (4.39) and (4.43), the recursive relation for learning K_{opt} is found to be

$$\hat{K}_{n+1} = \hat{K}_n + [\Delta_{n+1} - \hat{K}_n \delta_n] \delta_n^T \frac{W_{n+1}}{n+1} \quad (5.34)$$

and

$$W_{n+1} = \frac{n+1}{n\lambda} \{W_n - W_n \delta_n [n\lambda + \delta_n^T W_n \delta_n]^{-1} \delta_n^T W_n\} \quad (5.35)$$

The term λ which provides the exponential damping also determines the effective sample size of the estimation process. Both the ability of the algorithm to refine the estimates and to follow time variations are determined by λ . Unfortunately the two requirements are conflicting. The first implies a large value (near unity) of λ , and the second implies a small value (near zero) of λ . Thus, a compromise is necessary in selecting λ . Usually λ must be greater than 0.95 to ensure stability.

When the time evolution of $\underline{K}(n)$ is known with $Q(n)$ and $R(n)$ unknown, the procedure of Section 5.4.2 is recommended.

5.4.2 Known Time Evolution of $K(n)$

Let $\underline{K}(n)$ evolve in time in a manner described by Eq. (5.36)

$$\underline{K}(n+1) = F[\underline{K}(n)] + \omega(n) \quad (5.36)$$

where F is the $m \times m$ filter transition matrix and $\omega(n)$ is an $m \times 1$ random perturbation vector. In this case Eq. (4.39) is modified in a manner that is reminiscent of Kalman's filtering equation. This modification yields

$$\hat{\tilde{K}}_{n+1} = F(\hat{\tilde{K}}_n) + [\Delta_{n+1} - F(\hat{\tilde{K}}_n) \delta_n] \delta_n^T \frac{W_{n+1}}{n+1} \quad (5.37)$$

5.5 Nonlinear Dynamics

When the plant dynamics are nonlinear, the filter that is optimal in the sense that it minimizes the mean square error is unknown. In this section a different definition for optimality of the filter is given and adaptive algorithms which learn this optimal filter are presented. The Learning Criterion as given by Eq. (3.6)

$$E\{\delta_{j+1} \delta_j^T\} = [0] \quad \text{for all } j \leq n \quad (3.6)$$

of Theorem 3-3 is used as the definition for optimal filtering. The definition is selected because it is equivalent to optimal Kalman filtering when the dynamics are linear and because it forces the residual time series to be white. This latter fact implies that all the information possible has been extracted from this series. The system model considered is given by the equations

$$x(n+1) = f(x_n) + w(n) \quad (5.38)$$

$$z(n) = x(n) + v(n) \quad (5.39)$$

and the form of the filtering equation is specified to be

$$\hat{x}(n+1) = f(\hat{x}(n)) + K_n [z(n+1) - f(\hat{x}(n))] \quad (5.40)$$

Note that each of these equations is the same as their linear counterparts with the linear dynamics replaced with the nonlinear dynamics.

Equation (5.38) corresponds to a stochastic signal with a nonrational spectral density.

From Eq. (5.40) it is obvious that the residual time series is

$$\delta_n = z(n+1) - f(\hat{x}(n)) \quad (5.41)$$

Thus from the definition of optimal filtering given by Eq. (3.6), the filter is optimum if and only if it is chosen such that

$$E\{[z(j+1) - f(\hat{x}(j))][z(j) - f(\hat{x}(j-1))]\} = 0 \quad \text{for all } j \leq n \quad (5.42)$$

Two different classes of nonlinearities are now considered. The first type satisfies a global Lipschitz condition

$$|f(x_n) - f(x_m)| \leq \phi |x_n - x_m| \quad \text{for all } x_n, x_m$$

For this class of nonlinearities the adaptive algorithm for learning the filter that satisfies Eq. (5.42) is given by Eq. (4.21) where $H = 1$ and ϕ is the Lipschitz constant.

$$\hat{K}_{n+1} = \hat{K}_n + \frac{1}{n} \phi^{-1} \delta_{n+1} \delta_n W_{n+1} \quad (5.43)$$

The second type of nonlinearity has an inverse. The dynamics $f(x_n)$ must be such that

$$f^{-1}[f(x_n)] = x_n \quad \text{for all } x_n \quad (5.44)$$

For this case the stochastic approximation algorithm is heuristically modified as follows

$$\hat{K}_{n+1} = \hat{K}_n + \frac{1}{n} f^{-1}(\delta_{n+1} \delta_n) W_{n+1} \quad (5.45)$$

The engineering justification for the nonlinear algorithms of Eqns. (5.43) and (5.45) is that they "work". That is, they allow one to make "useful" predictions. Experimental confirmation of this statement is presented in the next chapter.

5.6 Summary

This chapter has derived algorithms for estimating useful quantities such as Σ , Γ , R and Q . This knowledge is necessary to solve the optimal smoothing problem. The restrictive condition of Chapter 4 requiring that H^{-1} exist has been lifted. The algorithms of Chapter 4 were then heuristically modified to accommodate special classes of time-varying and nonlinear dynamics. This chapter also concludes the theoretical aspects of the report. The next chapter contains a presentation of the experimental results of this study.

CHAPTER 6

NUMERICAL SIMULATION OF THE ADAPTIVE ALGORITHMS

6.1 Introduction and Organization of the Chapter

In Chapter 4 adaptive stochastic algorithms were derived for directly learning the optimal Kalman filter K_{opt} . Chapter 5 extended these results to include estimation of Σ , Γ , Q and R . These quantities are required for adaptive optimal smoothing. The restrictive assumption that H^{-1} exist was also lifted and the algorithms of Chapter 4 were modified to accommodate slowly time-varying systems statistics and certain classes of nonlinear dynamics.

In this chapter the adaptive algorithms of Chapter 4 and Sections 5.2 and 5.5 are applied to specific systems. The experimental results were obtained by simulating the different systems and the associated algorithms on a CDC 6400 computer. In interpreting these simulations, it is emphasized that double precision arithmetic was not used.

The first part of the chapter considers linear systems. Section 6.2 presents the experimental data for first order through fourth order dynamic systems. Nonlinear plants are discussed in Section 6.3. Section 6.4 applies the theory to the problem of estimating the states of a thermionic reactor.

6.2 Linear System Simulations

For convenience the dynamic system equations and adaptive algorithms to be studied are repeated here. The system dynamics and observations are defined by

$$\underline{x}(n+1) = \Phi \underline{x}(n) + \underline{w}(n) \quad (1.3)$$

$$\underline{z}(n) = H \underline{x}(n) + \underline{v}(n) \quad (1.4)$$

The filtering equation is

$$\hat{\underline{x}}(n) = \Phi \hat{\underline{x}}(n-1) + K[z(n) - H\Phi \hat{\underline{x}}(n-1)] \quad (2.19)$$

where the estimate $\hat{\underline{x}}(n)$ is optimal if and only if K is determined according to Eq. (2.20). The non-accelerated adaptive algorithm for learning the optimal filter K_{opt} when R and Q are unknown is given by

$$\hat{K}_{n+1} = \hat{K}_n + \frac{1}{n+1} [\Delta_{n+1} - \hat{K}_n \delta_n] \delta_n^T \quad (4.28)$$

where

$$\Delta_{n+1} = (H\Phi)^{-1} [z(n+1) - H\Phi \hat{\underline{x}}(n-1)]$$

and

$$\delta_n = \underline{z}(n) - H\Phi \hat{\underline{x}}(n-1)$$

The accelerated version of Eq. (4.28) is

$$\hat{K}_{n+1} = \hat{K}_n + \frac{1}{n+1} [\Delta_{n+1} - \hat{K}_n \delta_n] \delta_n^T W_{n+1} \quad (4.43)$$

where

$$W_{n+1} = \frac{n+1}{n} [W_n + W_n \delta_n (n + \delta_n^T W_n \delta_n)^{-1} \delta_n^T W_n] \quad (4.39)$$

The equations for estimating Σ , Γ , R and Q are

$$\hat{\Sigma}_n = \hat{K}_n W_n^{-1} (H^T)^{-1} \quad (5.19)$$

$$\hat{\Gamma}_n = [I - \hat{K}_n H] \hat{\Sigma}_n \quad (5.20)$$

$$\hat{R}_n = W_n^{-1} - H \hat{\Sigma}_n H^T \quad (5.21)$$

$$\hat{Q}_n = \hat{\Sigma}_n - \hat{\Phi} \hat{\Gamma}_n \hat{\Phi}^T \quad (5.16)$$

A listing of the computer program developed to perform these calculations may be found in the Appendix. Unless it is specified otherwise, both the plant perturbation noise $\underline{v}(n)$ and observation noise $\underline{w}(n)$ have Gaussian distributions.

6.2.1 First Order Dynamics

In this case Eqns. (2.20) and (2.21) can be solved explicitly for K_{opt}

$$K_{opt} = \frac{-(1 + Q/R - \phi^2)}{2\phi^2} + \left\{ \frac{(1 + Q/R - \phi^2)^2}{4\phi^4} + \frac{Q/R}{\phi^2} \right\}^{1/2} \quad (6.1)$$

EXAMPLE 6.1:

Given: $\phi = 0.5$, $H = 1.0$, $Q = R = 1.0$

For these parameters $K_{opt} = 0.53$. To begin the adaptive process, an initial value K_0 for the filter must be selected. In the scalar case $K_{opt} \in [0,1]$. Choosing $K_0 = 0.0$ corresponds to the one extreme of assuming the measurements are just noise, i.e., they contain no information. For this case the experimental results for both the non-accelerated

adaptive algorithms are plotted in Fig. 6.1. The solid line represents the optimum value of the filter. The filter initial condition of zero is denoted by a small circle on the 0 iteration line. The accelerated algorithm has converged to within 10% of its optimal value in 80 iterations. The non-accelerated algorithm requires 500 iterations to attain this accuracy. Both algorithms converge asymptotically to K_{opt} . This serves to verify experimentally the observations made in Chapter 4.

From Eqns. (5.16), (5.19), (5.20), and (5.21), Σ , Γ , $\bar{R}$, and Q may be estimated. Using these quantities, estimates of M , S , U , and J , which are required for optimal smoothing, may be obtained. These estimates after 1000 iterations of algorithm (4.43) are compared with their optimum values in Table 6.1.

TABLE 6.1

COMPARISON OF THE ESTIMATES AND THE OPTIMAL VALUES FOR EXAMPLE 6.1

Quantity	K	Σ	Γ	$\bar{R}$	Q	M	S	U	J, N=1
Optimal Value	.531	1.13	.531	1.0	1.0	.235	.1250	1.888	0.235
Estimate	.528	1.15	.542	1.04	.994	.236	.1245	1.834	0.236

To illustrate the insensitivity of the adaptive algorithms to the initial guess for the filter, the convergence properties are examined for $K_0 = 1.0$. This corresponds to the other extreme of assuming the measurements are perfect, i.e., they contain no noise. The experimental results for the accelerated and non-accelerated algorithms are shown in Fig. 6.2. Again the accelerated algorithm converges much more rapidly.

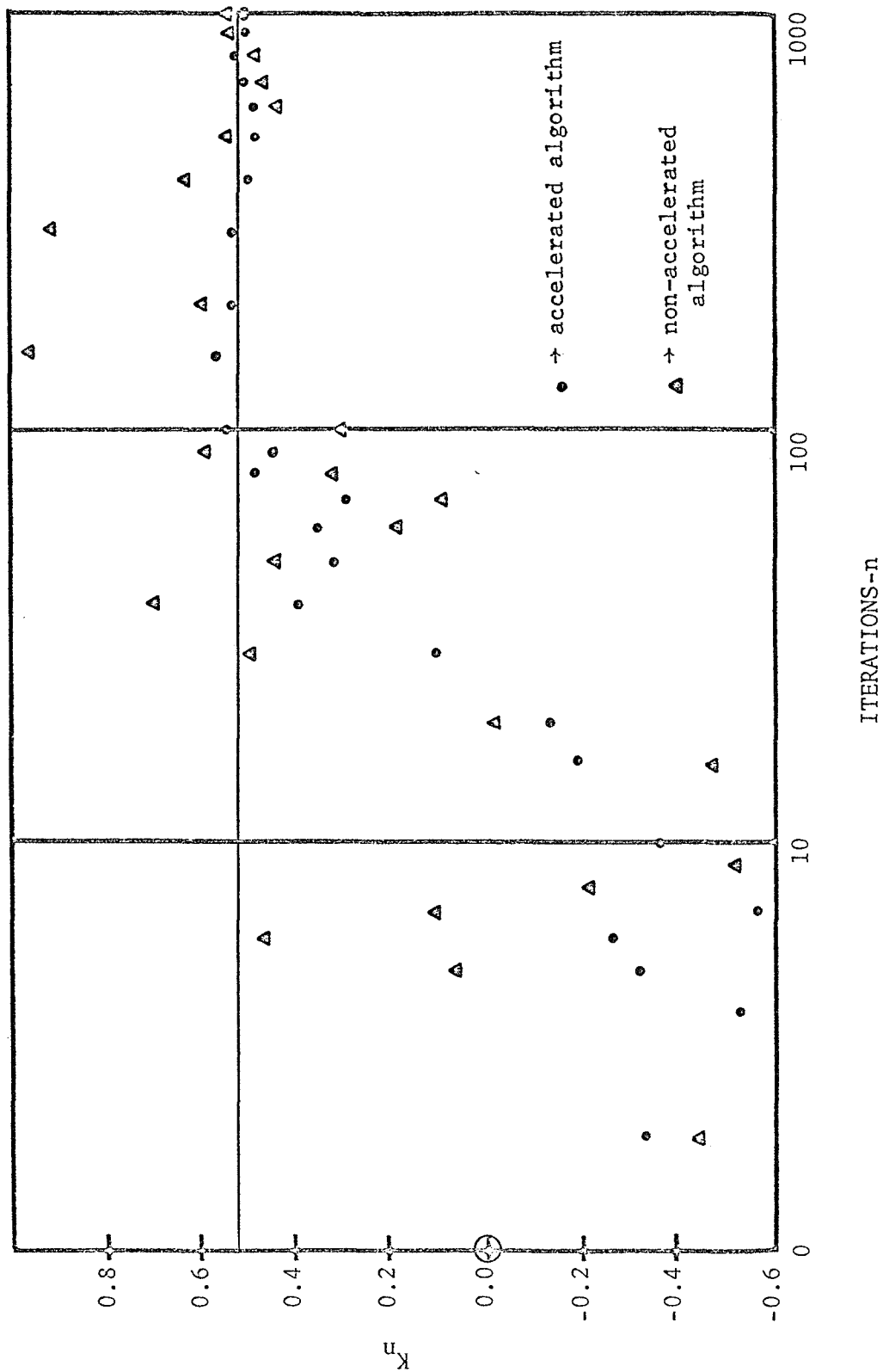


Fig. 6.1 Time Evolution of K_n for Example 6.1: $K_0 = 0.0$

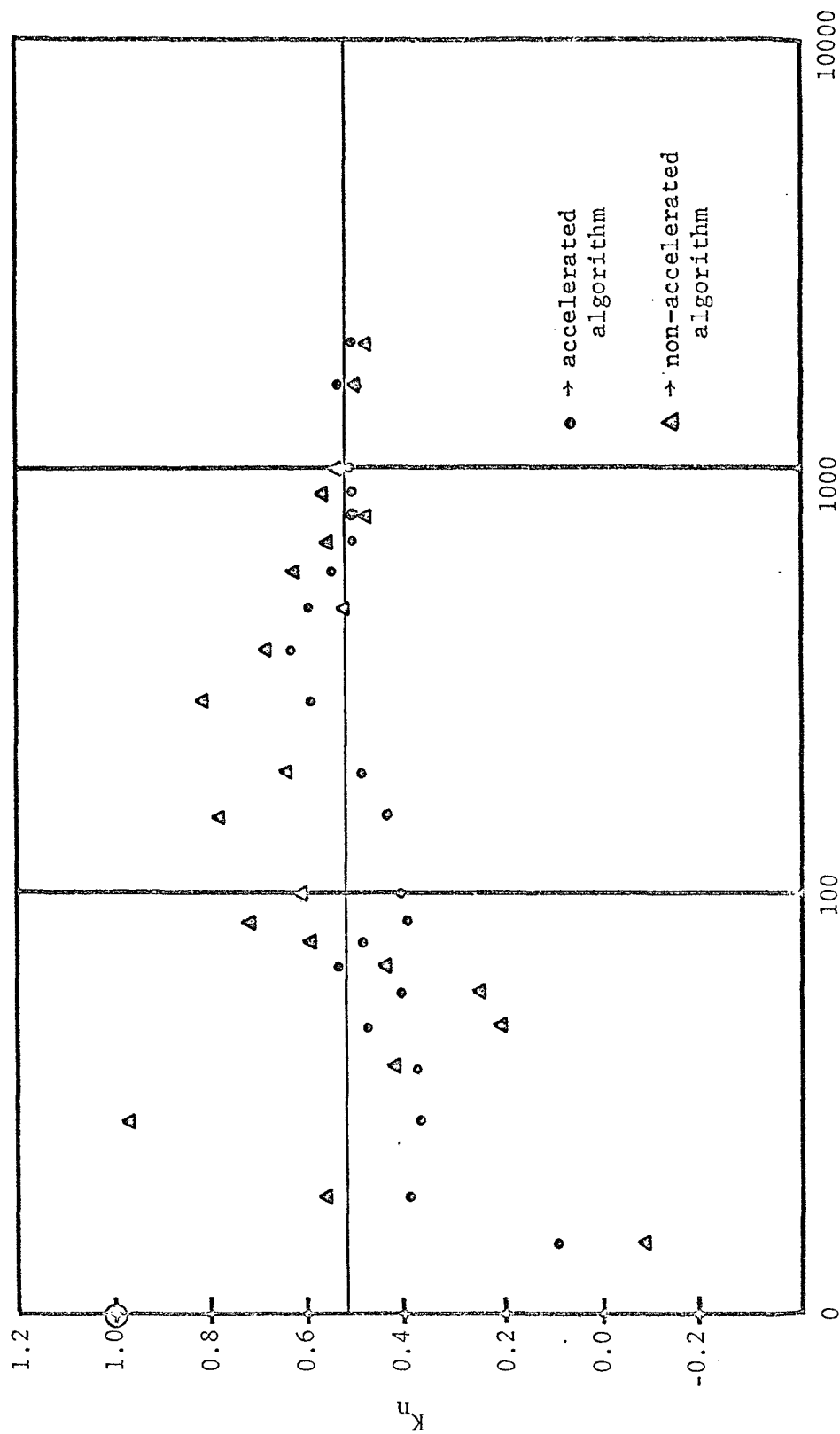


Fig. 6.2 Time Evolution of K_n for Example 6.1: $K_0 = 1.0$

EXAMPLE 6.2:

Given: $\phi = 0.5$, $H = 1.0$, $Q = 1.0$, $R = 4Q$

The purpose of this example is to demonstrate the power of the self-adaptive algorithm of Eqns. (4.43) and (4.39) to learn the optimum filter in a very noisy environment. From a communication theory viewpoint, Eq. (1.3) is simply the signal generating process and $\tilde{x}(n)$ is the signal. The signal to noise ratio for this problem is given by the ratio

$$\frac{\phi_x(o)}{\phi_v(o)}$$

where $\phi_x(\tau)$ denotes the autocorrelation function of x . $\phi_x(o)$ is

$$\begin{aligned}\phi_x(o) &= \frac{\phi_w(o)}{(1-\phi^2)} \\ &= \frac{Q}{(1-\phi^2)}\end{aligned}\tag{6.2}$$

Therefore, the signal to noise ratio is

$$\frac{\phi_x(o)}{\phi_v(o)} = \frac{Q}{(1-\phi^2)R}\tag{6.3}$$

For the parameters of this example, Eq. (6.3) has the numerical value of

$$\frac{\phi_x(o)}{\phi_v(o)} = \frac{1}{3}$$

or

$$\left. \frac{\phi_x(o)}{\phi_v(o)} \right|_{\text{db}} = -9.55 \text{ db.}$$

A signal to noise ratio of -9.55 db is extremely low. Still, the adaptive algorithm learns the optimal filter quite well. The results after 1,000 iterations are summarized in Table 6.2. The initial value of the filter was 1.0. Similar results were obtained for $K_0 = 0.0$.

TABLE 6.2

COMPARISON OF THE ESTIMATES AND THE OPTIMAL VALUES FOR EXAMPLE 6.2

Quantity	K	Σ	Γ	R	Q
Optimal Value	.236	1.24	.944	4.00	1.00
Estimate	.240	1.29	.961	4.08	1.04

Other first order problems were considered with ϕ being varied from 0.2 to 1.0 (the threshold of instability) and with different distributions functions for $\tilde{w}(n)$ and $\tilde{v}(n)$. In all cases K_n converged to K_{opt} within 1,000 iterations.

6.2.2 Second Order Dynamics

EXAMPLE 6.3:

Given:

$$\phi = \begin{bmatrix} 0.966 & 0.0 \\ 0.155 & 0.894 \end{bmatrix}, \quad H = I, \quad Q = \begin{bmatrix} 1.441 & 0.738 \\ 0.738 & 0.610 \end{bmatrix}, \quad R = I$$

For this simulation K_0 was chosen to be the null matrix [0.0]. The data for the adaptation process are shown in Fig. 6.3. Again the process has essentially converged within a small number of iterations (approximately

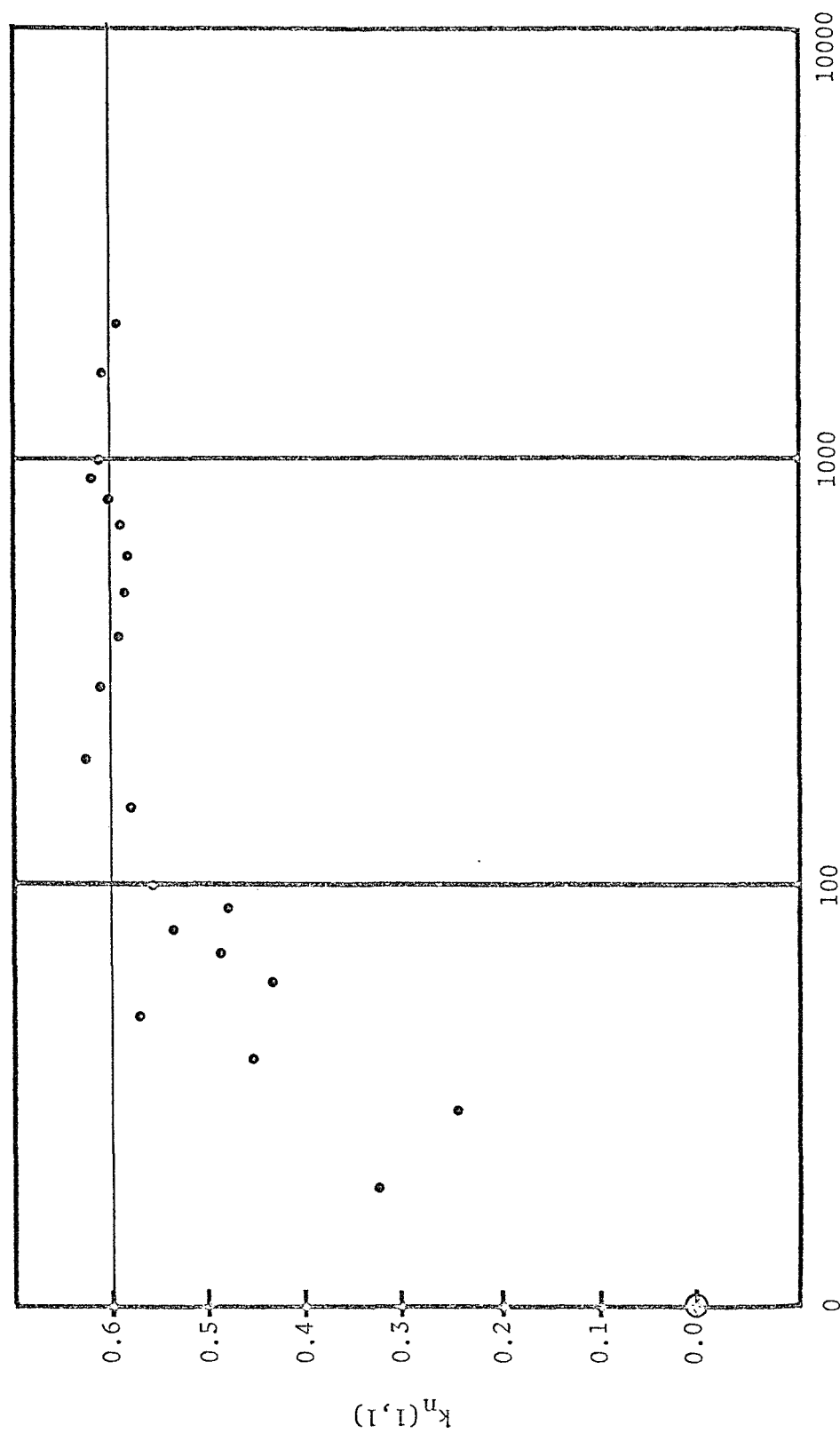


Fig. 6.3 Time Evolution of K_n for Example 6.3: Symmetry Forced

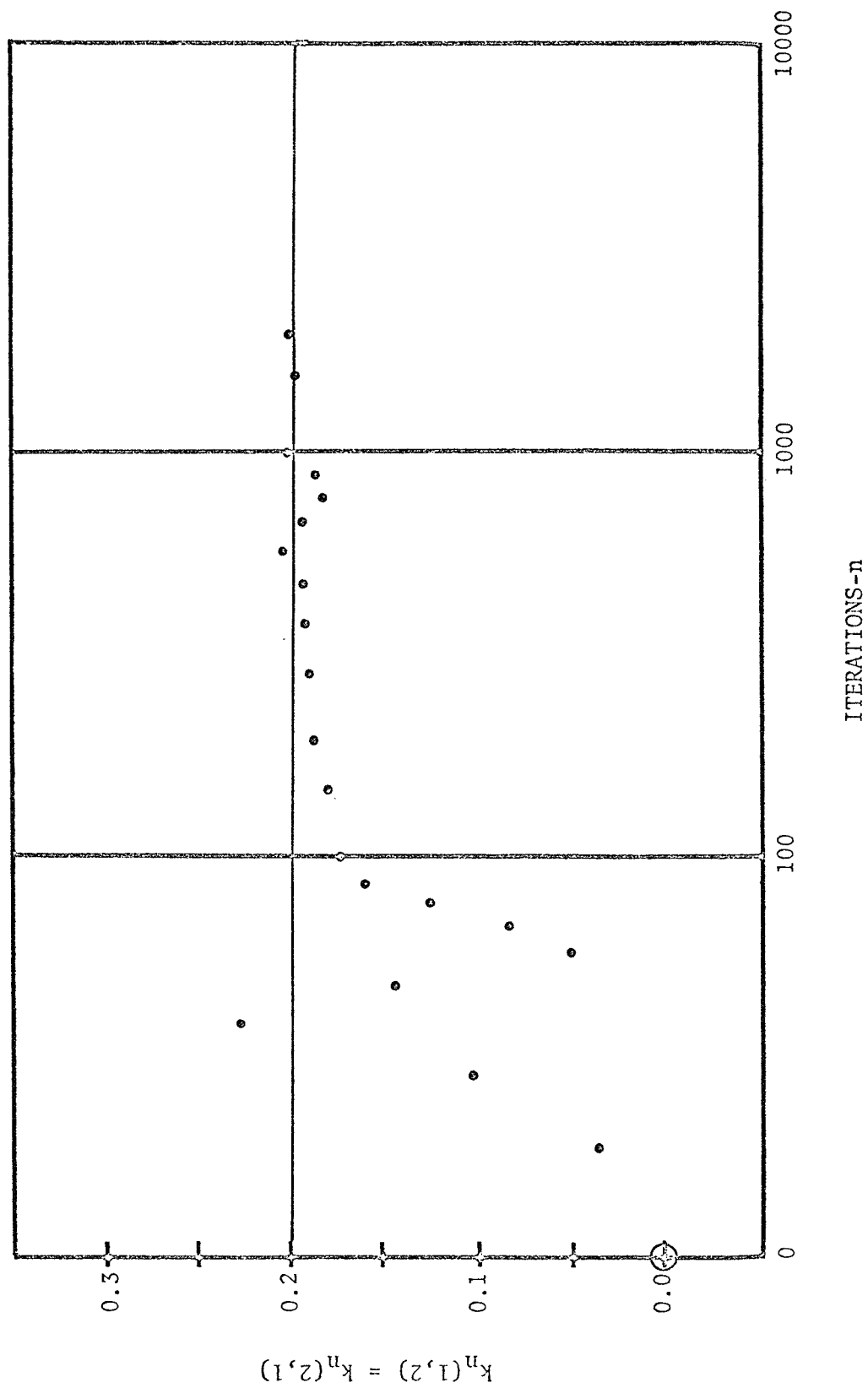


Fig. 6.3 Time Evolution of K_n for Example 6.3: Symmetry Forced (continued)

150). The estimates after 1000 iterations are compared with their optimal values below

$$\hat{K} = \begin{bmatrix} 0.591 & 0.202 \\ 0.202 & 0.397 \end{bmatrix} \quad K_{\text{opt}} = \begin{bmatrix} 0.600 & 0.200 \\ 0.200 & 0.400 \end{bmatrix}$$

$$\hat{\Sigma} = \begin{bmatrix} 1.893 & .960 \\ .960 & 1.009 \end{bmatrix} \quad \Sigma = \begin{bmatrix} 2.0 & 1.0 \\ 1.0 & 1.0 \end{bmatrix}$$

$$\hat{\Gamma} = \begin{bmatrix} .635 & .203 \\ .203 & .408 \end{bmatrix} \quad \Gamma = \begin{bmatrix} .60 & .20 \\ .20 & .40 \end{bmatrix}$$

$$\hat{R} = \begin{bmatrix} 1.164 & -.032 \\ -.032 & .985 \end{bmatrix} \quad R = \begin{bmatrix} 1.0 & 0.0 \\ 0.0 & 1.0 \end{bmatrix}$$

$$\hat{Q} = \begin{bmatrix} 1.301 & .688 \\ .688 & .611 \end{bmatrix} \quad Q = \begin{bmatrix} 1.441 & 0.738 \\ 0.738 & 0.610 \end{bmatrix}$$

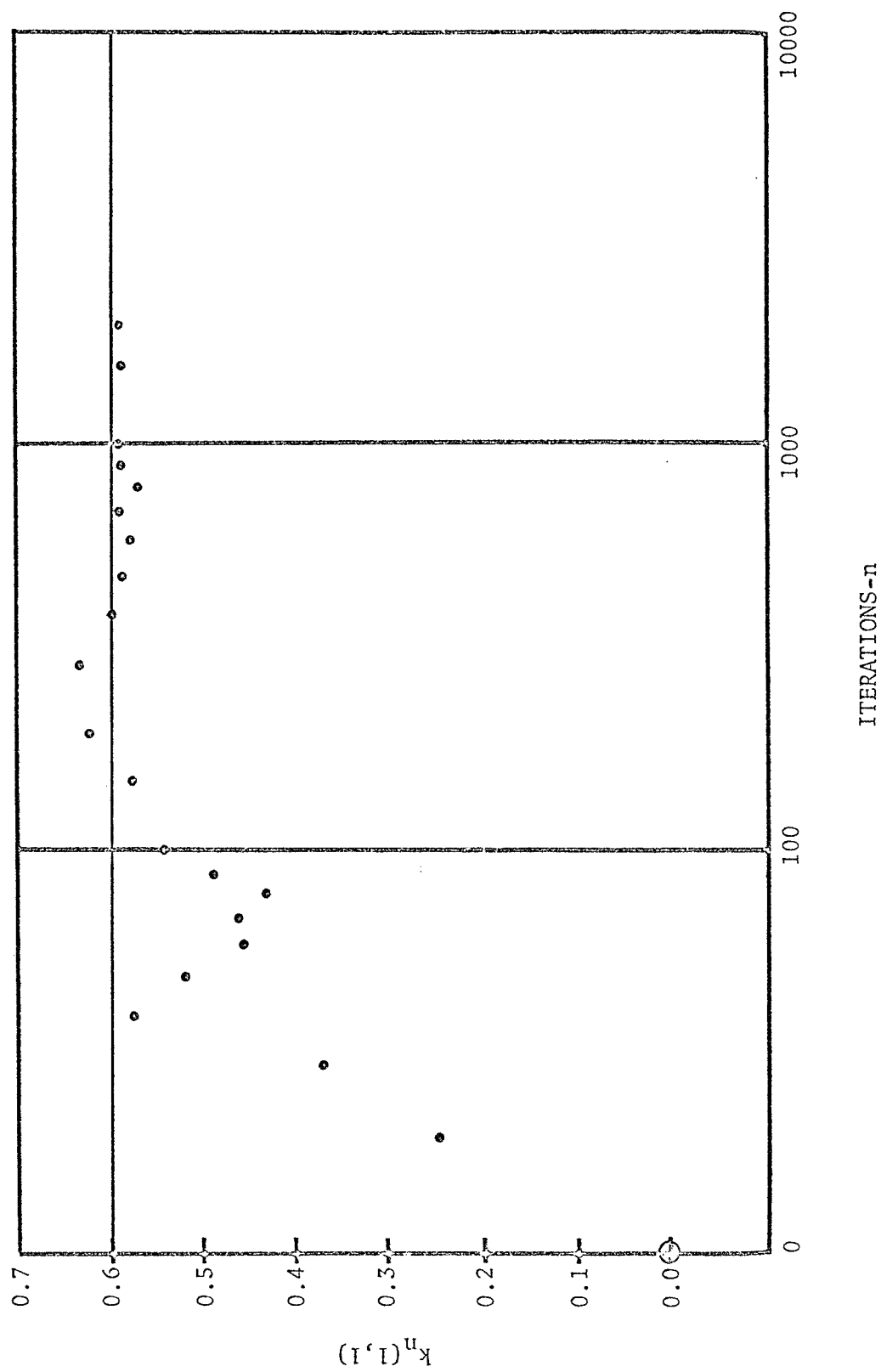
Note from Fig. 6.3 that symmetry has been forced on $k_n(1,2)$ and $k_n(2,1)$. The adaptation process for the same system and K_0 , but where symmetry is not forced, is illustrated in Fig. 6.4. As can be seen, $k_n(1,2)$ and $k_n(2,1)$ converge more rapidly when symmetry is forced.

6.2.3 Third Order Dynamics

EXAMPLE 6.4:

Given:

$$\Phi = \begin{bmatrix} 0.88 & 0.112 & 0.05 \\ 0.0 & 0.76 & 0.0 \\ 0.06 & 0.0 & 0.97 \end{bmatrix} \quad H = I$$



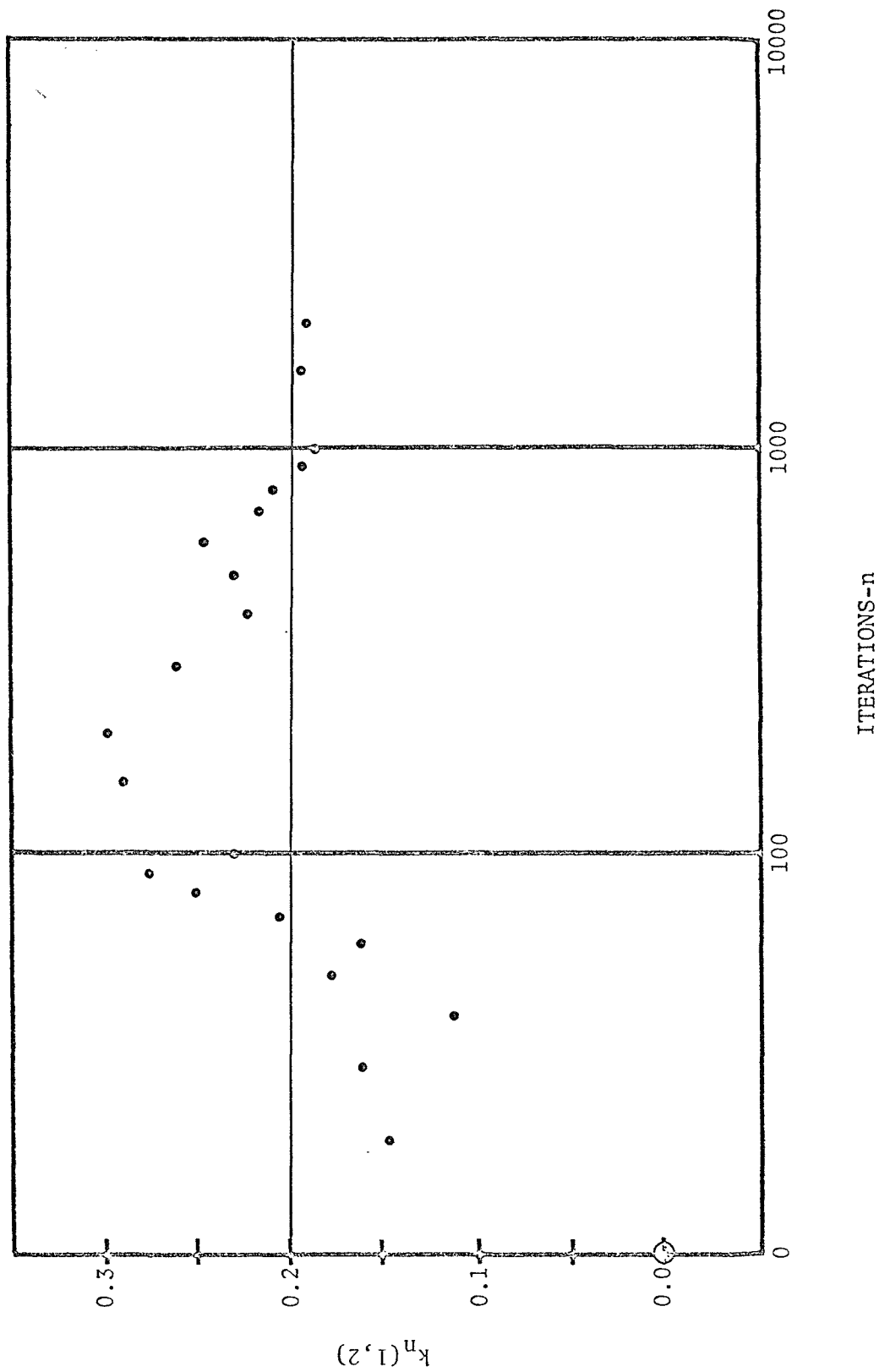


Fig. 6.4 Time Evolution of K_n for Example 6.3: Symmetry Not Forced (continued)

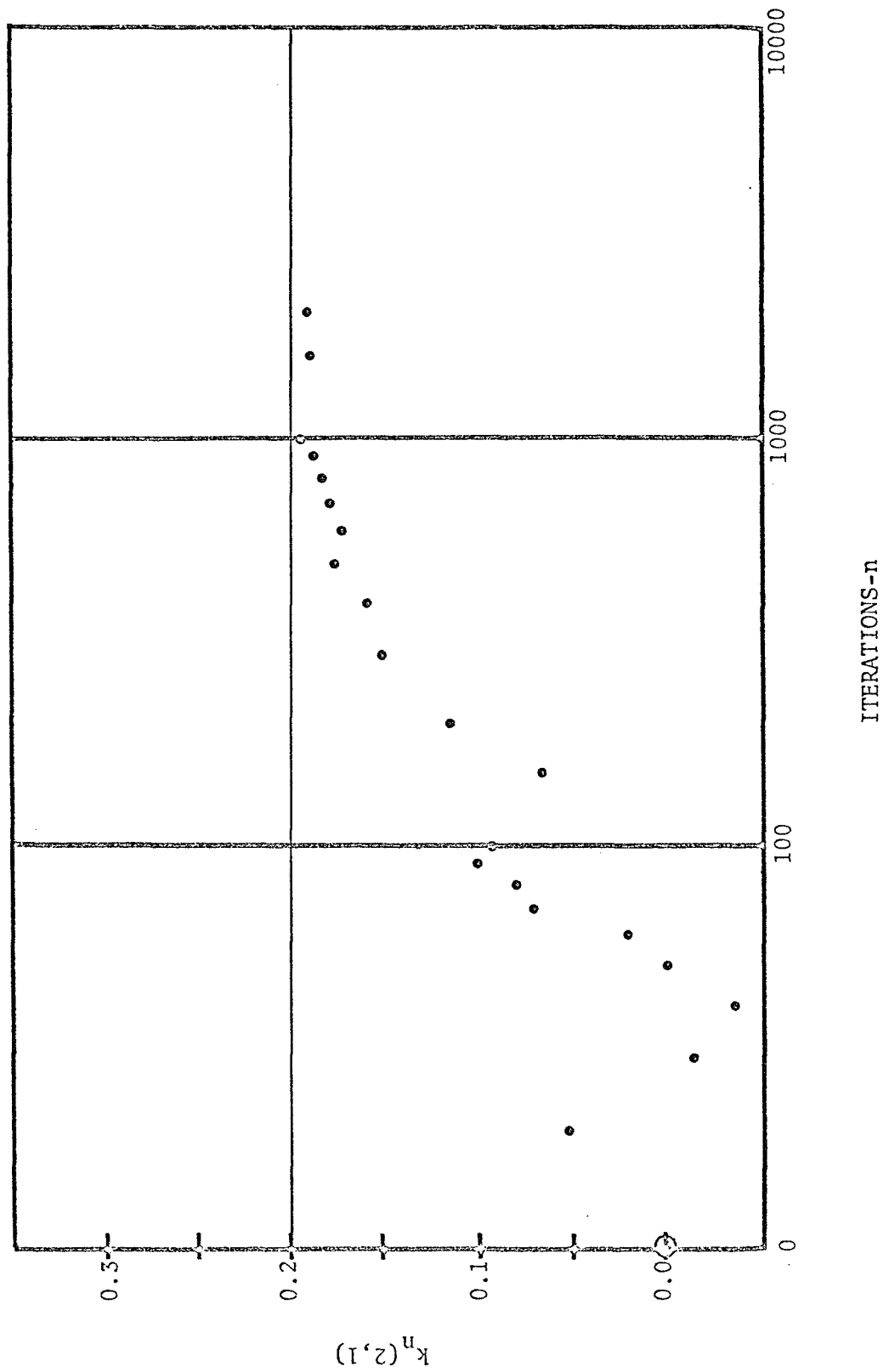


Fig. 6.4 Time Evolution of K_n for Example 6.3: Symmetry Not Forced (continued)

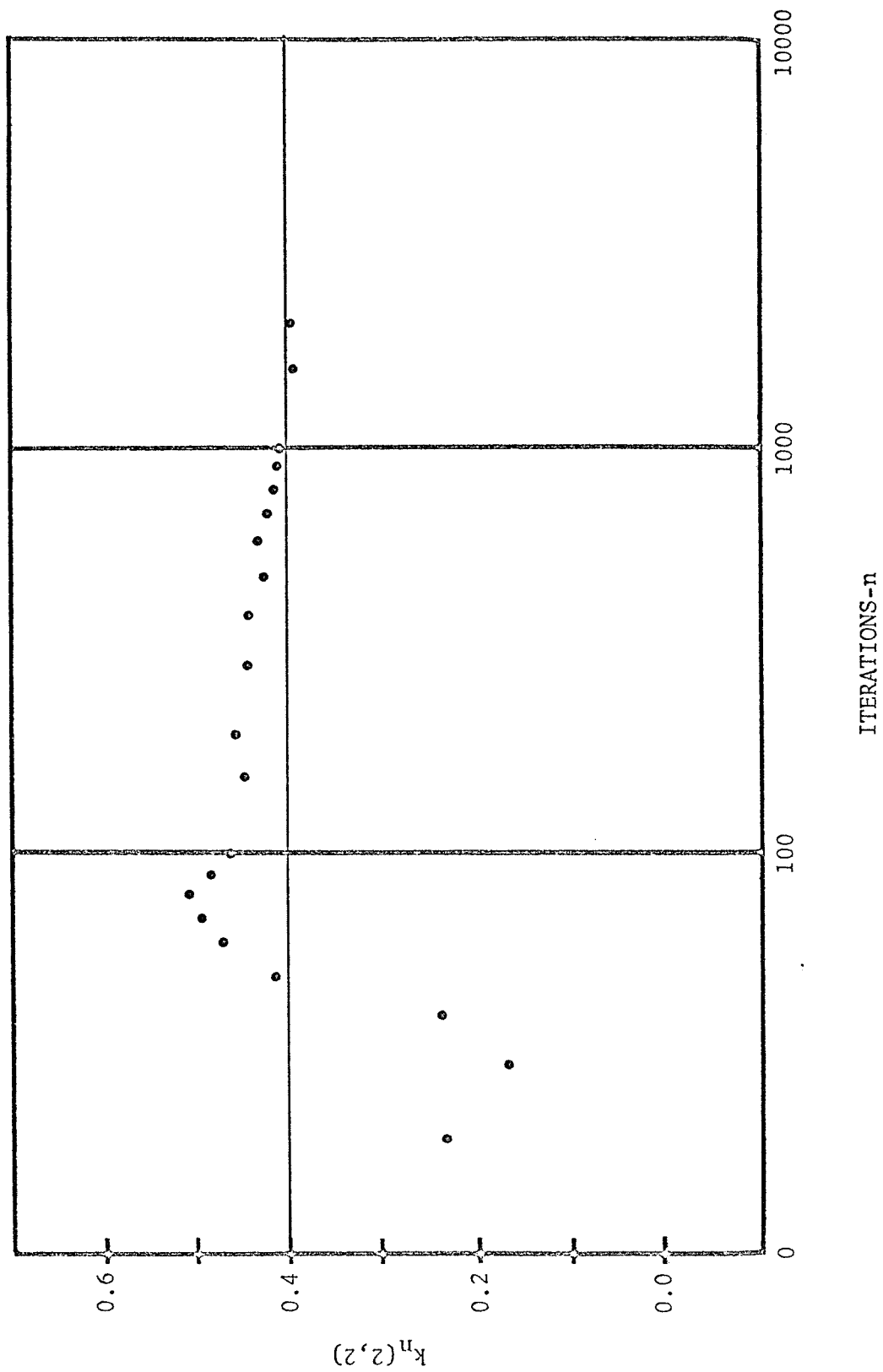


Fig. 6.4 Time Evolution of k_n for Example 6.3: Symmetry Not Forced (continued)

$$Q = \begin{bmatrix} 2.54 & 1.46 & 1.14 \\ 1.46 & 2.73 & 1.25 \\ 1.14 & 1.25 & 1.70 \end{bmatrix}$$

$$R = \begin{bmatrix} 1.11 & 1.83 & -.865 \\ 1.83 & 7.25 & 0.30 \\ -.864 & 0.30 & 2.00 \end{bmatrix}$$

The system model specified above was simulated for Gaussian, triangular and uniformly distributed noise processes $w(n)$ and $v(n)$. For Gaussian processes and with $K_0 = [0.0]$ the time evolution of the adaptive algorithm is shown in Fig. 6.5. One thousand iterations are now required for convergence. The results after 3,000 iterations are summarized below.

$$\hat{K} = \begin{bmatrix} 0.861 & -.162 & 0.269 \\ 0.256 & 0.220 & 0.241 \\ 0.281 & -.047 & 0.570 \end{bmatrix},$$

$$K_{opt} = \begin{bmatrix} 0.871 & -.160 & 0.260 \\ 0.267 & 0.227 & 0.234 \\ 0.271 & -.049 & 0.567 \end{bmatrix}$$

$$\hat{\Sigma} = \begin{bmatrix} 2.93 & 2.13 & 1.02 \\ 2.13 & 4.26 & 1.55 \\ 1.02 & 1.55 & 2.49 \end{bmatrix},$$

$$\Sigma = \begin{bmatrix} 3.0 & 2.0 & 1.0 \\ 2.0 & 4.0 & 1.5 \\ 1.0 & 1.5 & 2.5 \end{bmatrix}$$

$$\hat{\Gamma} = \begin{bmatrix} .478 & .561 & -.269 \\ .561 & 2.27 & .295 \\ -.269 & .295 & .868 \end{bmatrix},$$

$$\Gamma = \begin{bmatrix} .458 & .511 & -.282 \\ .511 & 2.21 & .309 \\ -.282 & .307 & .885 \end{bmatrix}$$

$$\hat{R} = \begin{bmatrix} 1.24 & 1.82 & -.915 \\ 1.82 & 7.03 & .270 \\ -.915 & .270 & 2.00 \end{bmatrix},$$

$$R = \begin{bmatrix} 1.11 & 1.83 & -.864 \\ 1.83 & 7.25 & .300 \\ -.864 & .300 & 2.00 \end{bmatrix}$$

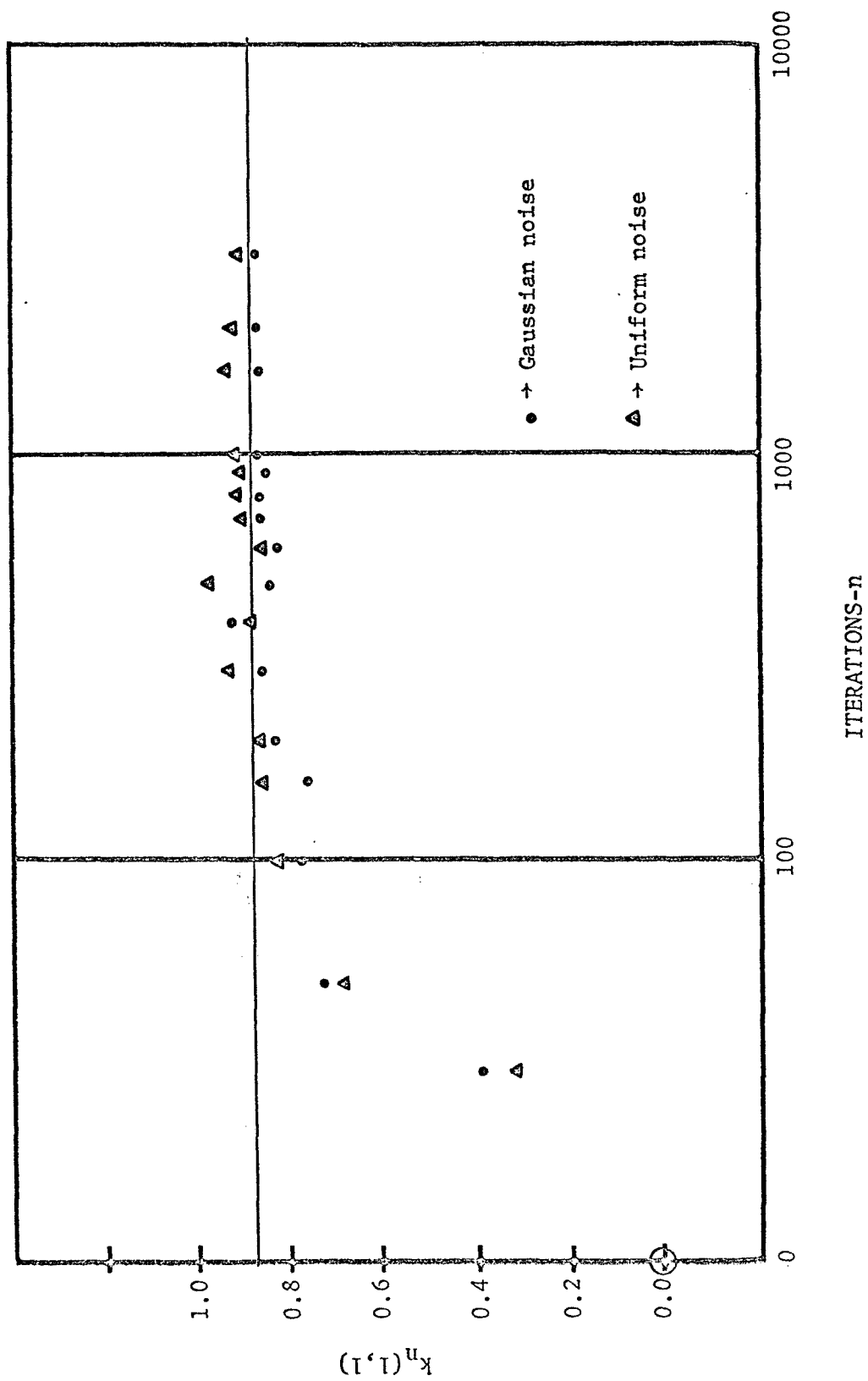


Fig. 6.5 Time Evolution of K_n for Example 6.4

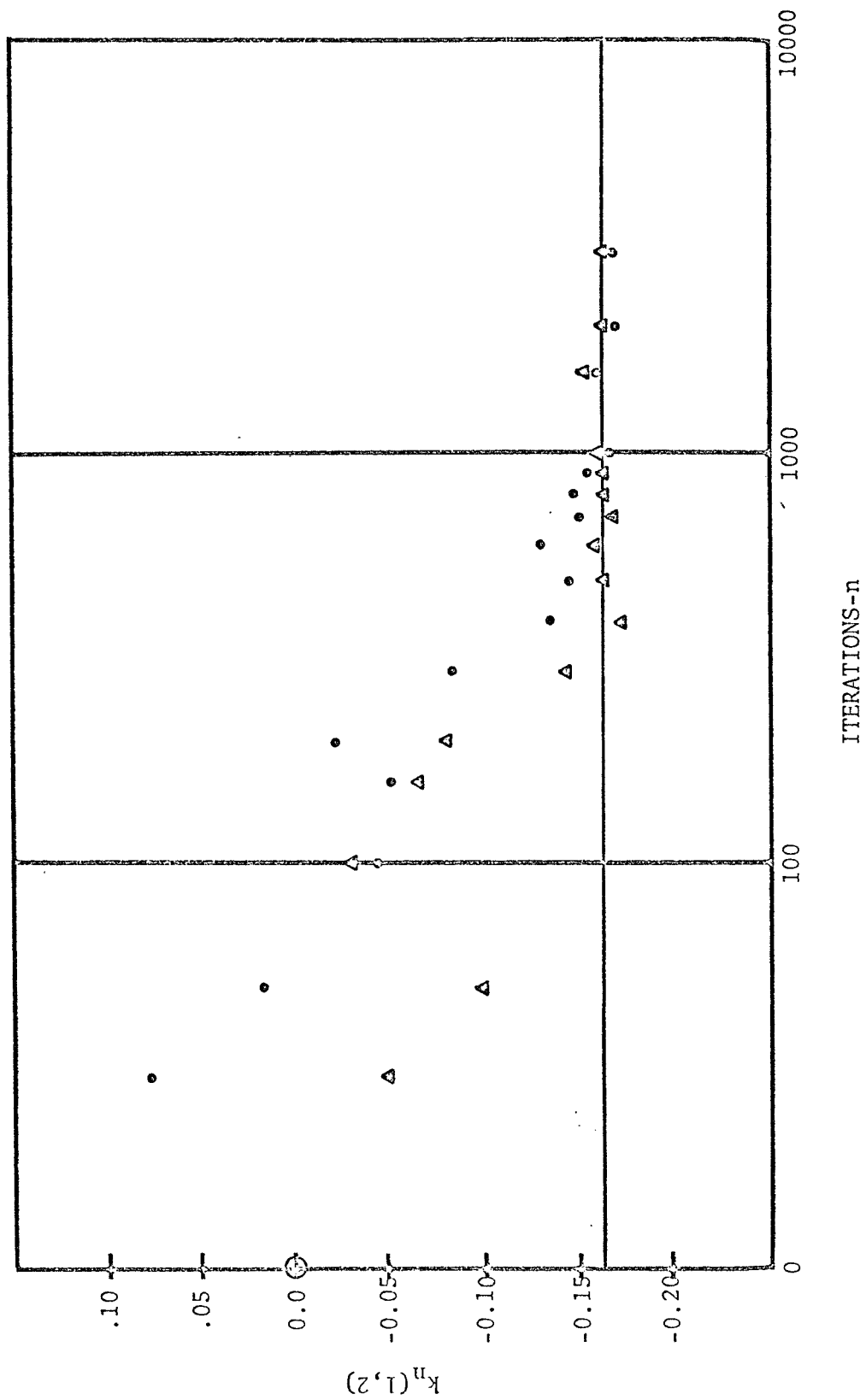


Fig. 6.5 Time Evolution of K_n for Example 6.4 (continued)

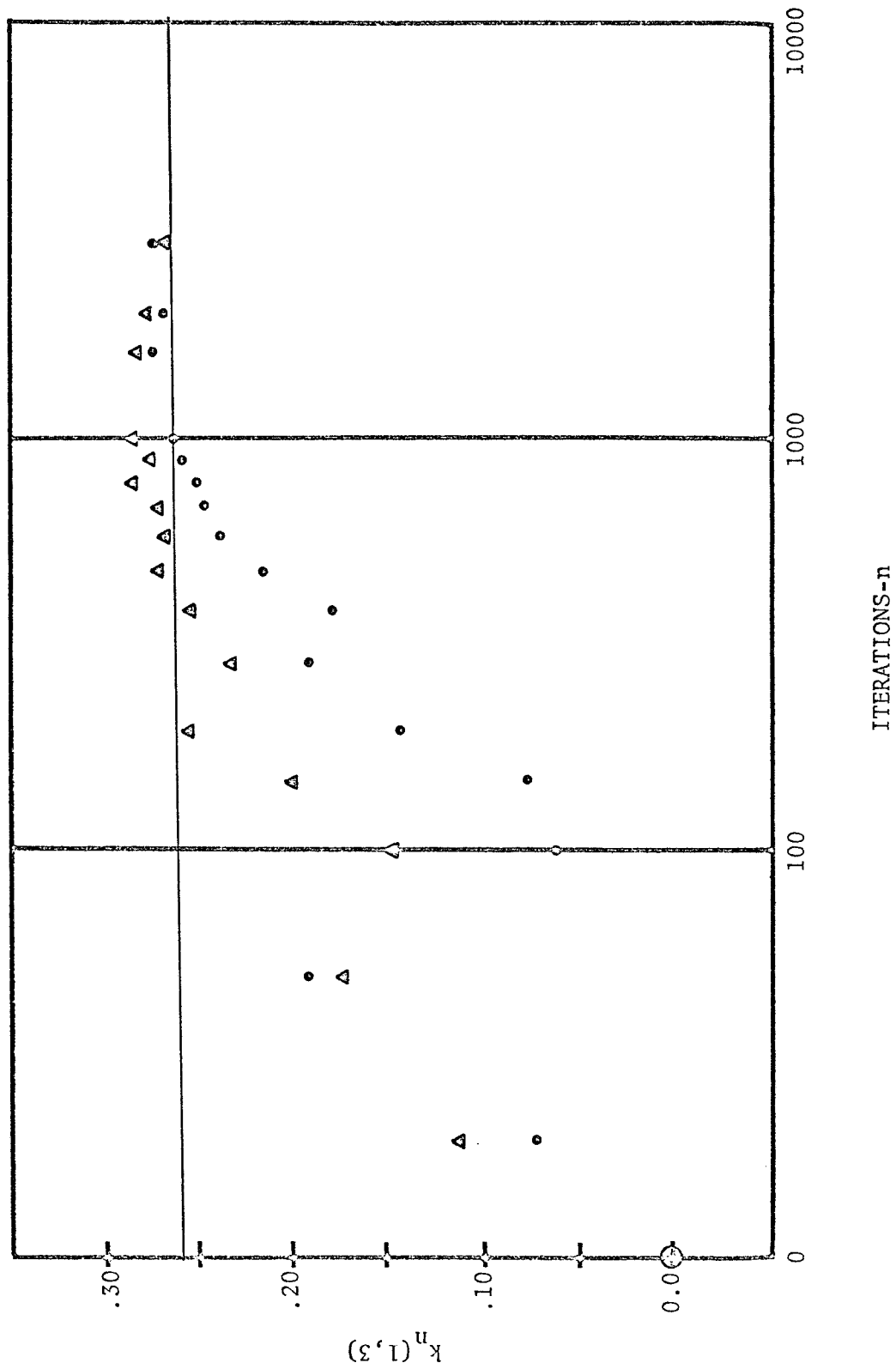


Fig. 6.5 Time Evolution of K_n for Example 6.4 (continued)

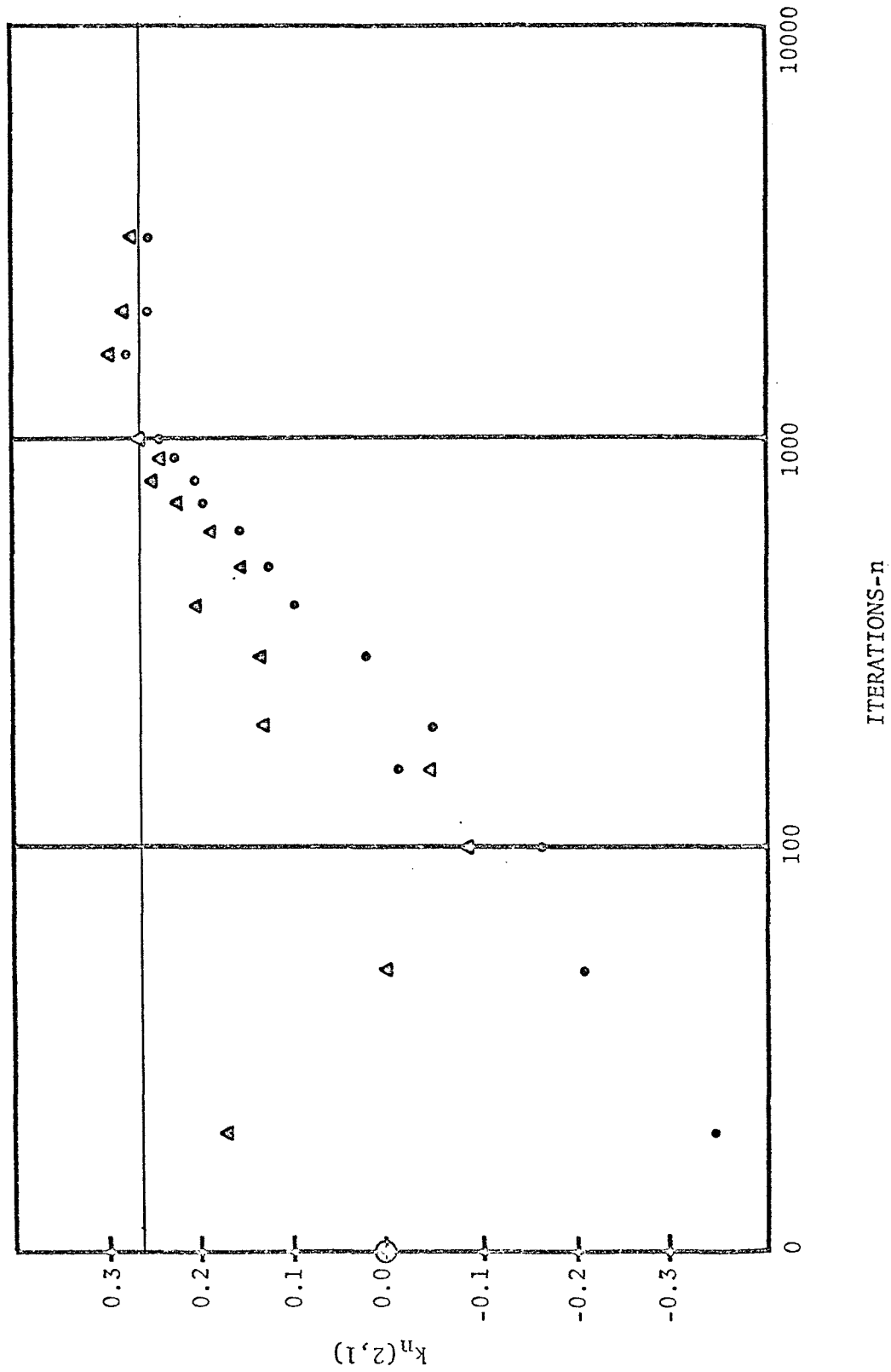
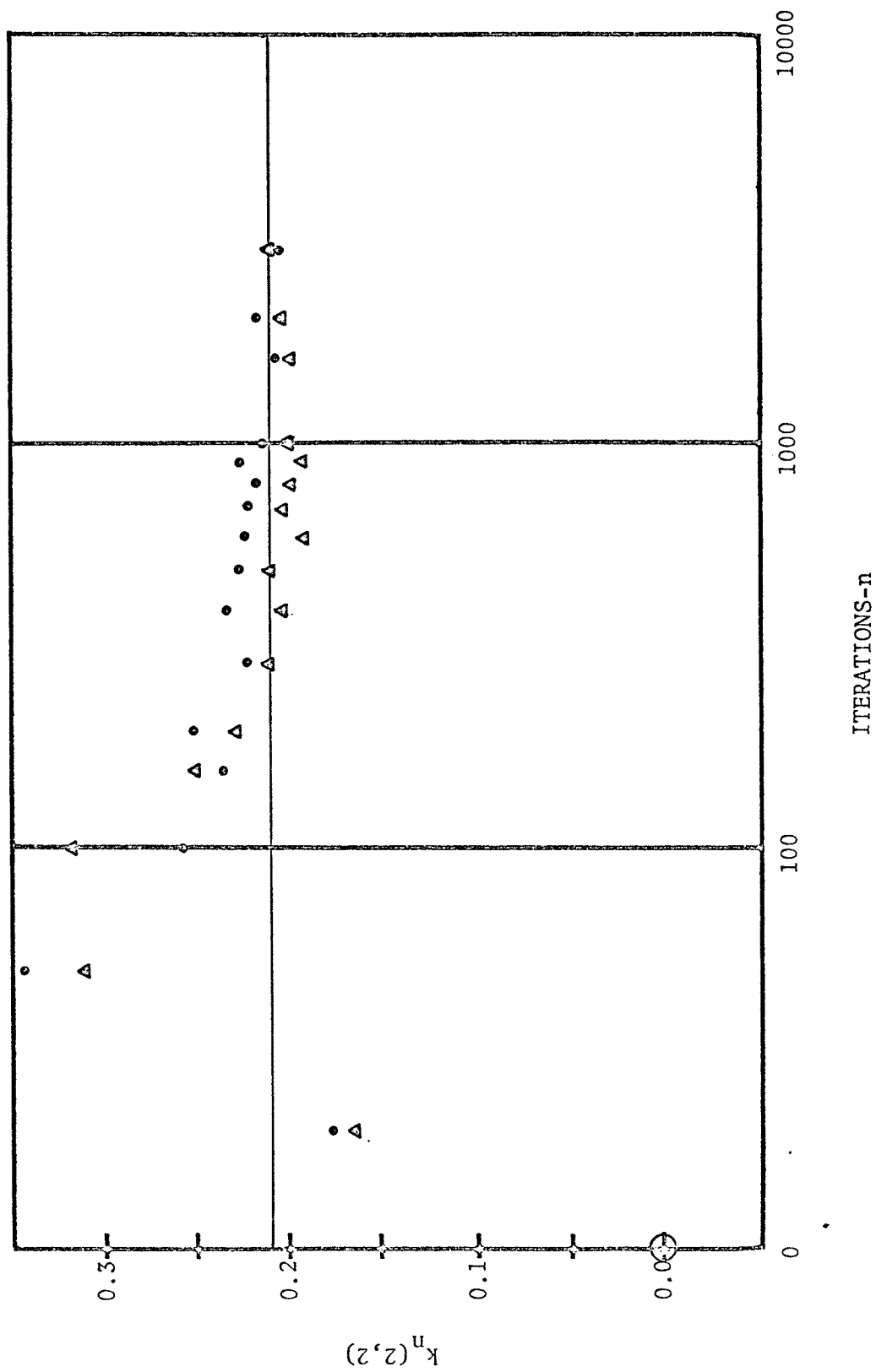


Fig. 6.5 Time Evolution of K for Example 6.4 (continued)

Fig. 6.5 Time Evolution of K_n for Example 6.4 (continued)

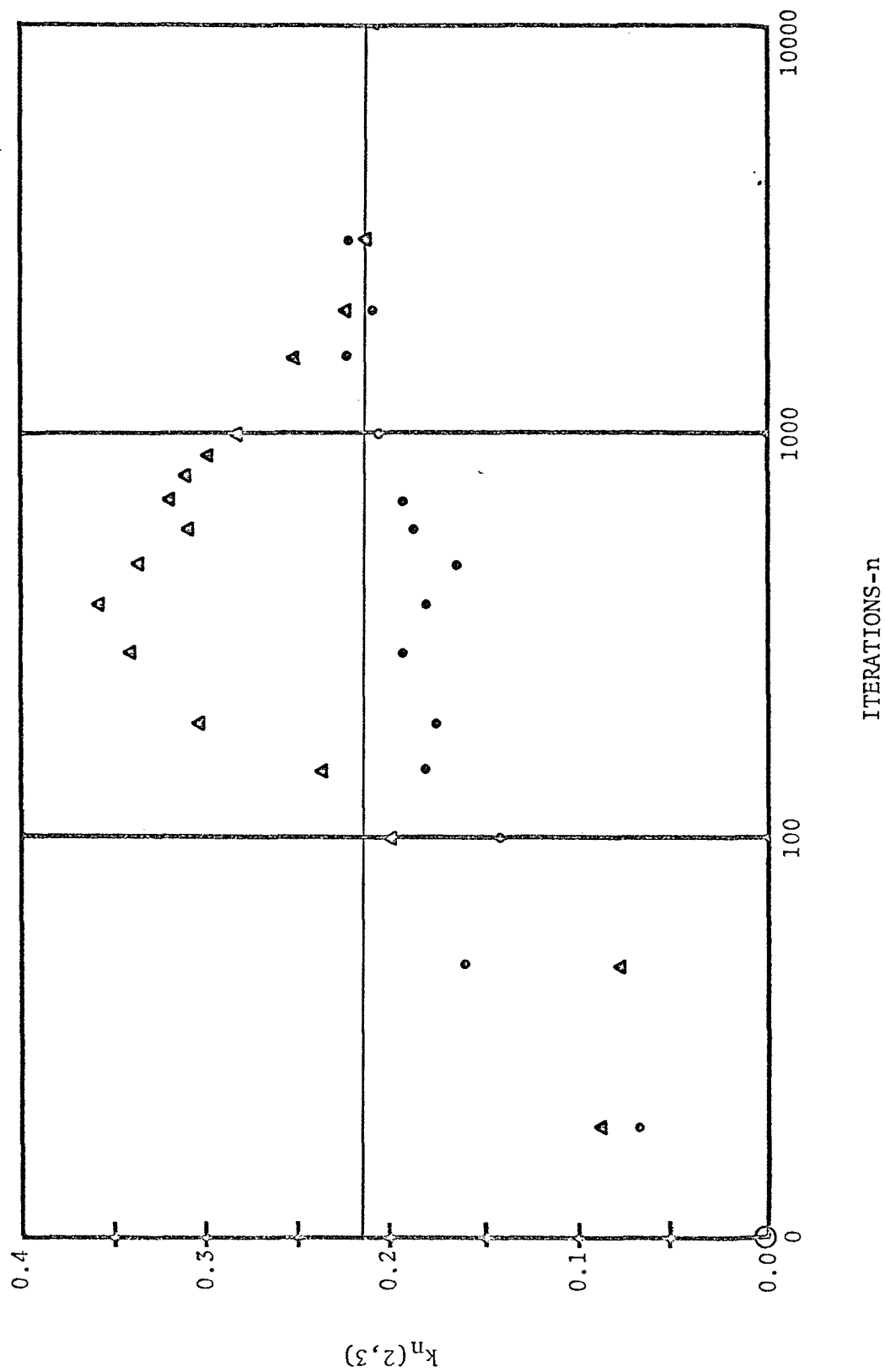
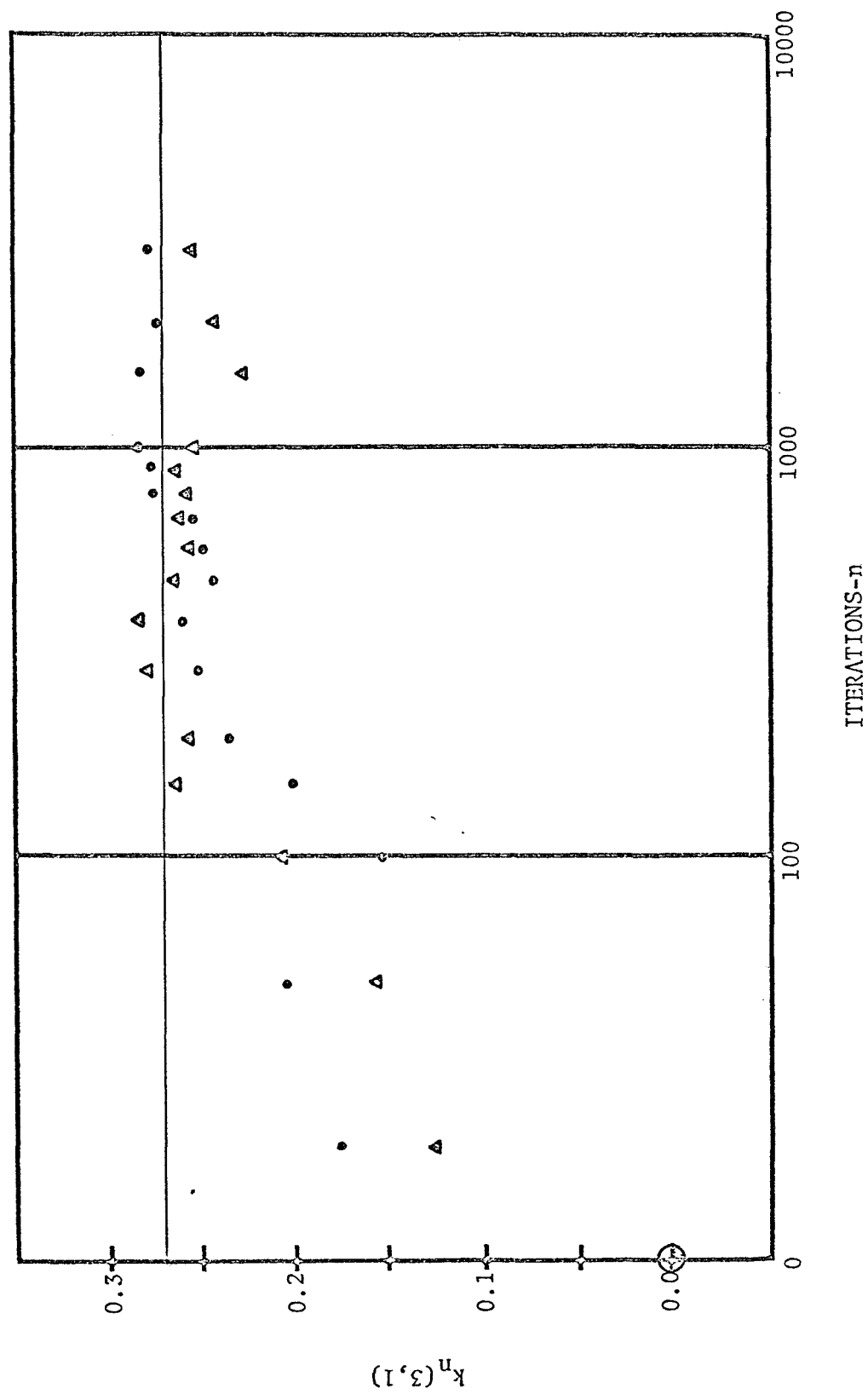


Fig. 6.5 Time Evolution of K_n for Example 6.4 (continued)

Fig. 6.5 Time Evolution of K_n for Example 6.4 (continued)

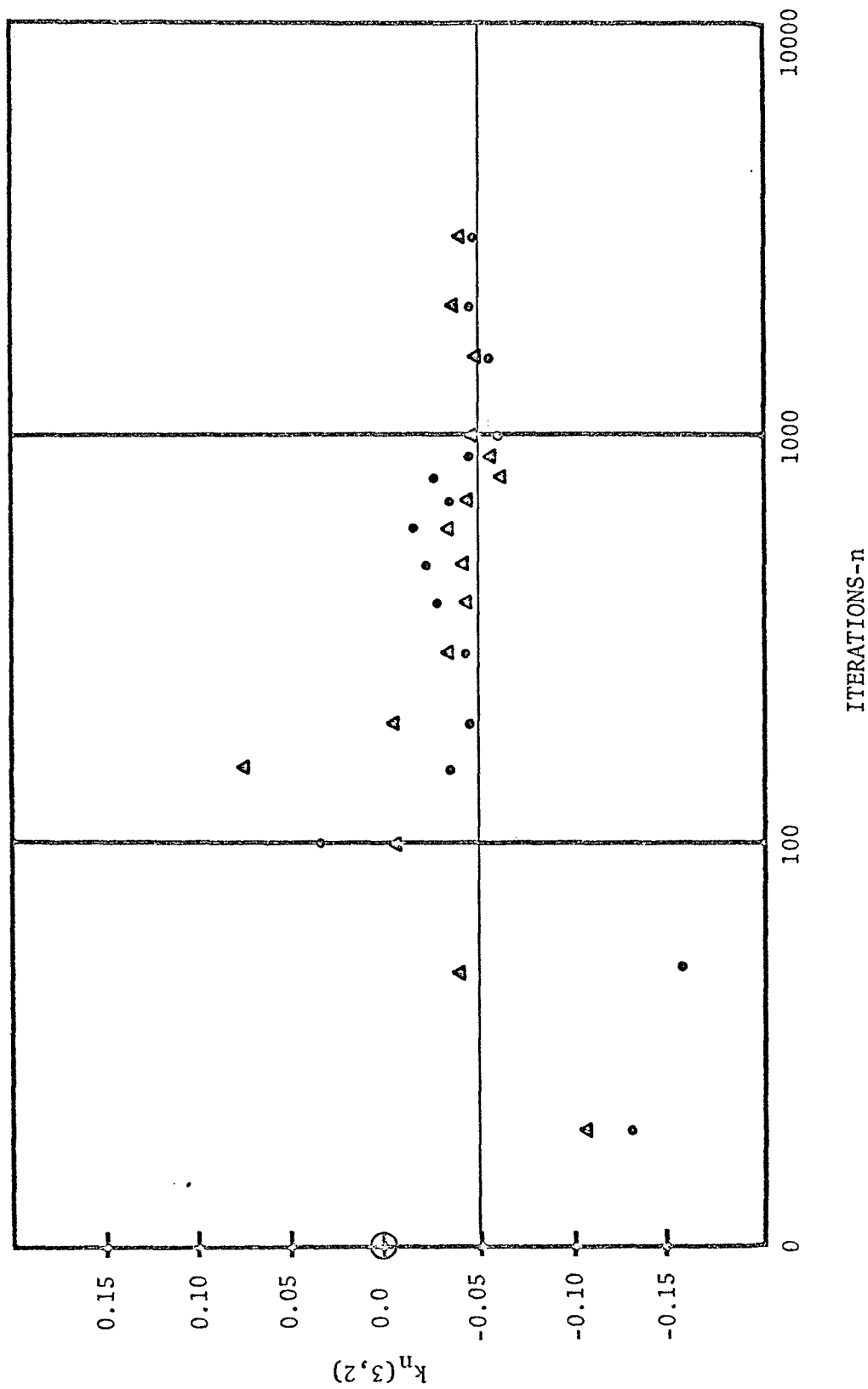


Fig. 6.5 Time Evolution of K_n for Example 6.4 (continued)

$$\hat{Q} = \begin{bmatrix} 2.40 & 1.55 & 1.15 \\ 1.55 & 2.95 & 1.30 \\ 1.15 & 1.30 & 1.70 \end{bmatrix}, \quad Q = \begin{bmatrix} 2.54 & 1.46 & 1.14 \\ 1.46 & 2.73 & 1.25 \\ 1.14 & 1.25 & 1.70 \end{bmatrix}$$

As can be seen the estimates are in excellent agreement with the optimal values. The mean square error (MSE) of the filtering process is defined as

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)][\underline{x}(n) - \hat{\underline{x}}(n)]^T\}$$

Since the error covariance matrix Γ is defined by Eq. (2.21) as

$$E\{[\underline{x}(n) - \hat{\underline{x}}(n)][\underline{x}(n) - \hat{\underline{x}}(n)]^T\}$$

Therefore, the mean square error is the sum of the major diagonal elements of Γ . Performing this calculation, the optimum and estimated mean square errors are

$$\hat{MSE} = 3.66, \quad MSE = 3.55$$

The actual M.S.E. for the filter given by $\hat{K}$ can be determined by calculating the actual Γ using Eq. (5.8). It is

$$(MSE)_{act} = 3.62$$

This indicates that the "learned" filter $\hat{K}$ produces a MSE only 1.97% larger than the optimal filter. The corresponding one-step prediction mean square errors (PMSE) are

$$\hat{PMSE} = 9.68, \quad (PMSE)_{act} = 9.55, \quad PMSE = 9.50$$

To demonstrate the insensitivity of the algorithm of Eq. (4.43) to the assumed initial value of the filter K_0 was chosen to be the identity matrix. After 1,000 iterations,

$$\hat{K} = \begin{bmatrix} 0.883 & -.171 & .254 \\ 0.253 & 0.239 & .258 \\ 0.260 & -.052 & .555 \end{bmatrix}$$

Theoretically, the adaptive algorithm is independent of the amplitude distributions of $w(n)$ and $v(n)$. Consequently, the final value of the filter should be the same. To test this experimentally, uniform and triangular distributions were used. The simulation results for 3,000 iterations with $K = [0.0]$ are

$$(\hat{K})_{\text{triangle}} = \begin{bmatrix} .881 & -.155 & .275 \\ .261 & .233 & .240 \\ .263 & -.047 & .554 \end{bmatrix}$$

$$(\hat{K})_{\text{uniform}} = \begin{bmatrix} .876 & -.163 & .267 \\ .270 & .223 & .230 \\ .249 & -.041 & .549 \end{bmatrix}$$

The case for uniformly distributed noise is also plotted in Fig. 6.5 for comparison with the Gaussian distributed case.

6.2.4 Fourth Order Dynamics

EXAMPLE 6.5:

Given:

$$\Phi = \begin{bmatrix} 0.87 & 0.0 & 0.0 & 0.0 \\ 0.0 & 0.69 & 0.0 & 0.0 \\ 0.0 & 0.0 & 0.94 & 0.0 \\ 0.0 & 0.0 & 0.0 & 0.75 \end{bmatrix}, \quad H = I$$

$$Q = \begin{bmatrix} 1.60 & 0.32 & 0.88 & 0.85 \\ 0.32 & 1.09 & 0.78 & 0.96 \\ 0.88 & 0.78 & 2.38 & 1.43 \\ 0.85 & 0.96 & 1.43 & 2.15 \end{bmatrix}$$

$$R = \begin{bmatrix} 6.89 & 3.42 & -.55 & 1.41 \\ 3.42 & 6.17 & -1.10 & -0.34 \\ -.55 & -1.10 & 1.21 & 0.33 \\ 1.41 & -0.34 & 0.33 & 1.25 \end{bmatrix}$$

$$Q = \begin{bmatrix} 1.60 & 0.32 & 0.88 & 0.85 \\ 0.32 & 1.09 & 0.78 & 0.96 \\ 0.88 & 0.78 & 2.38 & 1.43 \\ 0.85 & 0.96 & 1.43 & 2.15 \end{bmatrix}$$

To initialize the simulation process for this system, K_0 was again chosen to be the null matrix. The time evolution of the filter adaptation for Gaussian distributions is shown in Fig. 6.6. As before, in addition to learning K_{opt} , Σ , Γ , R , and Q were also estimated. The results after 4,000 iterations are given below

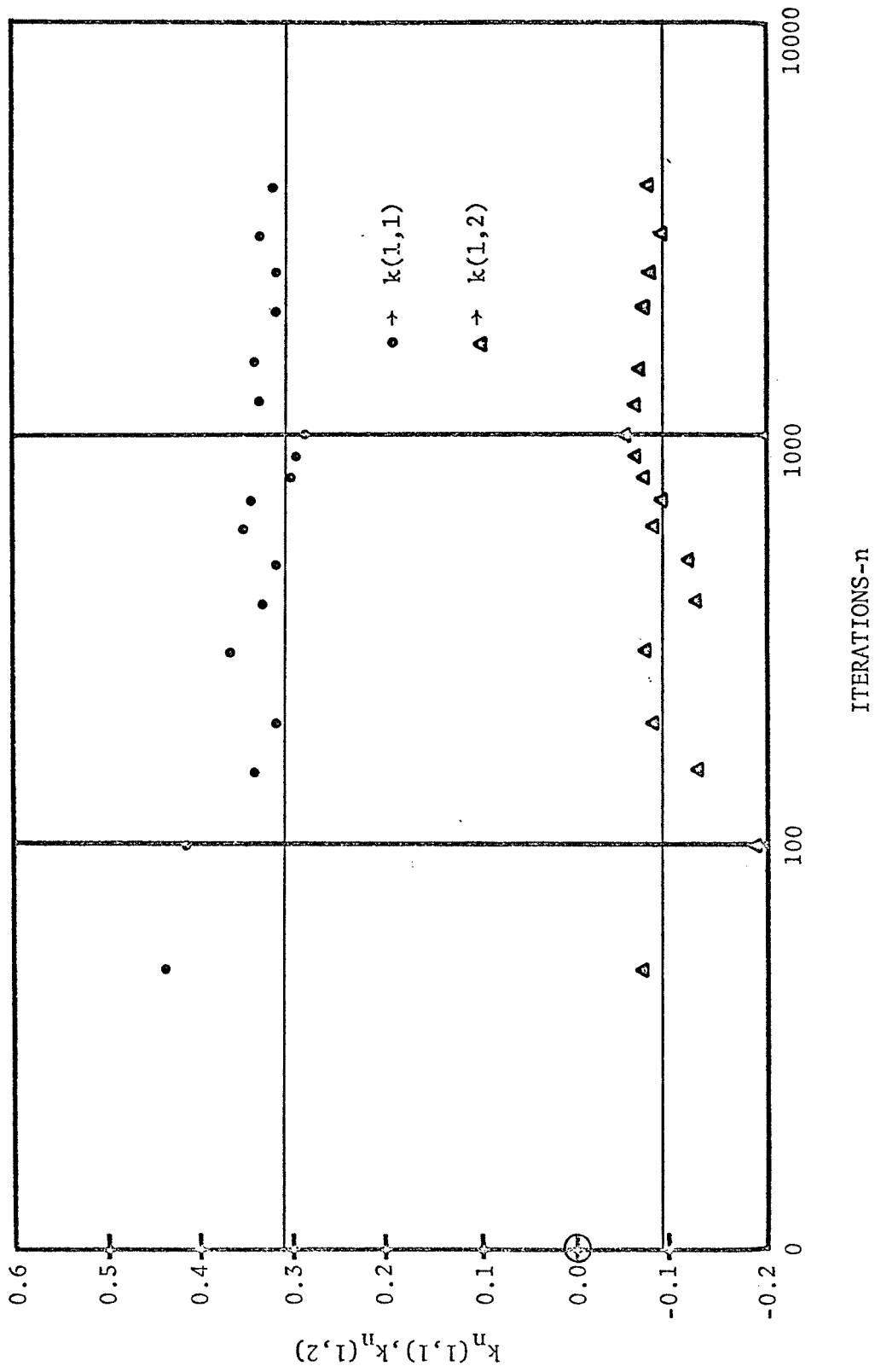


Fig. 6.6 Time Evolution of K_n for Example 6.5

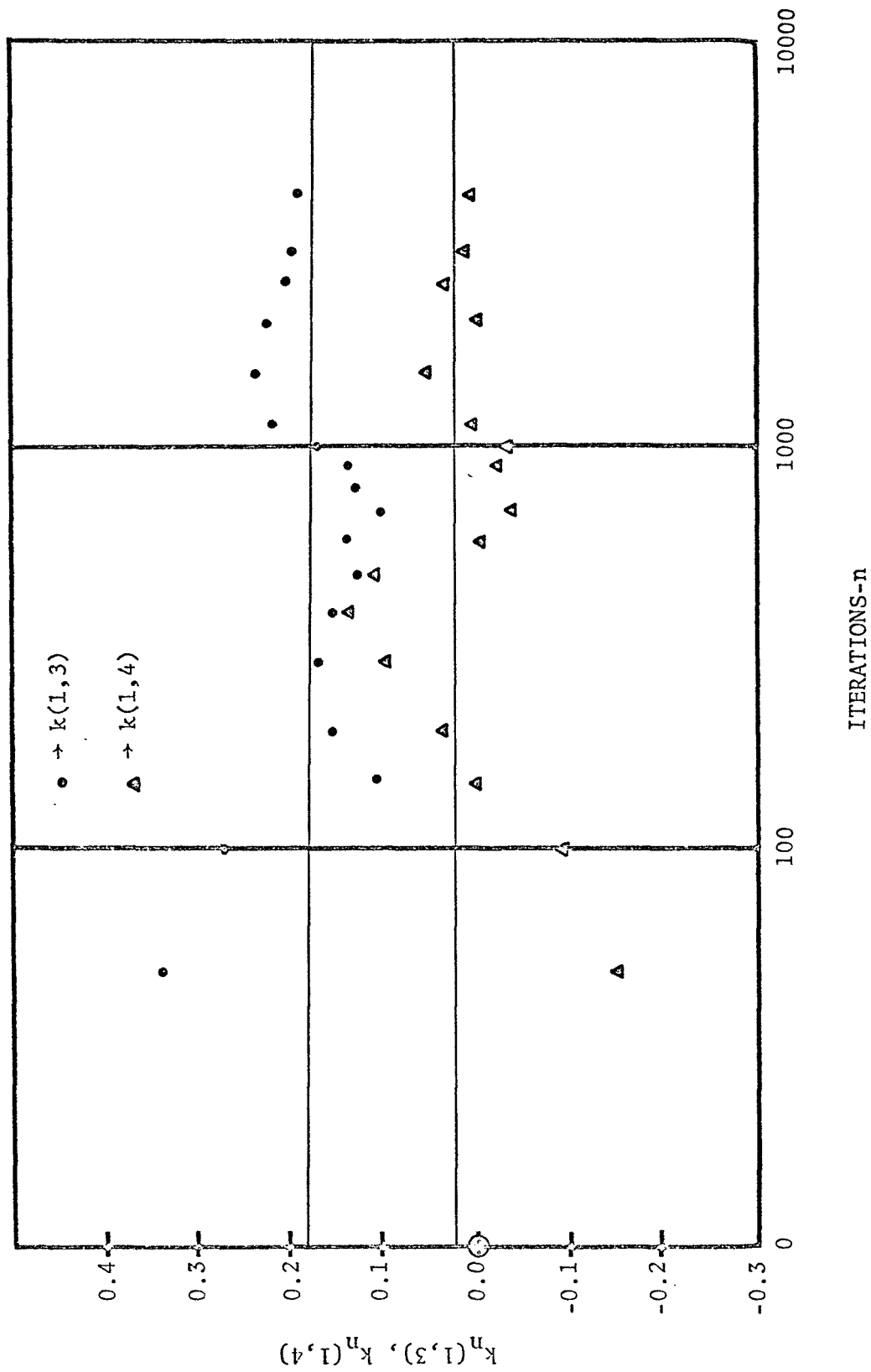


Fig. 6.6 Time Evolution of K_n for Example 6.5 (continued)

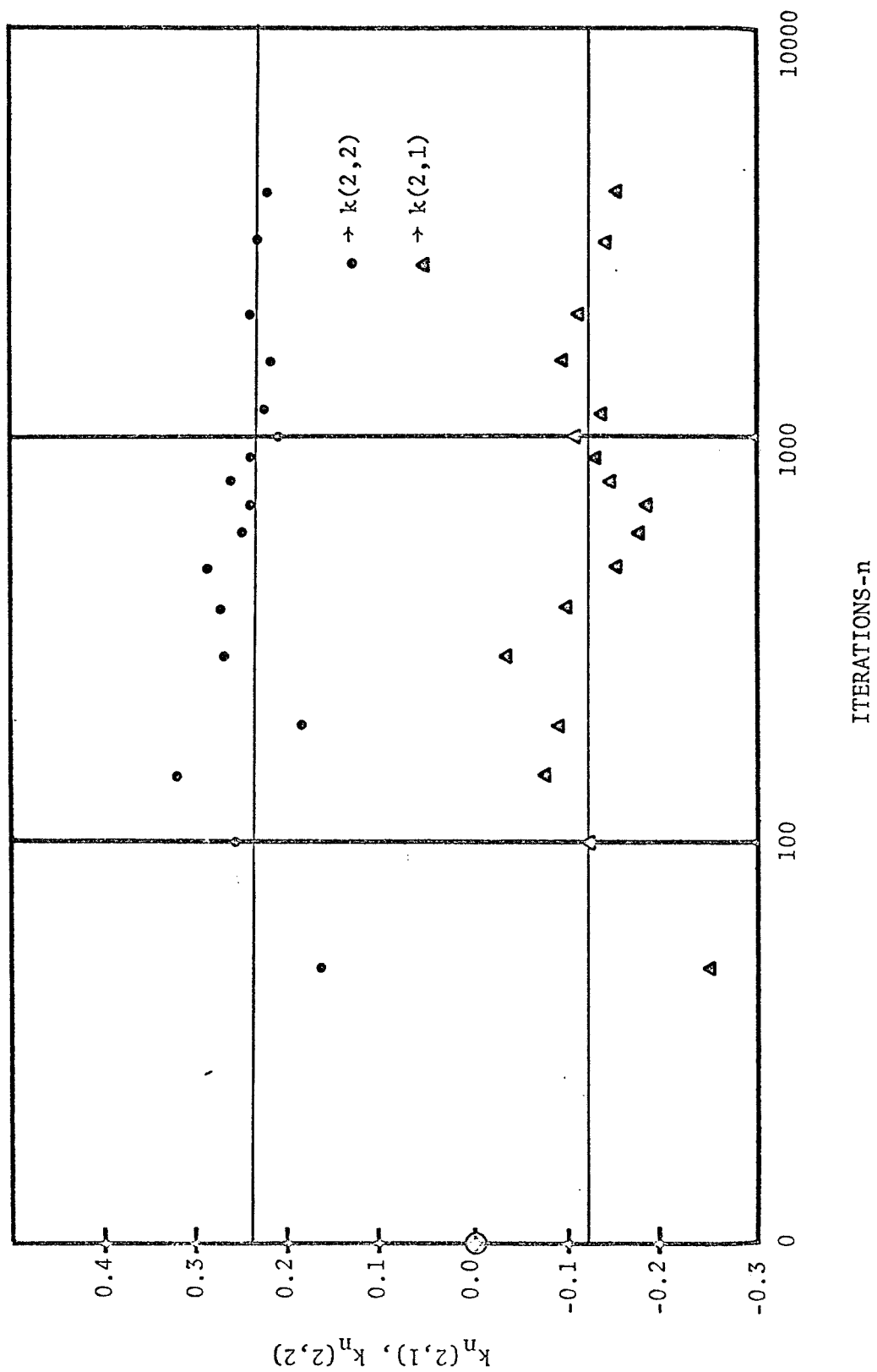


Fig. 6.6 Time Evolution of K_n for Example 6.5 (continued)

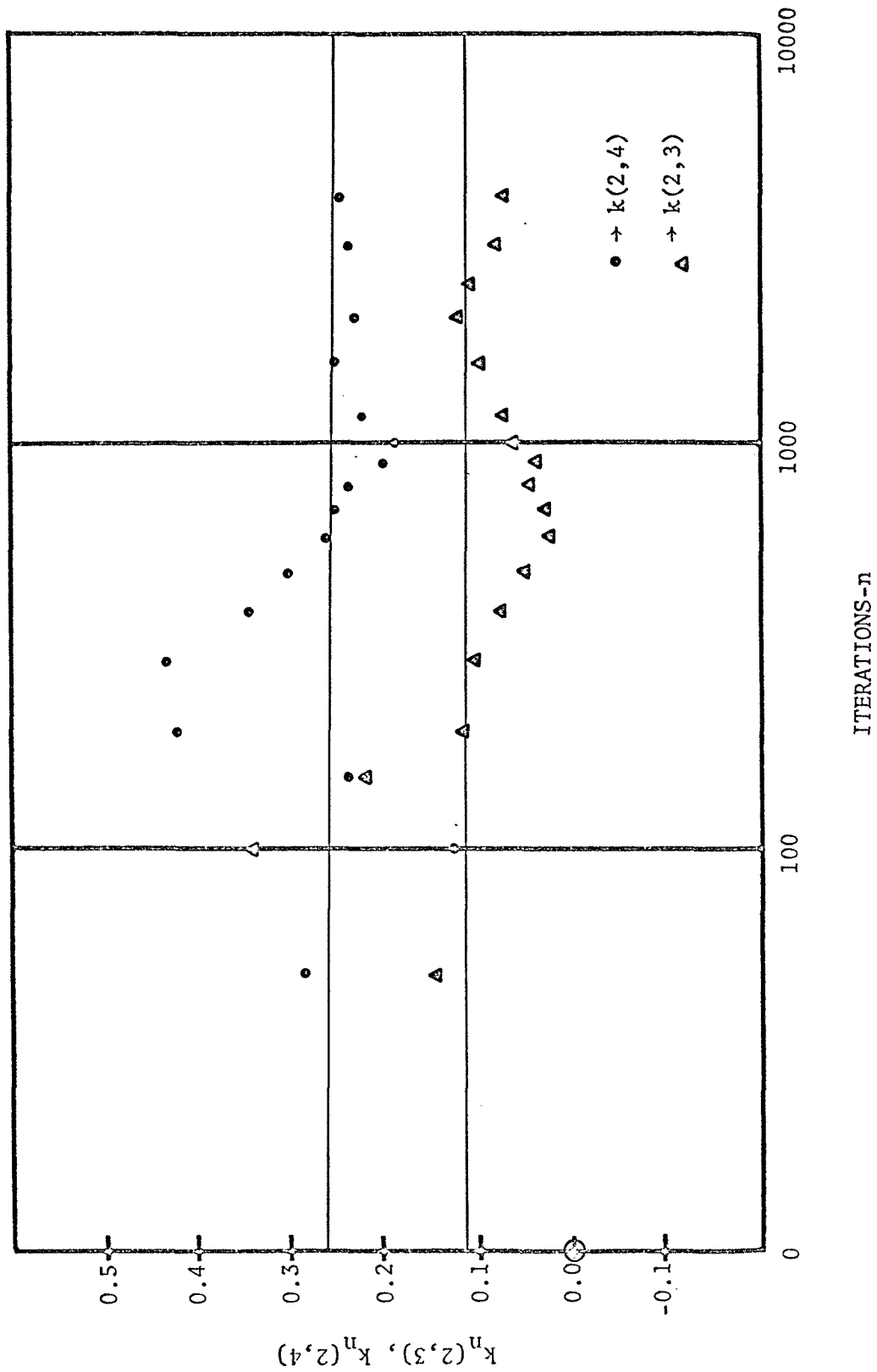


Fig. 6.6 Time Evolution of K_n for Example 6.5 (continued)

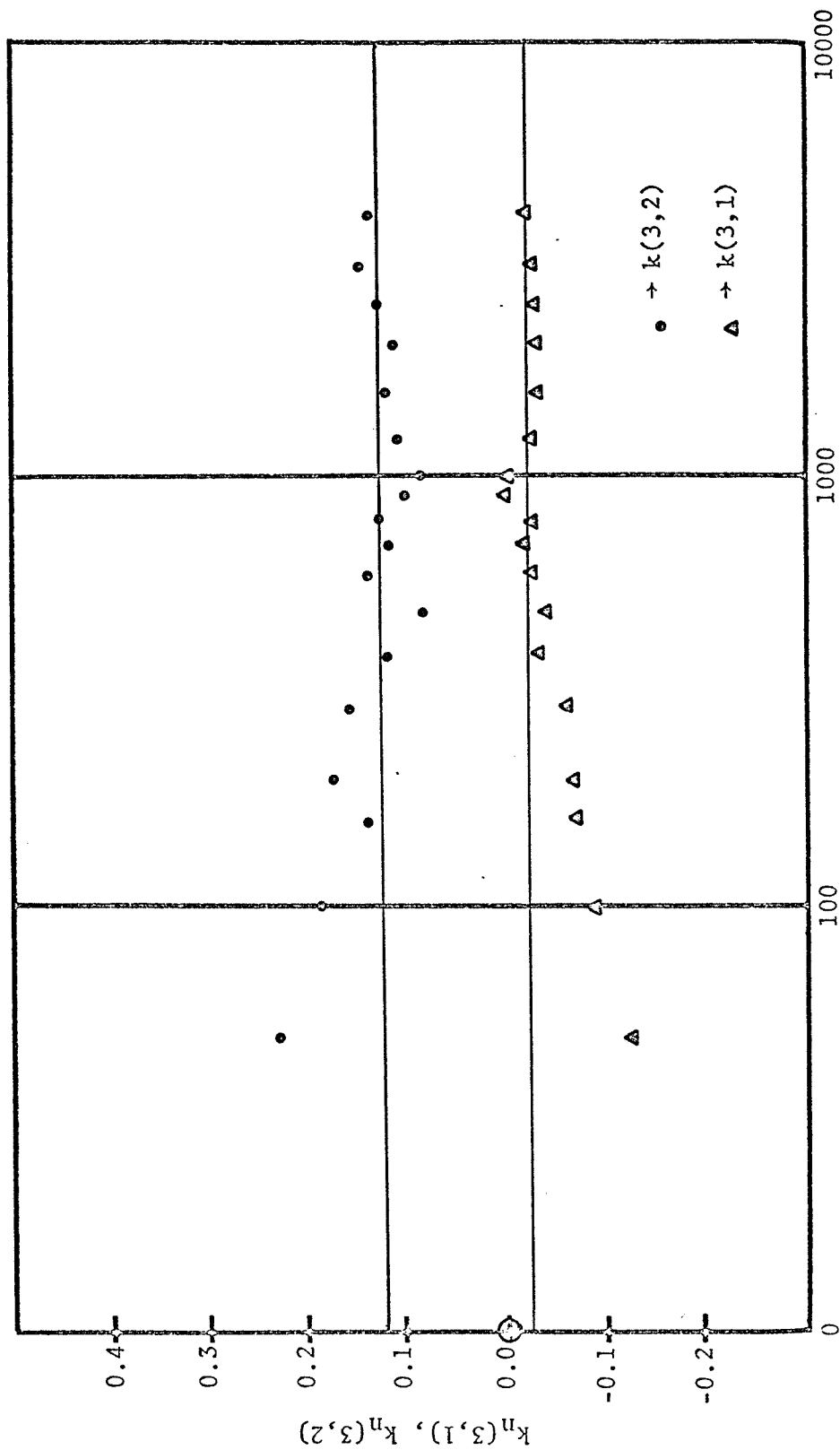


Fig. 6.6 Time Evolution of k_n for Example 6.5 (continued)

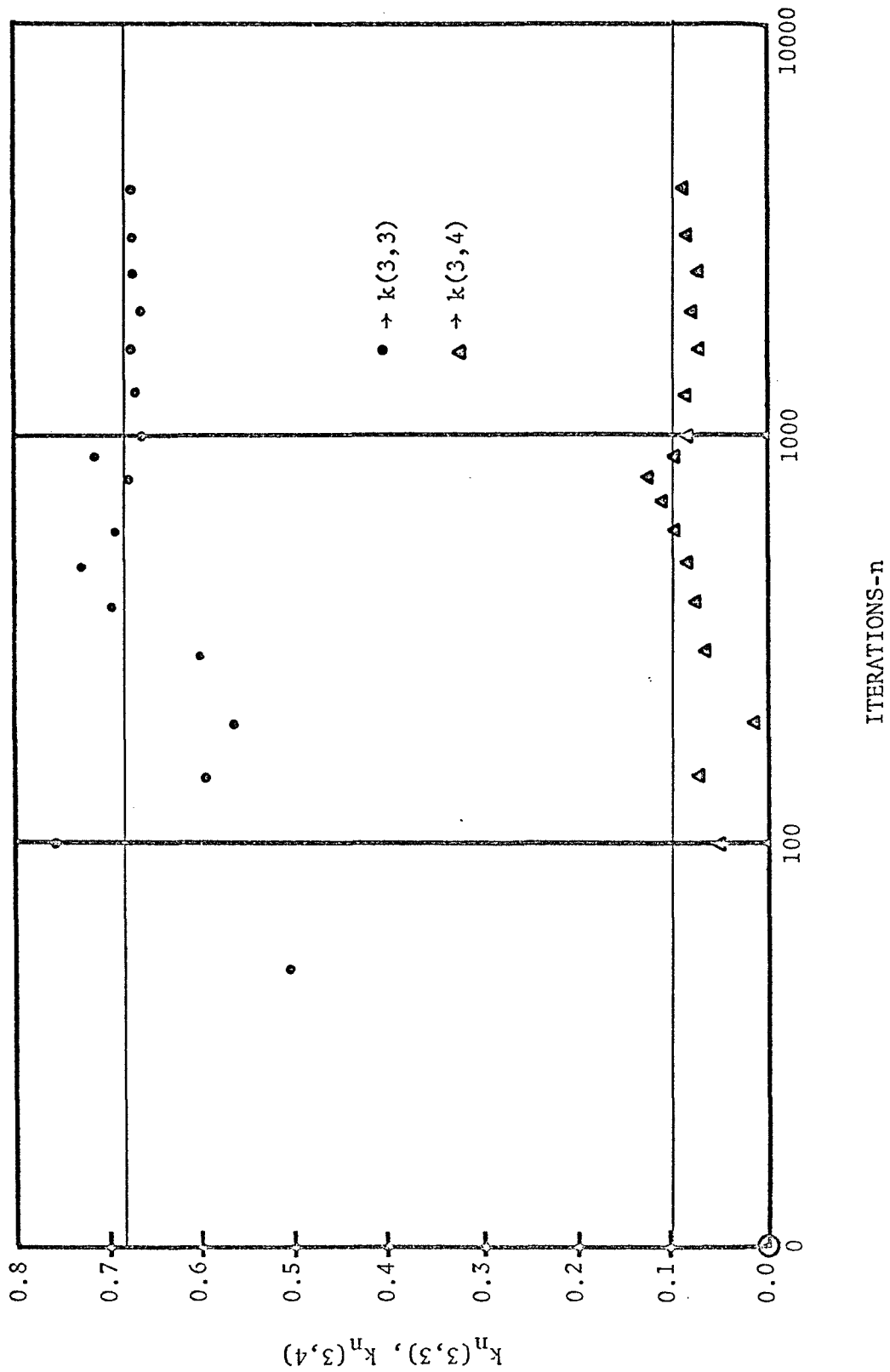


Fig. 6.6 Time Evolution of K_n for Example 6.5 (continued)

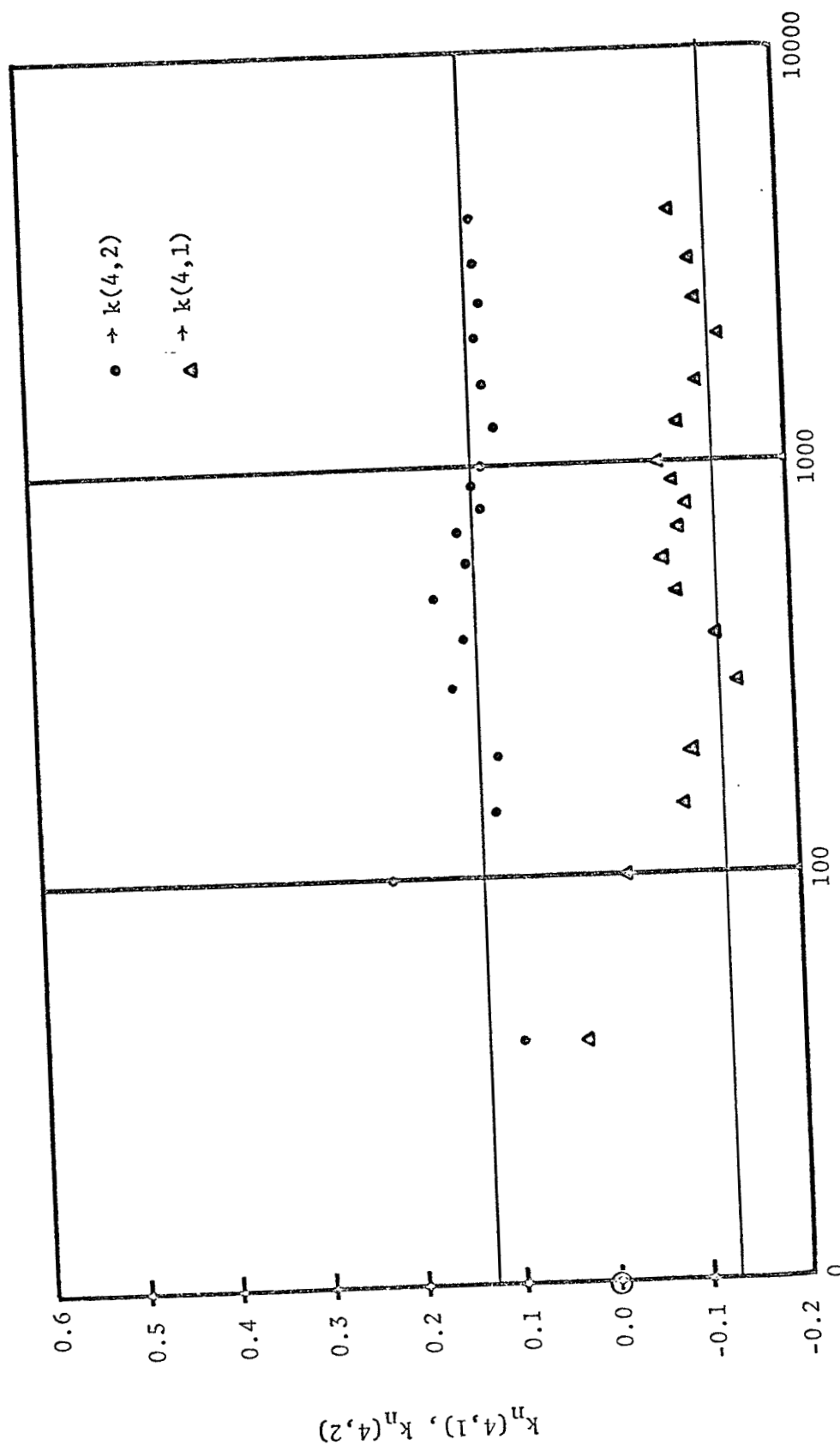


Fig. 6.6 Time Evolution of K_n for Example 6.5 (continued)

ITERATIONS-n

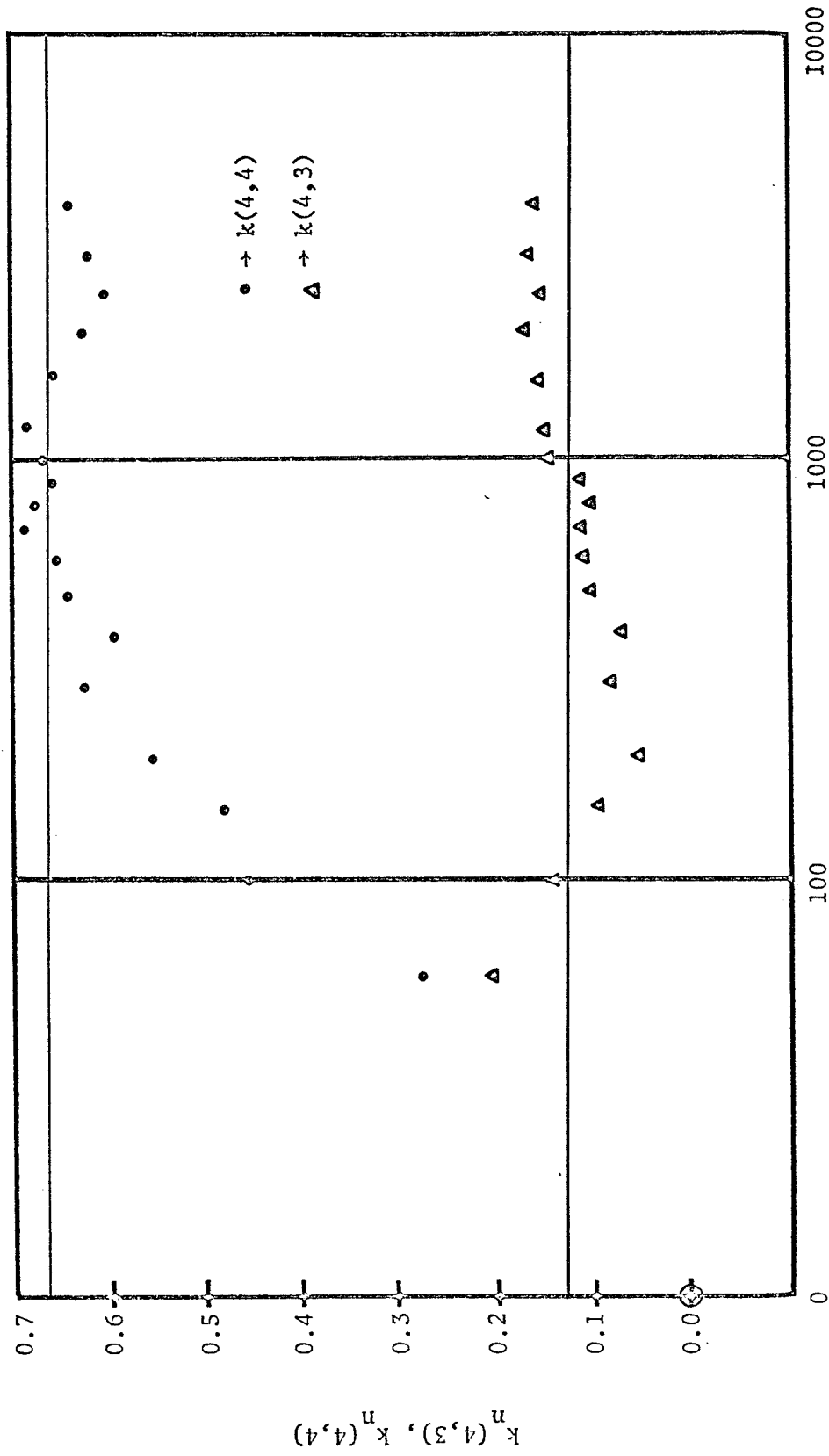


Fig. 6.6 Time Evolution of K_n for Example 6.5 (continued)

ITERATIONS-n

$$\hat{K} = \begin{bmatrix} .327 & -.086 & .208 & .030 \\ -.130 & .234 & .095 & .254 \\ -.004 & .130 & .679 & .096 \\ -.091 & .137 & .163 & .628 \end{bmatrix}, K_{\text{opt}} = \begin{bmatrix} .323 & -.095 & .181 & .037 \\ -.118 & .239 & .116 & .259 \\ -.006 & .123 & .686 & .102 \\ -.116 & .138 & .137 & .667 \end{bmatrix}$$

$$\hat{\Sigma} = \begin{bmatrix} 3.18 & .744 & 1.06 & 1.31 \\ .744 & 1.71 & .854 & 1.29 \\ 1.06 & .854 & 3.02 & 1.46 \\ 1.31 & 1.29 & 1.46 & 2.51 \end{bmatrix}, \Sigma = \begin{bmatrix} 3.0 & .5 & 1.0 & 1.2 \\ .5 & 1.5 & .75 & 1.0 \\ 1.0 & .75 & 3.0 & 1.5 \\ 1.2 & 1.0 & 1.5 & 2.5 \end{bmatrix}$$

$$\hat{\Gamma} = \begin{bmatrix} 1.94 & .399 & .154 & .560 \\ .399 & .835 & -.068 & .118 \\ .154 & -.068 & .700 & .060 \\ .560 & .118 & .060 & .579 \end{bmatrix}, \Gamma = \begin{bmatrix} 1.85 & .308 & .148 & .542 \\ .308 & .854 & -.048 & .081 \\ .148 & -.048 & .701 & .100 \\ .542 & .081 & .100 & .628 \end{bmatrix}$$

$$\hat{R} = \begin{bmatrix} 6.7 & 3.3 & -.43 & 1.34 \\ 3.3 & 6.00 & -1.06 & -.524 \\ -.43 & -1.06 & 1.20 & .032 \\ 1.34 & -.524 & .032 & 1.24 \end{bmatrix}, R = \begin{bmatrix} 6.89 & 3.42 & -.55 & 1.41 \\ 3.42 & 6.17 & -1.10 & -.335 \\ -.55 & -1.10 & 1.21 & .033 \\ 1.41 & -.335 & .033 & 1.25 \end{bmatrix}$$

$$\hat{Q} = \begin{bmatrix} 1.71 & .503 & .935 & .946 \\ .503 & 1.31 & .898 & 1.23 \\ .935 & .898 & 2.39 & 1.42 \\ .946 & 1.23 & 1.42 & 2.18 \end{bmatrix}, Q = \begin{bmatrix} 1.60 & .315 & .879 & .847 \\ .315 & 1.09 & .781 & .958 \\ .879 & .781 & 2.38 & 1.43 \\ .847 & .958 & 1.43 & 2.15 \end{bmatrix}$$

These estimates are again in good agreement with the optimal values, but as the dimension of the problem increases so does the number of iterations required to attain this agreement. By summing the diagonal elements of

the appropriate error covariance matrix, the corresponding mean square errors are calculated to be

$$\hat{MSE} = 4.05 \quad MSE = 4.04 \quad (MSE)_{ACT} = 4.11$$

Therefore, the "learned" filter $\hat{K}$ produces a MSE only 1.74% larger than the optimal filter assuming perfect a priori information. This example represents the best results obtained for fourth order problems. Many other examples were simulated, and in every case the actual mean square error was less than 2.75% greater than its optimum value.

6.3 Nonlinear Dynamics

In this section the nonlinear system algorithms of Section 5.5 are investigated. The first class of nonlinearities considered satisfy a global Lipschitz condition. The second class have a unique inverse. In both cases the form of the filtering equation is given by Eq. (5.40).

6.3.1 Nonlinearities That Satisfy a Global Lipschitz Condition

Consider the class of nonlinear systems

$\{f(x): |f(x_n) - f(x_m)| \leq \phi |x_n - x_m|, \text{ for all } x_n, x_m\}$. Then the adaptive algorithm, which learns the filter that is optimal in the sense that it satisfies Eq. (5.42), is given by Eq. (5.43)

$$\hat{K}_{n+1} = \hat{K}_n + \frac{1}{n} \phi^{-1} \delta_{n+1} \delta_n^T W_{n+1} \quad (5.43)$$

Because the optimal filter is unknown in the nonlinear case, the results of this section are compared with those for a linear system with the same ϕ and noise covariances Q and R . Also from the state equation (5.38), a little reflection reveals that the lower bound for the

prediction error covariance Σ is given by Q . Therefore, the plant perturbation covariance Q can be used as a reference for evaluating the performance of the adaptively "learned" filter.

EXAMPLE 6.6: Saturation Nonlinearity

Given: $\phi = 0.5$, $Q = R = 1$

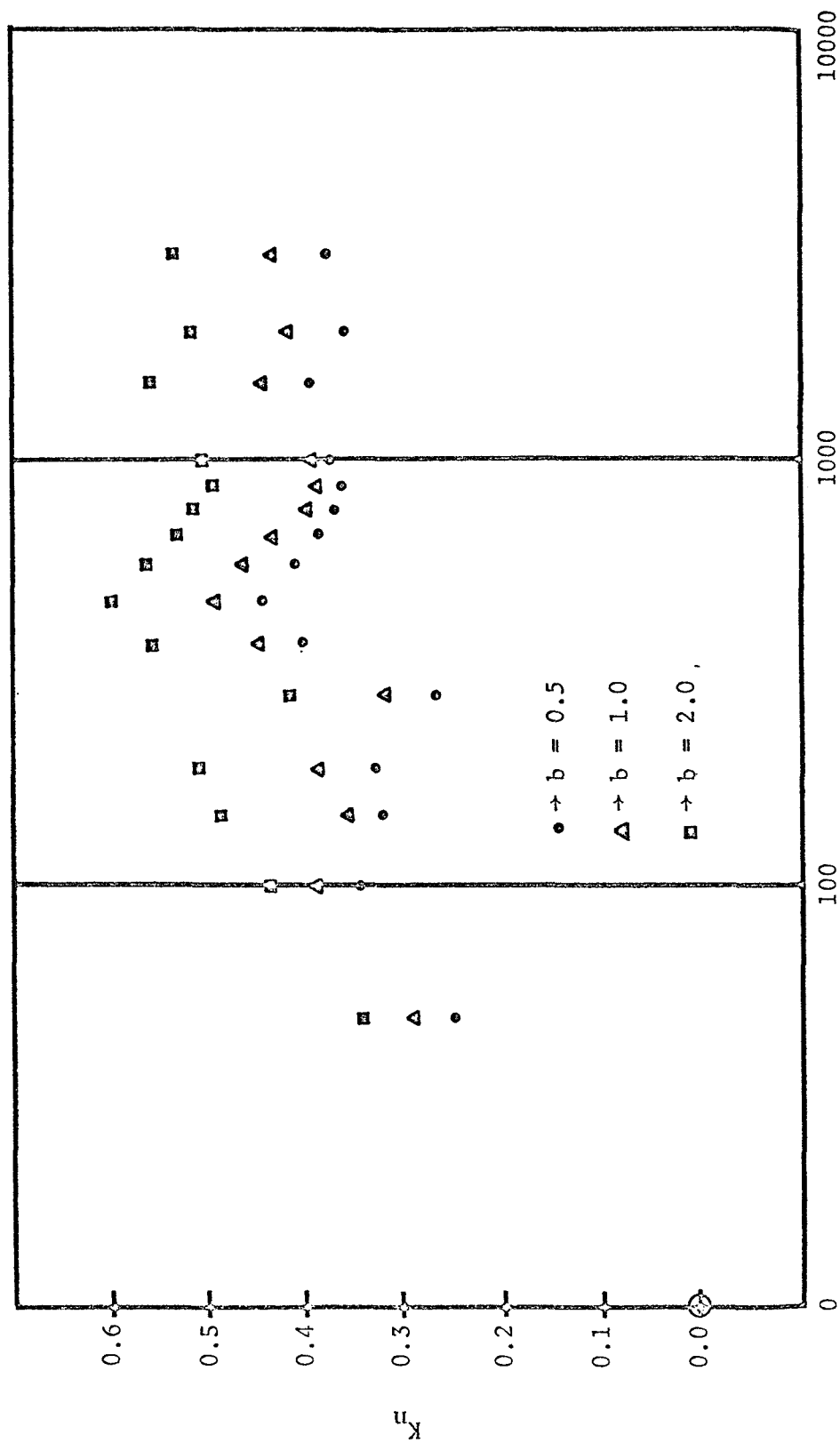
$$f(x) = \begin{cases} \phi b & \text{for } x > b \\ \phi x & \text{for } |x| \leq b \\ -\phi b & \text{for } x < -b \end{cases} \quad (5.44)$$

For this problem it is obvious that if b is increased without bound, then the "learned" filter should approach the value obtained for the linear system. This does occur and is illustrated in Fig. 6.7 for $b = 0.5$, 1.0 , and 2.0 respectively. For these three systems the sample mean square error Γ and the sample one-step predicted mean square error were calculated for both the filtered and non-filtered estimates of the state $x(n)$. By non-filtered, it is meant that the measurements are used as if they are noise free. This corresponds to an implied filter magnitude of 1.0 . The non-filtered one-step prediction of $\hat{x}(n+1)$ given $z(n)$ is then determined by $f(z(n))$. The data for 3,000 iterations are summarized in Table 6.3.

TABLE 6.3

COMPARISON OF THE ADAPTIVELY FILTERED AND NON-FILTERED RESULTS
FOR EXAMPLE 6.6

Quantity	b = 0.5			b = 1.0			b = 2.0		
	K	Γ	Σ	K	Γ	Σ	K	Γ	Σ
Non-Filtered	1.0	.855	1.03	1.0	.72	1.12	1.0	.69	1.18
Adaptively Filtered	.387	.529	1.01	.43	.537	1.05	.526	.538	1.07
Optimal Linear	.531	.531	1.13	.531	.531	1.13	.531	.531	1.13

Fig. 6.7 Time Evolution of K_n for Example 6.6

ITERATIONS-n

Observe from Table 6.3 that as the saturation threshold b is increased, the value of the adaptively "learned" filter approaches the K_{opt} for the corresponding linear system, i.e., $b = \infty$. Note also that for $b = 0.5$, the mean square error 0.855 associated with the non-filtered estimate is 61.8% larger than that for the adaptively filtered estimate. The computed Γ and Σ for the adaptive filter are also smaller than those for the optimal linear filter case. This is rather interesting. And it is not unexpected since when the system is in saturation, it is not difficult to accurately predict its next state much of the time. The other point worthy of mention is the lower bound on Σ . Since $Q = 1.0$, the absolute lower bound on Σ is 1.0 which is very close to the calculated Σ .

EXAMPLE 6.7: Saturation Nonlinearity

Given: $\phi = 0.5$, $b = 1.0$

The purpose of this example is to examine the effect of changing Q and R while keeping their ratio constant. In a linear system such a change does not alter K_{opt} . In a nonlinear system this is no longer true. Figure 6.8 illustrates this effect for $R = 4Q$ with $Q = 1$ and $Q = 4$. As is to be expected, the larger Q creates a smaller "learned" filter. This follows since the system is operating in the saturated region more often. The experimental calculations after 5000 iterations are listed in Table 6.4. The increase in the non-filtered estimation mean square error Γ over that for adaptive filtering is 343%. By comparing the other quantities in Table 6, it seems that the non-linear adaptive algorithm developed in Section 5.5 does indeed "work" well.

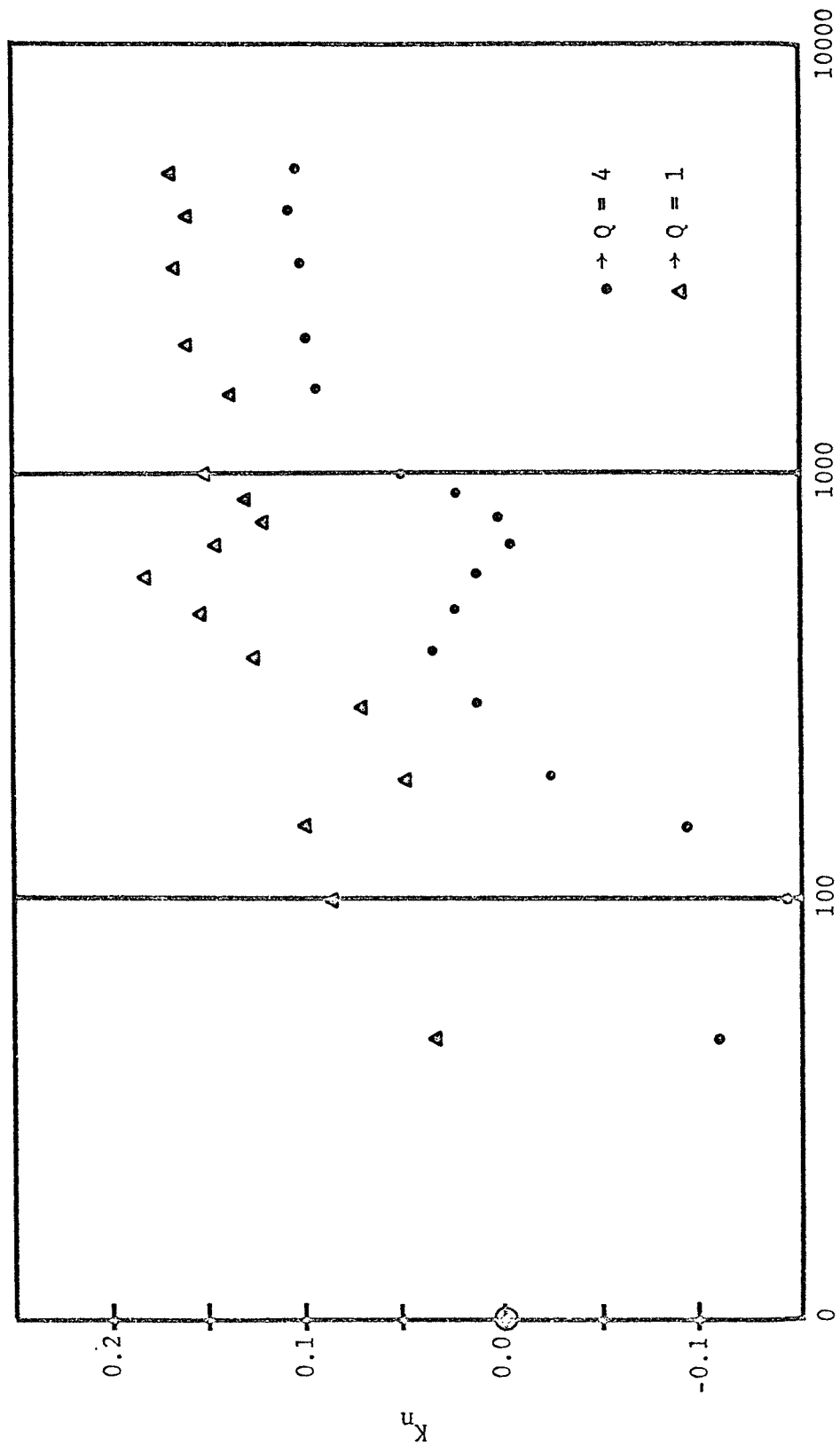


Fig. 6.8 Time Evolution of K_n for Example 6.7

ITERATIONS--n

TABLE 6.4

COMPARISON OF THE ADAPTIVELY FILTERED AND NON-FILTERED RESULTS
FOR EXAMPLE 6.7

Quantity	Q = 1.0			Q = 4.0		
	K	Γ	Σ	K	Γ	Σ
Non-Filtered	1.0	3.76	1.06	1.0	20.0	4.23
Adaptively Filtered	.165	0.85	1.02	.103	4.52	4.13
Optimal Linear	.236	0.94	1.24	.236	3.76	4.94

EXAMPLE 6.8: Dead Zone Nonlinearity

Given: $\phi = 0.5$, $b = 0.25$

$$f(x) = \begin{cases} \phi(x-b) & \text{for } x > b \\ 0 & \text{for } |x| \leq b \\ -\phi(x-b) & \text{for } x < -b \end{cases} \quad (5.45)$$

Algorithm (5.43) was programmed for the dead zone nonlinearity of Eq. (5.45). The results for $Q = R = 1.0$ are shown in Fig. 6.9. The sample mean square errors Γ and Σ after 3000 iterations are 0.534 and 1.1 respectively. The absolute lower bound on Σ is 1.0 and 0.531 is the optimal value for Γ in a linear system. Figure 6.9 also contains a plot for $Q = R = 2.0$.

EXAMPLE 6.9: Dead Zone Nonlinearity

Given: $\phi = 0.5$, $b = 0.25$, $Q = 1.0$, $R = 4.0$

The time evolution of K_n for these parameters is illustrated in Fig. 6.10 and a comparison of the adaptively filtered and non-filtered results is given in Table 6.5.

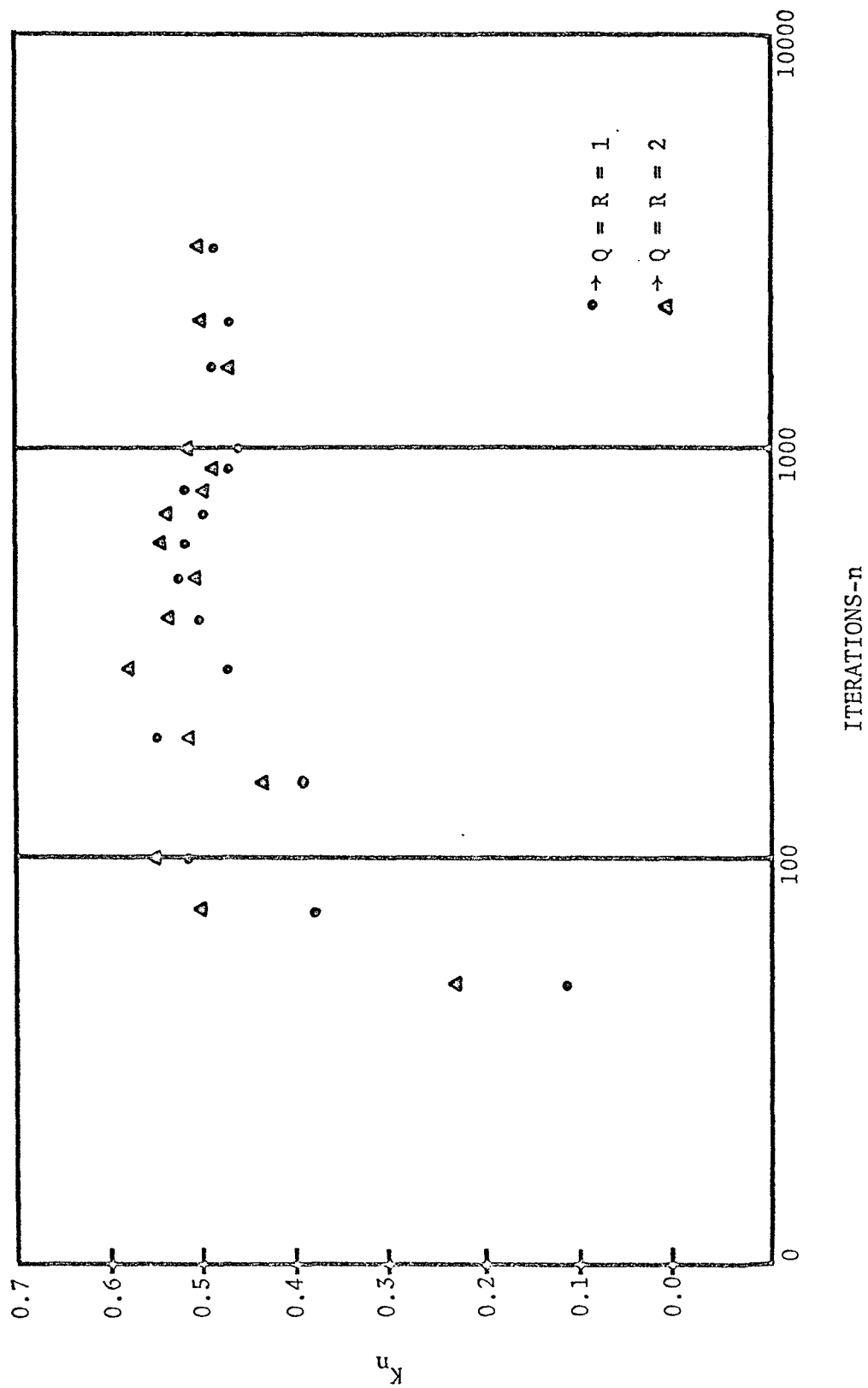


Fig. 6.9 Time Evolution of K_n for Example 6.8

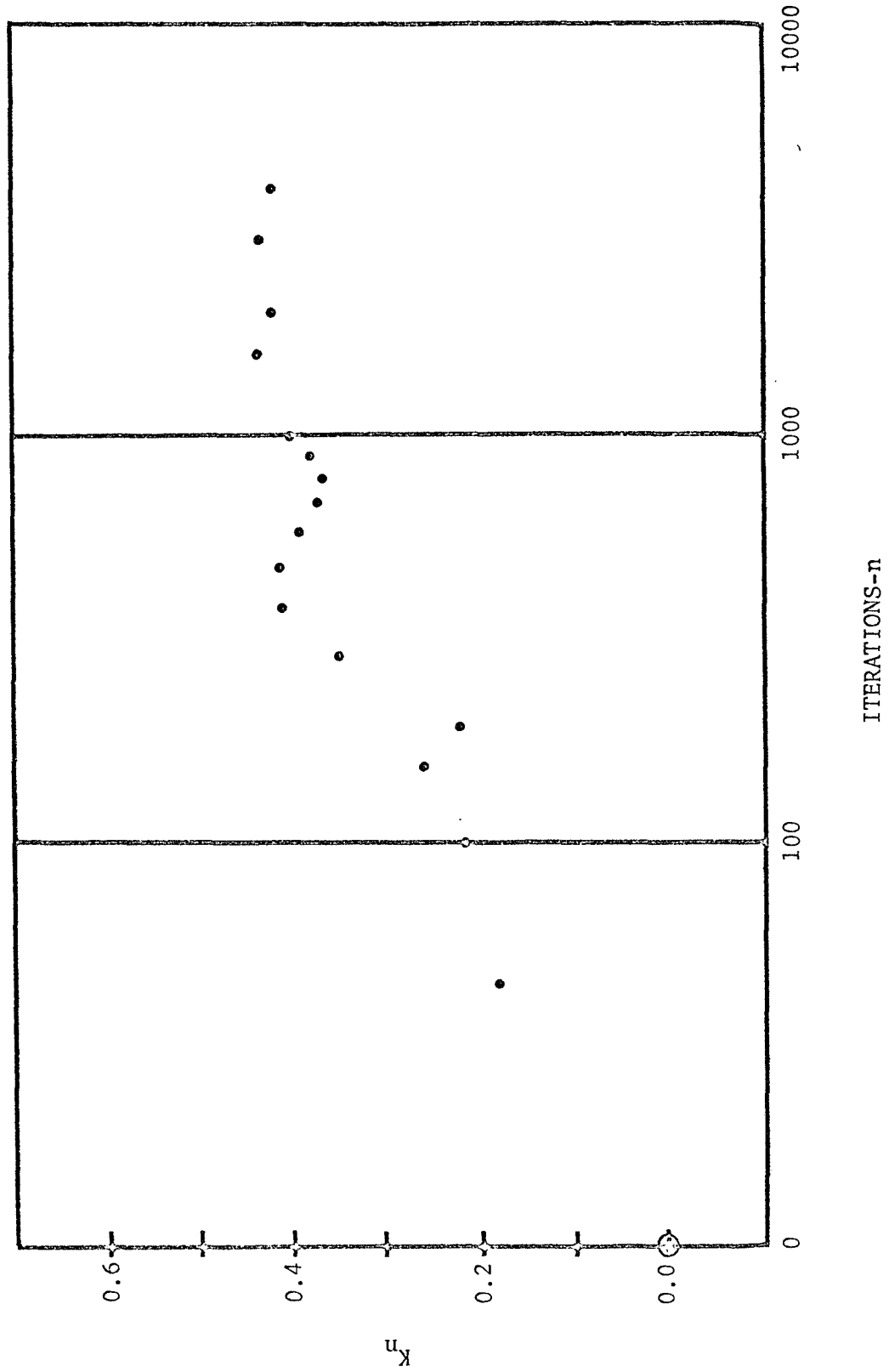


Fig. 6.10 Time Evolution of K_n for Example 6.9

TABLE 6.5

COMPARISON OF THE ADAPTIVELY FILTERED AND NON-FILTERED RESULTS
FOR EXAMPLE 6.9.

Quantity	K	Γ	Σ
Non-Filtered	1.0	1.90	1.75
Adaptively Filtered	.41	1.1	1.18

As in the saturation case, the adaptively filtered Γ and Σ are far superior to the non-filtered Γ and Σ . The reduction in Γ is 72.7%. Again the filtered value of Σ approaches its absolute lower bound of 1.0.

6.3.2 Nonlinearities That Have Unique Inverses

Consider the set of nonlinear functions $\{f(x)=y:f^{-1}(y) = x\}$. Then the adaptive stochastic algorithm, which learns the filter that is optimal in the sense of Eq. (5.42), is given by Eq. (5.45).

$$\hat{K}_{n+1} = \hat{K}_n + \frac{1}{n} f^{-1} (\delta_{n+1} \delta_n) W_{n+1} \quad (5.45)$$

EXAMPLE 6.10: Cubic Nonlinearity

Given: $\phi = 0.1$, $Q = 0.64$, $R = 4.0$

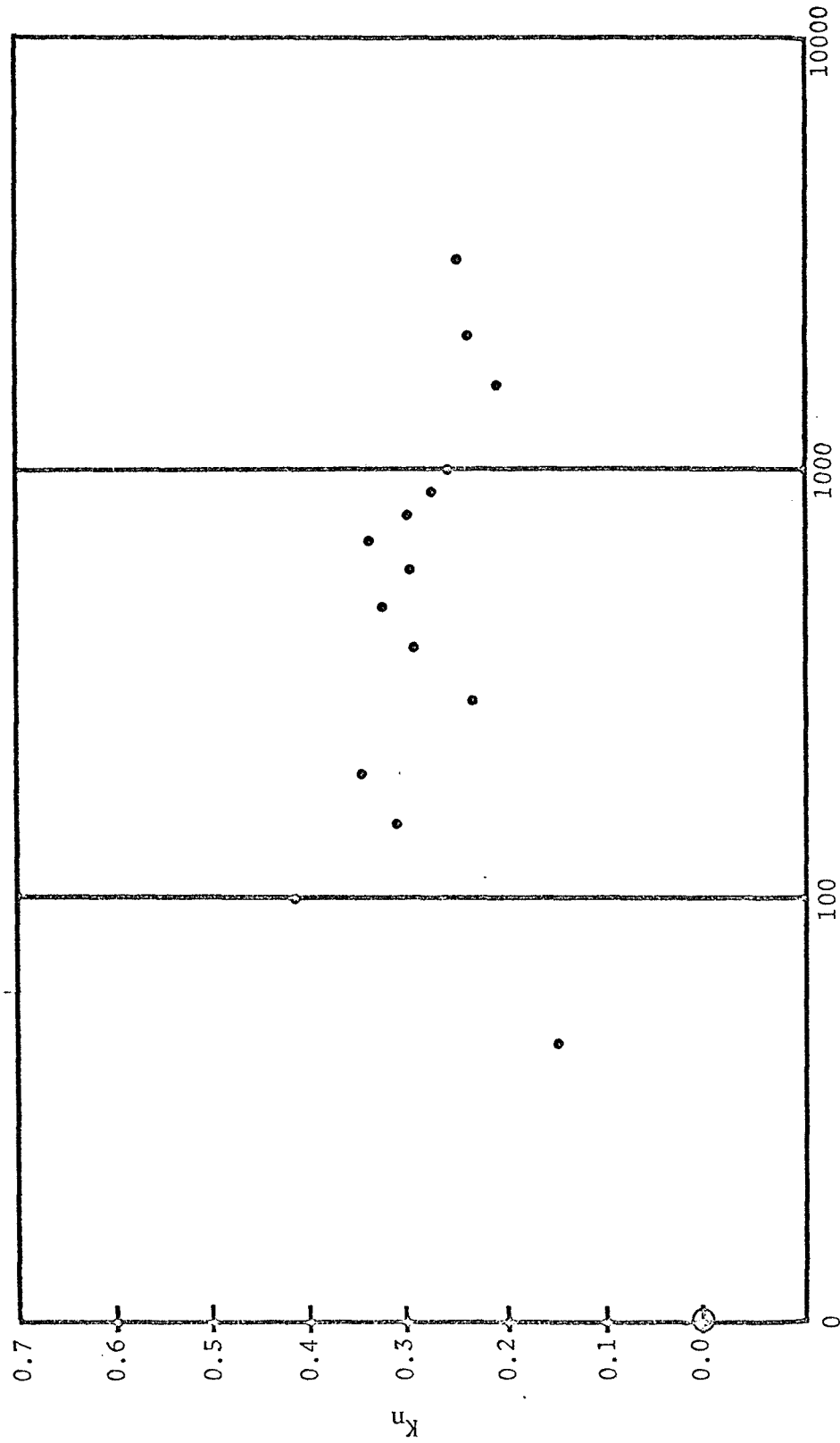
$$f(x) = \phi x^3$$

The behavior of the adaptive algorithm for $K_0 = 0.0$ is shown in Fig. 6.11. The "learned" filter is 0.232. The error covariances after 3,000 iterations are presented in Table 6.6.

TABLE 6.6

COMPARISON OF THE ADAPTIVELY FILTERED AND NON-FILTERED RESULTS
FOR EXAMPLE 6.10.

Quantity	K	Γ	Σ
Non-Filtered	1.0	2.70	.881
Adaptively Filtered	.232	.652	.681

Fig. 6.11 Time Evolution of K_n for Example 6.10

Observe that the filtered Γ is over four times smaller than the non-filtered Γ .

EXAMPLE 6.11: Cubic Nonlinearity

Given: $\phi = 0.1$, $Q = 0.25$, $R = 1.0$

The time evolution of the filter learning process for this problem is illustrated in Fig. 6.12. Its value after 4000 iterations is 0.205. The mean square errors are compared in Table 6.7.

TABLE 6.7

COMPARISON OF THE ADAPTIVELY FILTERED AND NON-FILTERED RESULTS
FOR EXAMPLE 6.11

Quantity	K	Γ	Σ
Non-Filtered	1.0	.742	.525
Adaptively Filtered	.205	.231	.252

The absolute lower bound on Σ is 0.25 and the Σ for the adaptively "learned" filter is only 0.252 which is twice as small as that for the non-filtered case. Also the filtered mean square error is 221% smaller than the non-filtered Γ . The conclusion is that the nonlinear algorithm of Eq. (5.45) gives very good results for the examples considered.

6.4 Application to a Thermionic Reactor Model

The thermionic reactor system considered in this section utilizes the linearized prompt-jump approximation to a simplified point reactor kinetics model. It is similar to the model treated in previous annual reports [Brehm, et al., 1968]. The reactor system is described by the following set of differential equations

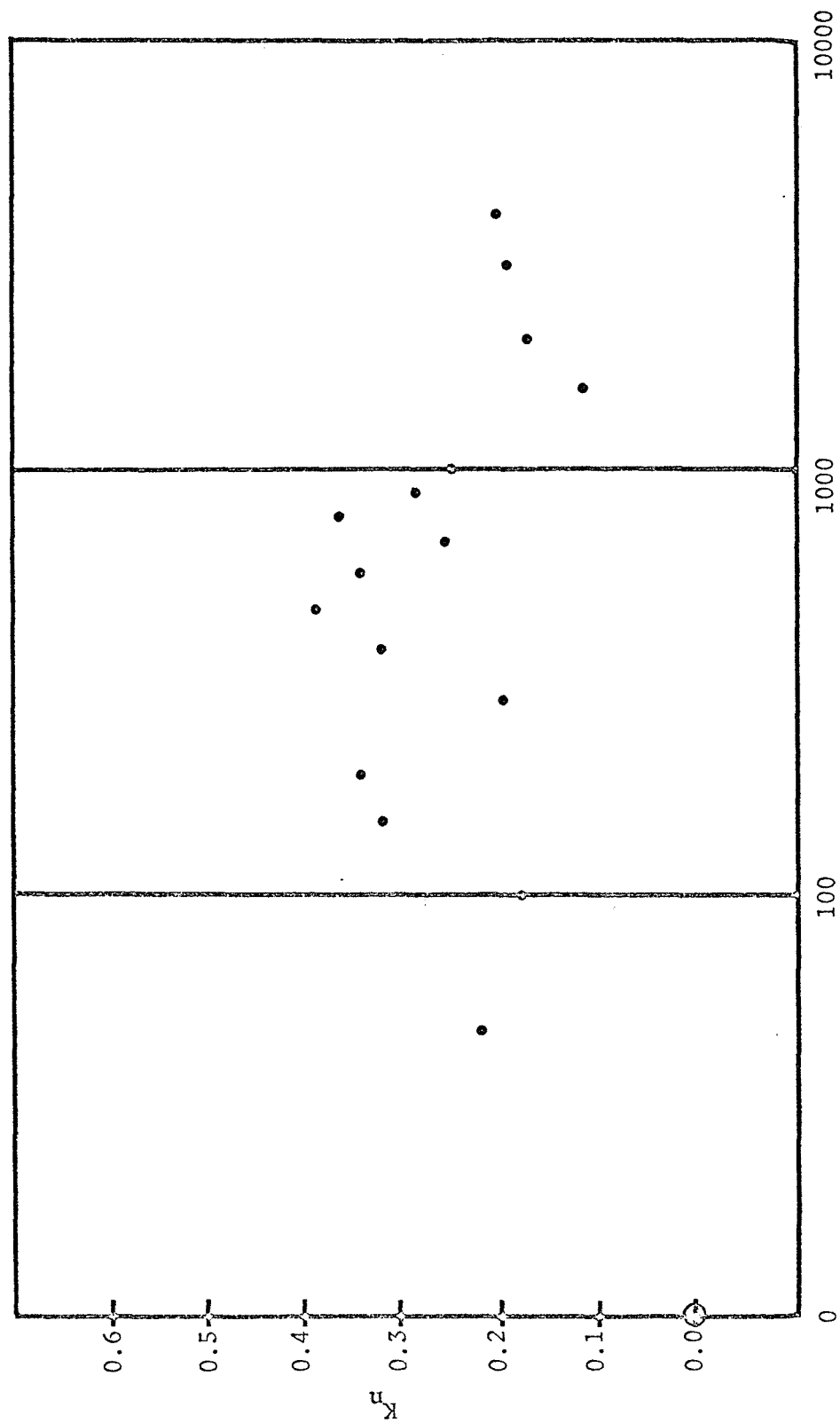


Fig. 6.12 Time Evolution of K_n for Example 6.11

$$\dot{x}_1 = -n_o \xi_1 x_2 - n_o \xi_2 x_3 + u_1(t) + \omega_1(t) \quad (6.4a)$$

$$\dot{x}_2 = b x_1 - v_1 x_2 + \omega_2(t) \quad (6.4b)$$

$$\dot{x}_3 = v_2 x_2 - v_3 x_3 + u_2(t) + \omega_3(t) \quad (6.4c)$$

where

x_1 is the neutron power density state variable

x_2 is the emitter surface temperature state variable

x_3 is the collector surface temperature state variable

$u_1(t)$ is the reactivity control input

$u_2(t)$ is the coolant control input

n_o is the equilibrium neutron power density

ξ_1 and ξ_2 are the reactivity temperature feedback coefficients

b , v_1 , v_2 , v_3 are the reciprocal system time constants and are explained in detail by Brehm, et al. [1968]

ω_1 , ω_2 , and ω_3 represent the random plant disturbance inherent in the reactor.

The measurement system for the state variables is described by

$$y_1(t) = h_1 x_1(t) + v_1(t) \quad (6.5a)$$

$$y_2(t) = h_2 x_2(t) + v_2(t) \quad (6.5b)$$

$$y_3(t) = h_3 x_3(t) + v_3(t) \quad (6.5c)$$

where

v_1 , v_2 , and v_3 represent the observation noise

h_1 , h_2 , and h_3 are the measurement transformation constants.

By multiplying each of the observation equations by the inverse of the appropriate measurement transformation, the following form is obtained.

$$z_1(t) = x_1(t) + v_1(t) \quad (6.6a)$$

$$z_2(t) = x_2(t) + v_2(t) \quad (6.6b)$$

$$z_3(t) = x_3(t) + v_3(t) \quad (6.6c)$$

Note that since the emitter surface temperature $x_2(t)$ can only be measured indirectly, it will be very noisy in general. By employing standard state variable notation Eqns. (6.4) and (6.6) may be compactly written

$$\dot{\underline{x}}(t) = A\underline{x}(t) + B\underline{u}(t) + \underline{\omega}(t) \quad (6.4)$$

$$\underline{z}(t) = \underline{x}(t) + \underline{v}(t) \quad (6.6)$$

where the state, control, and output vectors are respectively

$$\underline{x}(t) = \begin{bmatrix} x_1(t) \\ x_2(t) \\ x_3(t) \end{bmatrix} \quad \underline{u}(t) = \begin{bmatrix} u_1(t) \\ u_2(t) \end{bmatrix} \quad \underline{z}(t) = \begin{bmatrix} z_1(t) \\ z_2(t) \\ z_3(t) \end{bmatrix}$$

and the associated system matrices are

$$A = \begin{bmatrix} 0 & -n_o \xi_1 & -n_o \xi_2 \\ b & -v_1 & 0 \\ 0 & -v_2 & -v_3 \end{bmatrix} \quad B = \begin{bmatrix} 1 & 0 \\ 0 & 0 \\ 0 & 1 \end{bmatrix}$$

The remaining vectors $\underline{\omega}(t)$ and $\underline{v}(t)$ represent the plant perturbation noise and measurement noise terms respectively.

By invoking the Separation Theorem, the control term $\underline{u}(t)$ may be disregarded in learning the optimal Kalman filter. Thus, Eq. (6.4) becomes

$$\dot{\underline{\tilde{x}}}(t) = A\underline{\tilde{x}}(t) + \underline{\omega}(t) \quad (6.7)$$

To proceed with the development, numerical values must be assigned to the elements of the A matrix. Nominal values for these parameters are

$$\begin{array}{lll} b = 0.25 & \xi_1 = 0.5 & v_1 = 0.05 \\ n_0 = 80 & \xi_2 = 1.3 & v_2 = 0.55 \\ & & v_3 = 5.00 \end{array}$$

Using these values and a sampling time of 0.01 sec, the one-step state transition matrix equation describing the time evolution of (6.7) is

$$\underline{\tilde{x}}(n+1) = \Phi \underline{\tilde{x}}(n) + \underline{w}(n) \quad (6.8)$$

where

$$\Phi = \begin{bmatrix} 0.998 & -0.39 & -1.02 \\ 0.0025 & 0.985 & -0.004 \\ 0.0 & -0.0052 & 0.95 \end{bmatrix}$$

and

$\underline{w}(n)$ is zero mean random perturbation matrix whose covariance kernel is unknown.

The system defined by Eqns. (6.6) and (6.8) where the covariance matrices R and Q of $\underline{v}(n)$ and $\underline{w}(n)$ respectively are unknown, represents the simulation model of this section. The adaptive process Eqns. (4.39)

and (4.43) for learning the optimal filtering matrix for this system is initialized by setting K_0 equal to the null matrix. The experimental behavior of the adaptive algorithms for different values of Q and R was studied. Its success is illustrated by the representative example presented below.

With

$$Q = \begin{bmatrix} 2.14 & 1.73 & 1.39 \\ 1.73 & 2.01 & 1.75 \\ 1.39 & 1.75 & 2.13 \end{bmatrix}$$

and

$$R = \begin{bmatrix} 7.97 & 1.84 & 2.42 \\ 1.84 & 6.06 & -1.08 \\ 2.42 & -1.08 & 1.46 \end{bmatrix}$$

the optimal filter matrix is

$$K_{\text{opt}} = \begin{bmatrix} .154 & .133 & .203 \\ -.134 & .384 & .509 \\ -.235 & .219 & .827 \end{bmatrix}$$

The corresponding estimates after 3,000 iterations are

$$\hat{Q} = \begin{bmatrix} 2.08 & 1.60 & 1.35 \\ 1.60 & 2.17 & 1.66 \\ 1.35 & 1.66 & 2.08 \end{bmatrix}, \quad \hat{R} = \begin{bmatrix} 7.86 & 1.87 & 2.30 \\ 1.87 & 6.01 & -1.01 \\ 2.30 & -1.01 & 1.41 \end{bmatrix}$$

$$\hat{K} = \begin{bmatrix} .158 & .115 & .194 \\ -.151 & .368 & .500 \\ -.220 & .211 & .818 \end{bmatrix}$$

The minimum MSE for this thermionic system is 3.90 while the actual MSE resulting from the use of the "learned" filter $\hat{K}$ is 3.98. This means that $\hat{K}$ produces a MSE only 2.05%.

The excellent results presented above demonstrate the success obtained from a formal application of the adaptive learning theory to the state estimation problem in a thermionic reactor. The model used was deliberately simplified because it allowed greater fiscal economy in simulation without seriously degrading the illustrative effects. The assumption that all the state variables are available as noisy measurements can be eliminated by using the algorithm of Section 5.3.3.

6.5 Conclusion

The conclusion drawn from the examples simulated in the chapter is that the adaptive stochastic algorithms derived in Chapter 4 for learning the optimum Kalman filter perform excellently in practice. Many other systems were simulated and the results were of similar quality. Also the predicted acceleration of convergence obtained by using the weighting sequence defined by Eq. (4.39) was verified experimentally.

These same algorithms were modified in Chapter 5 to handle certain classes of nonlinear systems. Again the experimental results were quite good. However, a well-defined reference for performance evaluation does not exist in the nonlinear case.

This chapter concludes the contribution of this dissertation. A short summary of the entire thesis and suggestions for further research are the topics of Chapter 7.

CHAPTER 7

EPILOGUE

7.1 Summary

Kalman filter theory is a complete and powerful solution to the linear stochastic state estimation problem. However, to apply the technique requires complete statistical knowledge of the random environment. Specifically, the plant perturbation covariance Q and the observation noise covariance R must be known a priori. In the real world such extensive information is generally not available. To alleviate this difficulty in applying Kalman filter theory, stochastic algorithms are developed which directly learn the optimal stationary discrete time Kalman filter without any knowledge of Q and R .

The structure of the report is as follows. After motivating the subject in Chapter 1, Wiener and Kalman filter theory are reviewed and compared in Chapter 2. This theory provides the foundation for the unsupervised learning criterion of Chapter 3. There the learning criterion is shown to be a necessary and sufficient condition for optimal filtering. In Chapter 4 the adaptive stochastic algorithms, required to accomplish the adaptation indicated by the learning criterion, are derived. The theory of Stochastic Approximation is invoked to prove probability one convergence of the algorithms to the optimum Kalman filter. Additions and extensions to these results are presented in Chapter 5. Methods for estimating the additional quantities necessary

for optimal smoothing are developed. The most important contribution of Chapter 5 is the lifting of the restriction of Chapter 4 requiring that the inverse of the observation matrix exist. The stochastic algorithms are also modified to accommodate slowly time varying statistics and certain classes of nonlinear systems. The experimental results obtained by simulating the algorithm of Chapter 4 and Section 5.5 on a CDC 6400 digital computer are given in Chapter 6. In particular, the theory is successfully applied to the adaptive state estimation problem in a thermionic reactor. This data confirms the theoretical expectations and predictions of the first part of the study.

7.2 Recommendations for Further Study

In this report only sampled data were considered. The author conjectures that analogous results exist for continuous data. This should be a fruitful research topic.

At a lower level it would be interesting to apply the results of Section 5.3.3 to examples where the inverse of the observation matrix does not exist.

Another topic worthy of investigation would be to embed the nonlinear technique of Section 5.5 into a general control system containing a single nonlinearity such as a cubic or saturating amplifier.

APPENDIX A

LISTING OF THE SIMULATION PROGRAM

The following pages comprise a listing of the computer program developed to simulate the algorithms of Section 6.2.


```

145 FORMAT(/// * V MATRIX*)
DO 150 I=1,ORDER
  READ(5,90) (V(I,J),J=1,ORDER)
150 WRITE(6,87) (V(I,J),J=1,ORDER)
C  DEFINE INITIAL CONDITIONS
  DUMMY=RNRF(RANSET)
  DO 210 I=1,ORDER
210  RNSEQV(I)=GNOISE(N)
    CALL MATRIX(20,ORDER,ORDER,1,V,10,RNSEQV,10,RNSEQV,10)
    CALL MATRIX(21,ORDER,1,0,X,10,RNSEQV,10,XHAT,10)
    DO 230 I=1,ORDER
230  RNSEQW(I)=GNOISE(N)
    CALL MATRIX(20,ORDER,ORDER,1,W,10,RNSEQW,10,RNSEQW,10)
    CALL MATRIX(20,ORDER,ORDER,1,USTM,10,X,10,X,10)
    CALL MATRIX(21,ORDER,1,0,X,10,RNSEQW,10,X,10)
    DO 240 I=1,ORDER
240  RNSEQV(I)=GNOISE(N)
    CALL MATRIX(20,ORDER,ORDER,1,V,10,RNSEQV,10,RNSEQV,10)
    CALL MATRIX(21,ORDER,1,0,X,10,RNSEQV,10,Z,10)
    CALL MATRIX(20,ORDER,ORDER,1,USTM,10,XHAT,10,XHAT,10)
    CALL MATRIX(22,ORDER,1,0,Z,10,XHAT,10,DELTA1,10)
    CALL MATRIX(20,ORDER,ORDER,1,FILTER,10,DELTA1,10,STORE1,10)
    CALL MATRIX(21,ORDER,1,0,XHAT,10,STORE1,10,XHAT,10)
    DO 270 I=1,ORDER
    DO 270 J=1,ORDER
270  STMINV(I,J)=STM(I,J)
    CALL MATRIX(10,ORDER,ORDER,0,STMINV,10,DUMMY)
    NOINTS=NOPRINT=0
    IF(ITOT.LE.0) GO TO 300
C  FILTER INCREMENT PROCESS
    DO 2700 I=1,200
2700  SUM1(I,1)=0.0
    DO 2720 I=1,ITOT
    DO 2710 J=1,ORDER
    RNSEQV(J)=GNOISE(N)
2710  RNSEQW(J)=GNOISE(N)
    CALL MATRIX(20,ORDER,ORDER,1,V,10,RNSEQV,10,RNSEQV,10)
    CALL MATRIX(20,ORDER,ORDER,1,W,10,RNSEQW,10,RNSEQW,10)
    CALL MATRIX(20,ORDER,ORDER,1,STM,10,X,10,STORE1,10)
    CALL MATRIX(21,ORDER,1,0,STORE1,10,RNSEQW,10,X,10)
    CALL MATRIX(21,ORDER,1,0,X,10,RNSEQV,10,Z,10)
    CALL MATRIX(20,ORDER,ORDER,1,STM,10,XHAT,10,STORE1,10)
    CALL MATRIX(22,ORDER,1,0,Z,10,STORE1,10,DELTA2,10)
    CALL MATRIX(0,ORDER,1,0,DELTA1,10,DELWGT,1)
    CALL MATRIX(20,ORDER,1,ORDER,DELTA1,10,DELWGT,1,STORE1,10)
    CALL MATRIX(21,ORDER,ORDER,0,SUM1,10,STORE1,10,SUM1,10)
    CALL MATRIX(20,ORDER,1,ORDER,DELTA2,10,DELWGT,1,STORE1,10)
    CALL MATRIX(21,ORDER,ORDER,0,SUM12,10,STORE1,10,SUM12,10)
    CALL MATRIX(20,ORDER,ORDER,1,FILTER,10,DELTA2,10,STORE1,10)
    CALL MATRIX(20,ORDER,ORDER,1,STM,10,XHAT,10,XHAT,10)
    CALL MATRIX(21,ORDER,1,0,XHAT,10,STORE1,10,XHAT,10)
2720  CALL MATRIX(1,ORDER,1,0,DELTA2,10,DELTA1,10)
    CALL MATRIX(1,ORDER,ORDER,0,SUM1,10,WEIGHT,10)
    CALL MATRIX(10,ORDER,ORDER,0,WEIGHT,10,DUMMY)
    CALL MATRIX(20,ORDER,ORDER,ORDER,SUM12,10,WEIGHT,10,FDELTAK,10)
    CALL MATRIX(20,ORDER,ORDER,ORDER,STMINV,10,FDELTAK,10,DELTAK,10)
    CALL MATRIX(21,ORDER,ORDER,0,FILTER,10,DELTAK,10,FILTER,10)
    WRITE(6,2730) ITOT

```

```

2730 FORMAT(//RESULTS OF OPTIONAL FILTER INCREMENT,* I7 * ITERATIONS*
1 /// * SUM1 MATRIX* )
DO 2740 I=1,ORDER
2740 WRITE(6,87) (SUM1(I,J),J=1,ORDER)
WRITE(6,2750)
2750 FORMAT(// * SUM12 MATRIX*)
DO 2760 I=1,ORDER
2760 WRITE(6,87) (SUM12(I,J),J=1,ORDER)
WRITE(6,480)
DO 2770 I=1,ORDER
2770 WRITE(6,87) (WEIGHT(I,J),J=1,ORDER)
WRITE(6,2780)
2780 FORMAT(// * DELTAK MATRIX*)
DO 2790 I=1,ORDER
2790 WRITE(6,87) (DELTAK(I,J),J=1,ORDER)
WRITE(6,460)
DO 3000 I=1,ORDER
3000 WRITE(6,87) (FILTER(I,J),J=1,ORDER)
C FILTER ADAPTATION PROCESS
300 DO 320 I=1,ORDER
RNSEQV(I)=GNOISE(N)
320 RNSEQW(I)=GNOISE(N)
CALL MATRIX(20,ORDER,ORDER,1,V,10,RNSEQV,10,RNSEQV,10)
CALL MATRIX(20,ORDER,ORDER,1,W,10,RNSEQW,10,RNSEQW,10)
CALL MATRIX(20,ORDER,ORDER,1,OSTM,10,X,10,X,10)
CALL MATRIX(21,ORDER,1,0,X,10,RNSEQW,10,X,10)
CALL MATRIX(21,ORDER,1,0,X,10,RNSEQV,10,Z,10)
CALL MATRIX(20,ORDER,ORDER,1,USTM,10,XHAT,10,STORE1,10)
CALL MATRIX(22,ORDER,1,0,Z,10,STORE1,10,DELTA2,10)
CALL MATRIX(20,ORDER,ORDER,ORDER,ORDER,OSTM,10,FILTER,10,FILTER,10)
CALL MATRIX(23,ORDER,1,ORDER,DELTA1,10,WEIGHT,10,DELTWT,1)
CALL MATRIX(20,ORDER,1,ORDER,DELTA2,10,DELTWT,1,STORE1,10)
CALL MATRIX(21,ORDER,ORDER,0,FILTER,10,STORE1,10,FILTER,10)
CALL MATRIX(20,ORDER,ORDER,ORDER,ORDER,OSTMINV,10,FILTER,10,FILTER,10)
IF(IS.EQ.1H ) GO TO 350
CALL MATRIX(0,ORDER,ORDER,0,FILTER,10,STORE1,10)
DO 340 I=1,ORDER
DO 340 J=1,ORDER
340 FILTER(I,J)=(FILTER(I,J)+STORE1(I,J))/2.0
350 CALL MATRIX(20,1,ORDER,1,DELTWT,1,DELTA1,10,STORE1,10)
DIVISOR=STORE1(1,1)+LAMBDA
CALL MATRIX(20,ORDER,1,ORDER,DELTA1,10,DELTWT,1,STORE1,10)
CALL MATRIX(20,ORDER,ORDER,ORDER,WEIGHT,10,STORE1,10,STORE1,10)
DL=LAMBDA*DIVISOR
DO 375 I=1,ORDER
DO 370 J=1,ORDER
370 WEIGHT(I,J)=WEIGHT(I,J)/LAMBDA+STORE1(I,J)/DL
375 DELTA1(I)=DELTA2(I)
CALL MATRIX(20,ORDER,ORDER,1,OSTM,10,XHAT,10,XHAT,10)
CALL MATRIX(20,ORDER,ORDER,1,FILTER,10,DELTA2,10,STORE1,10)
CALL MATRIX(21,ORDER,1,0,XHAT,10,STORE1,10,XHAT,10)
NOINTS=NOINTS+1
NOPRINT=NOPRINT+1
IF(NOINTS.LT.ITER.AND.NOPRINT.LT.IX) GO TO 300
WRITE(6,410) NOINTS,DIVISOR,TITLE,(DELTA1(I),I=1,ORDER)
410 FORMAT(*11ITERATION* 18,8X *DIVISOR =* E12.4,8X,8A10 ///
1 * DELTA1 VECTOR * // 10E13.4)
WRITE(6,430) (X(I),I=1,ORDER)

```

```

430 FORMAT(// * X VECTOR * // 10E13,4)
   WRITE(6,440) (Z(I),I=1,ORDER)
440 FORMAT(// * Z VECTOR * // 10E13,4)
   WRITE(6,450) (XHAT(I),I=1,ORDER)
450 FORMAT(// * XHAT VECTOR * // 10E13,4)
   WRITE(6,460)
460 FORMAT(// * FILTER MATRIX *)
   DO 470 I=1,ORDER
470   WRITE(6,87) (FILTER(I,J),J=1,ORDER)
   WRITE(6,480)
480 FORMAT(// * WEIGHT MATRIX *)
   DO 490 I=1,ORDER
490   WRITE(6,87) (WEIGHT(I,J),J=1,ORDER)
   NOPRINT=0
   IF(NPOINTS.LT.ITER) GO TO 300
C FINAL CALCULATIONS
   DIV=ITOT+ITER
   CALL MATRIX(10,ORDER,ORDER,0,WEIGHT,10,DUMMY)
   DO 5000 I=1,ORDER
   DO 5000 J=1,ORDER
5000   CHAT(I,J)=WEIGHT(I,J)/DIV
   CALL MATRIX(20,ORDER,ORDER,ORDER,FILTER,10,CHAT,10,SIGMA,10)
   CALL MATRIX(22,ORDER,ORDER,ORDER,CHAT,10,SIGMA,10,RHAT,10)
   CALL MATRIX(22,ORDER,ORDER,ORDER,IDENTY,10,FILTER,10,STORE1,10)
   CALL MATRIX(20,ORDER,ORDER,ORDER,STORE1,10,SIGMA,10,CERROR,10)
   CALL MATRIX(0,ORDER,ORDER,ORDER,STM,10,STORE1,10)
   CALL MATRIX(20,ORDER,ORDER,ORDER,CERROR,10,STORE1,10,QHAT,10)
   CALL MATRIX(20,ORDER,ORDER,ORDER,STM,10,QHAT,10,STORE1,10)
   CALL MATRIX(22,ORDER,ORDER,ORDER,SIGMA,10,STORE1,10,QHAT,10)
   WRITE(6,5010) TITLE
5010 FORMAT(*FINAL CALCULATIONS * 30X,8A10 /// * CHAT MATRIX * )
   DO 5020 I=1,ORDER
5020   WRITE(6,87) (CHAT(I,J),J=1,ORDER)
   WRITE(6,5030)
5030 FORMAT(// * SIGMA MATRIX *)
   CALL MATRIX(0,ORDER,ORDER,ORDER,SIGMA,10,STORE1,10)
   DO 1001 I=1,ORDER
   DO 1001 J=1,ORDER
1001   SIGMA(I,J) = (SIGMA(I,J) + STORE1(I,J)) / 2.0
   DO 5040 I=1,ORDER
5040   WRITE(6,87) (SIGMA(I,J),J=1,ORDER)
   WRITE(6,5050)
5050 FORMAT(// * CERROR MATRIX *)
   CALL MATRIX(0,ORDER,ORDER,ORDER,CERROR,10,STORE1,10)
   DO 1002 I=1,ORDER
   DO 1002 J=1,ORDER
1002   CERROR(I,J) = (CERROR(I,J)+STORE1(I,J)) / 2.0
   DO 5060 I=1,ORDER
5060   WRITE(6,87) (CERROR(I,J),J=1,ORDER)
   WRITE(6,5070)
5070 FORMAT(// * RHAT MATRIX *)
   CALL MATRIX(0,ORDER,ORDER,ORDER,RHAT,10,STORE1,10)
   DO 1003 I=1,ORDER
   DO 1003 J=1,ORDER
1003   RHAT(I,J) = (RHAT(I,J) + STORE1(I,J)) / 2.0
   DO 5080 I=1,ORDER
5080   WRITE(6,87) (RHAT(I,J),J=1,ORDER)
   WRITE(6,5090)

```

```

5090 FORMAT (// * QHAT MATRIX *)
      CALL MATRIX (0,ORDER,ORDER,0,QHAT,10,STORE1,10)
      DO 1004 I=1,ORDER
      DO 1004 J=1,ORDER
1004 QHAT(I,J) = (QHAT(I,J)+STORE1(I,J)) /2.0
      DO 5100 I=1,ORDER
5100 WRITE(6,87) (QHAT(I,J),J=1,ORDER)
      DO 1010 I=1,ORDER
      DO 1010 J=1,ORDER
      STORE1(J,I) = W(I,J)
      STORE2(J,I) = V(I,J)
1010 CALL MATRIX (20,ORDER,ORDER,ORDER,W,10,STORE1,10,Q,10)
      CALL MATRIX (20,ORDER,ORDER,ORDER,V,10,STORE2,10,R,10)
      PRINT 1011
1011 FORMAT (// * Q MATRIX *)
      DO 1012 I=1,ORDER
1012 WRITE (6,87) (Q(I,J),J=1,ORDER )
      PRINT 1013
1013 FORMAT (// * R MATRIX * )
      DO 1014 I=1,ORDER
C
C      CALCULATE CLAMBDA
1014 WRITE(6,87) (R(I,J),J=1,ORDER)
      DO 1020 I=1,ORDER
      DO 1020 J=1,ORDER
      STORE1(J,I) = 0.0
      IF (I.EQ.J) STORE1(J,I) = 1.0
1020 CONTINUE
      DO 1021 I=1,ORDER
      DO 1021 J=1,ORDER
1021 STORE2(I,J) = STORE1(I,J) - FILTER(I,J)
      CALL MATRIX (20,ORDER,ORDER,ORDER,STORE2,10,SIGMA,10,STORE1,10)
      DO 1023 I=1,ORDER
      DO 1023 J=1,ORDER
1023 STORE3(J,I) =STORE2(I,J)
      CALL MATRIX (20,ORDER,ORDER,ORDER,STORE1,10,STORE3,10,STORE2,10)
      CALL MATRIX (20,ORDER,ORDER,ORDER,FILTER,10,R,10,STORE1,10)
      DO 1024 I=1,ORDER
      DO 1024 J=1,ORDER
1024 STORE3(I,J) = FILTER(J,I)
      CALL MATRIX (20,ORDER,ORDER,ORDER,STORE1,10,STORE3,10,STORE1,10)
      DO 1025 I=1,ORDER
      DO 1025 J=1,ORDER
1025 STORE2(J,I) = STORE2(J,I) + STORE1(J,I)
      PRINT 1026
1026 FORMAT (// * CLAMBDA MATRIX *)
      DO 1027 I=1,ORDER
C
C      CALCULATE NSIGMA
1027 WRITE(6,87) (STORE2(I,J),J=1,ORDER)
      CALL MATRIX (20,ORDER,ORDER,ORDER,STM,10,STORE2,10,STORE1,10)
      DO 1030 I=1,ORDER
      DO 1030 J=1,ORDER
1030 STORE3(J,I) = STM(I,J)
      CALL MATRIX (20,ORDER,ORDER,ORDER,STORE1,10,STORE3,10,STORE2,10)
      DO 1035 I=1,ORDER
      DO 1035 J=1,ORDER
1035 STORE2(J,I) = STORE2(J,I) + Q(J,I)

```

```
      WRITE(6,1036)
1036  FORMAT (// 'HSIGMA MAIRIX* )
      DO 1038 I=1,ORDER
1038  WRITE(6,1037) (STORE2(I,J),J=1,ORDER)
      GO TO 40
      END
```

APPENDIX B
TONELLI'S THEOREM

THEOREM B.1:

Given the relation

$$x(n,v) = E_w \left\{ \sum_{j=-\infty}^n y(n,j,v,w) \right\}$$

then the interchange of the expectation and summation

$$x(n,v) = \sum_{j=-\infty}^n E_w \{ y(n,j,v,w) \}$$

is valid if either

$$(i) \quad E_w \left\{ \sum_j |y(n,j,v,w)| \right\} < \infty$$

$$(ii) \quad \sum_{j=-\infty}^n E_w \{ |y(n,j,v,w)| \} < \infty$$

This is a special case of Tonelli's Theorem [see for example Royden, 1967, p. 234].

LIST OF REFERENCES

- Albert, A. A. and L. A. Gardner. Stochastic Approximation and Nonlinear Regression, M.I.T. Press, Cambridge, Massachusetts, 1967.
- Balakrishnan, A. V. "Report on Progress in Information Theory in the U.S.A., 1960-1963: Prediction and Filtering," IEEE Trans. on Information Theory, IT-9 (October 1963), pp. 237-239.
- Brehm, R. L., D. L. Hetrick, J. G. Guppy, K. D. Kerns, and T. R. Schmidt. "Dynamics of a Fast Reactor with In-Core Thermionic Converters, Second Annual Report," Engineering Experimental Station, University of Arizona, Tucson, Arizona, 1968.
- Bryson, A. E., and D. E. Johansen. "Linear Filtering for Time-Varying Systems Using Measurements Containing Colored Noise," IEEE Trans. on Automatic Control, AC-10 (January 1965), pp. 4-10.
- Bryson, A. E., and R. K. Mehra. "Linear Smoothing Using Measurements Containing Correlated Noise with an Application to Inertial Navigation," IEEE Trans. on Automatic Control, AC-13 (October 1968), pp. 496-503.
- Deutsch, R. Estimation Theory, Prentice-Hall, Englewood Cliffs, N. J., 1965.
- Doob, J. L. Stochastic Processes, John Wiley and Sons, New York, 1963.
- Dvoretzky, A. "On Stochastic Approximation," Third Berkeley Symposium on Math. Statistics and Probability, 1956.
- Fu, K. S. Sequential Methods in Pattern Recognition and Machine Learning, Academic Press, New York, 1968.
- Gladyshev, E. G. "On Stochastic Approximation," Theory of Probability and Its Applications, Vol. X, No. 2 (1965), pp. 275-278.
- Hampton, R. L. T. Stochastic Approximation and Its Engineering Applications, Engineering Experiment Station, The University of Arizona, December, 1967.
- Hancock, J. C. and P. A. Wintz. Signal Detection Theory, McGraw-Hill, New York, 1966.
- Heffes, H. "The Effect of Erroneous Models on the Kalman Filter Response," IEEE trans. on Automatic Control, AC-11 (July 1966), pp. 541-543.
- Hilborn, C. G., and D. G. Lainiotis. "Optimal Estimation in the Presence of Unknown Parameters," IEEE Trans. on System Science and Cybernetics, SSC-5 (January 1969) pp. 38-43.

- Ho, Y. C. and C. Blaydon. "On the Abstraction Problem in Pattern Classification," Proceedings of the 1966 National Electronics Conference, (1966), pp. 857-862.
- Joseph, P. D., and J. Tou. "On Linear Control Theory," AIEE Trans. on Applications and Industry, 80 (1961), pp. 193-196.
- Kailath, T. "Adaptive Matched Filters," Chapter 6 in Mathematical Optimization Techniques, Richard Bellman, Ed., University of California Press, Los Angeles, 1963.
- Kalman, R. E. "A New Approach to Linear Filtering and Prediction Problems," J. Basic Engr., 82 D (1960), pp. 33-45.
- Kalman, R. E., and R. S. Bucy. "New Results in Linear Filtering and Prediction Theory," J. Basic Engr., 83 D (1961), pp. 95-108.
- Kolmogorov, A. N. "Interpolation and Extrapolation von Stationären Zufälligen Folgen," Bull. Acad. Sci. (USSR), Ser. Math., 5 (1941), pp. 3-14.
- Korn, G. A. Random-Process Simulation and Measurement, McGraw-Hill, New York, 1966.
- Lee, R. C. K. Optimal Estimation, Identification and Control, M.I.T. Press, Cambridge, Mass., 1964.
- Magill, D. T. "Optimal Adaptive Estimation of Sampled Stochastic Processes," IEEE Trans. on Automatic Control, AC-10 (October 1965), pp. 434-439.
- Meditch, J. S. "Orthogonal Projection and Discrete Optimal Linear Smoothing," S.I.A.M. Journal on Control, Vol. 5, 1967.
- _____. Stochastic Optimal Linear Estimation and Control, McGraw-Hill, New York, 1969.
- Motyka, P. R. and J. A. Cadzow. "The Factorization of Discrete Process Spectral Matrices," Proceedings of the National Electronics Conference, 23 (1967), pp. 66-73.
- Nishimura, T. "On the A Priori Information in Sequential Estimation Problems," IEEE Trans. on Automatic Control, AC-11 (April 1966), pp. 197-204.
- _____. "Error Bounds of Continuous Kalman Filters and the Application to Orbit Determination Problems," IEEE Trans. on Automatic Control, AC-12 (June 1967), pp. 268-275.

- Robbins, H., and S. Monro. "A Stochastic Approximation Method," Ann. Math. Stat., Vol. 22, No. 1, 1951.
- Royden, H. L., Real Analysis, Macmillan, New York, 1963.
- Sage, A. P. Optimum Systems Control, Prentice-Hall, Englewood Cliffs, N. J., 1968.
- Schultz, D. G., and J. L. Melsa. State Functions and Linear Control Systems, McGraw-Hill, New York, 1967.
- Shellenbarger, J. C. "Estimation of Covariance Parameters for an Adaptive Kalman Filter," Proceedings of the 1966 N.E.C., (1966) pp. 698-702.
- _____. "A Multivariate Learning Technique for Improved Dynamic System Performance," Proceedings of the 1967 N.E.C., (1967), pp. 146-151.
- Sherman, S. "A Theorem on Convex Sets With Applications," Ann. Math. Stat., 26 (1955), pp. 763-767.
- _____. "Non-Mean-Square Error Criteria," IRE Trans. on Information Theory, IT-4 (September 1958), pp. 125-126.
- Soong, T. T. "On A Priori Statistics in Minimum Variance Estimation Problems," J. Basic Engineering, 87 D (1965), pp. 109-112.
- Sorenson, H. W. "Kalman Filtering Techniques," Chapter 6 in Advances in Control Systems, Vol. 3, C. T. Leondes, Ed., Academic Press, New York, 1966.
- Stumper, F. L. "A Bibliography of Information Theory, Communication and Cybernetics," IRE Trans. on Information Theory, IT-1 (September 1955), pp. 35-37.
- Van Trees, H. L. Detection, Estimation and Modulation Theory, John Wiley and Sons, New York, 1968.
- Viterbi, A. J. Principles of Coherent Communication, McGraw-Hill, New York, 1966.
- Wiener, N. The Extrapolation, Interpolation, and Smoothing of Stationary Time Series, John Wiley and Sons, New York, 1949.