

**A NEW ALGORITHM
FOR
STOCHASTIC APPROXIMATION**

A NEW ALGORITHM
FOR
STOCHASTIC APPROXIMATION

by
Michael Paul Griscik, B. Eng.

A Thesis

Submitted to the Faculty of Graduate Studies
in Partial Fulfilment of the Requirements
for the Degree
Master of Engineering

McMaster University

April 1970

Master of Engineering
Electrical Engineering

McMaster University
Hamilton, Ontario

Title: A NEW ALGORITHM
FOR
STOCHASTIC APPROXIMATION

Author: Michael Paul Griscik, B. Eng.
(McMaster University)

Supervisor: N. K. SINHA, B. Sc. Eng. (Banaras)
Ph. D. (University of Manchester)

Number of
Pages: vii, 172

Scope and A review of Stochastic Approximation and the
Contents: major contributions to the area is made. A
proof of convergence for the algorithm is developed.
An optimization is attempted on the rate of con-
vergence problem and the uniqueness problem is
faced. An alternative proof of convergence is
given as an independent check on the first one.
Simulation results are present in light of the
theory developed, and conclusions, limitations
and recommendations are presented.

Acknowledgements

It would not have been possible to perform this work without the financial support obtained from the National Research Council of Canada and from appointments as an assistant instructor at McMaster University. My thanks go to Dr. C. K. Campbell, Dr. S. S. Haykin, Dr. R. Kitai and Dr. N. K. Sinha for their assistance in obtaining this support. Special thanks are extended to my supervisor, Dr. Sinha, for introducing me to this topic and for his guidance and encouragement in the preparation of this work.

I would like to thank my father and mother for their concern and encouragement, and also my father and mother-in-law for their sincere consideration.

Most important of all, I would like to thank my wife, Ann, and our son, Gregory, for their patience, understanding and unwavering faith during the progress of this work.

TABLE OF CONTENTS

<u>Chapter</u>	<u>Page</u>
I Introduction	1
1.1 Introductory Definition	1
1.2 Formulation of the Principal Problem in the Theory of Approximation	1
1.3 Scope and Advantages of Stochastic Approximation	3
1.4 Major Contributions	5
1.5 Algorithms	6
1.6 Preview	7
II Major Contributions to Stochastic Approximation. A Historical Review	10
2.1 Founding Contributors	10
2.2 The Robbins - Monro Technique	11
2.3 The Kiefer - Wolfowitz Method	15
2.4 The Work of A. Dvoretzky	20
2.5 Generalization and Applications	23
III Development of a Stochastic Approximation Algorithm	27
3.1 Background	27
3.2 Convergence	30
3.3 Development	32
3.4 Gamma Sequences	39
3.5 Consistency and Bias	41
IV Optimization of Convergence	46
4.1 Preamble for Optimization	46
4.2 The Discrete Maximum Principle for Optimization	47
4.3 Identification with Optimal Control Theory	51
4.4 Boundary Value Problem	56
4.5 Uniqueness of the Optimum Gamma Sequence	58

V	Simulation Results	59
	5.1 Structure of the Simulations	59
	5.2 Illustration of Simulation Results	62
	5.3 Summary of Results	87
VI	Alternative Proof of Convergence	93
	6.1 Alternative Approach	93
	6.2 Statement of Dvoretzky's Theorem	94
	6.3 Convergence using Dvoretzky's Theorem	97
	6.4 Summary	113
VII	Conclusions and Future Work	114
	References	119
	List of Symbols	123

Appendices

A	Relationship Between Reinforced Learning and Stochastic Approximation	126
B	Proof of Two Stochastic Approximation Algorithms	137
C	The Discrete Maximum Principle	151
D	Dvoretzky's Theorem on Stochastic Approximation	157
	D-1 Introduction	157
	D-2 Theorem	158
	D-3 The Extension	159
	D-4 Generalizations	160
E	Cauchy's Integral Test	163
F	Noise, Spectra, and Autocorrelation	167

LIST OF FIGURES, TABLES AND GRAPHS

		<u>Page</u>
Figure (1.1)	Schematic Illustration of Principal Problem	2
Figure (2.2.1)	Schematic Illustration of First Assumption for Robbins - Monro Technique	12
Figure (2.3.1)	Illustration of Requirement on Derivative of $m(\alpha)$ in Kiefer - Wolfowitz Method	16
Figure (2.3.2)	Locus of $m'(\alpha)$	16
Figure (3.1)	Schematic Illustration of Learning Control System	28
Table (5.2.1)		63
Graph (5.2.1) to (5.2.8)	A.P., S.S.E. and T.S.S.E. as a function of n	64 - 71
Table (5.2.2)		72
Graph (5.2.9) to (5.2.16)	A.P., S.S.E. and T.S.S.E. as a function of n	73 - 80
Graph (5.2.17) to (5.2.21)	Sum of Square Errors (S.S.E.) as a function of Λ ratio	82 - 86
Graph (5.2.22) to (5.2.25)	The Λ ratio of equi-utility for algorithms (5.1.4) and (5.1.5) based on S.S.E. and T.S.S.E. evaluations as a function of sample autocorrelation delay	88 - 91
Graph (5.2.26)	Normalized Expected Mean Square Error as a function of n	92
Figure (A-1)	T-maze Rat Experiment	126
Figure (E-1)	Schematic Illustration of Locus of $f(x)$	163
Graph (F-1)	Autocorrelation Function of Pseudo-Random noise	169

Graph (F-2)	Fast Fourier Power Spectra of Noise	170
to (F-4)		- 172

CHAPTER I

Introduction

1.1 Introductory Definition

Stochastic approximation is a method dealing with algorithms converging to some sought value if, as a result of the stochastic nature of the problem, the observations involve errors. The methods of most interest and value are those that are self-correcting in the sense that a mistake tends to diminish to zero in the limit. The convergence to a desired value is of some specified nature, for example, mean-square convergence.

1.2 Formulation of the Principal Problem in the Theory of Approximation

The main problem in the theory of approximation can be stated as follows: suppose that two functions $f(P)$ and $F(P; A_1, \dots, A_n)$ of the point $P \in \beta$ are defined within a certain point set β in a space of any number of dimensions. Here $F(P; A_1, \dots, A_n)$ depends on a certain number of parameters A_r , $r=1, \dots, n$. It is to so determine the parameters $A_1, \dots, A_n$ that the deviation of the function $F(P; A_1, \dots, A_n)$ from the function $f(P)$ in β will be a minimum. The distance between the function $f(P)$ and $F(P; A_1, \dots, A_n)$ must be defined;

generally, the Euclidian metric is most convenient. Once the metric is defined it is desired to estimate the value of the A_r , $r=1, \dots, n$, parameters for which the expected value of the metric satisfies some condition, such as, taking on a minimum or maximum value, or equaling some fixed value.

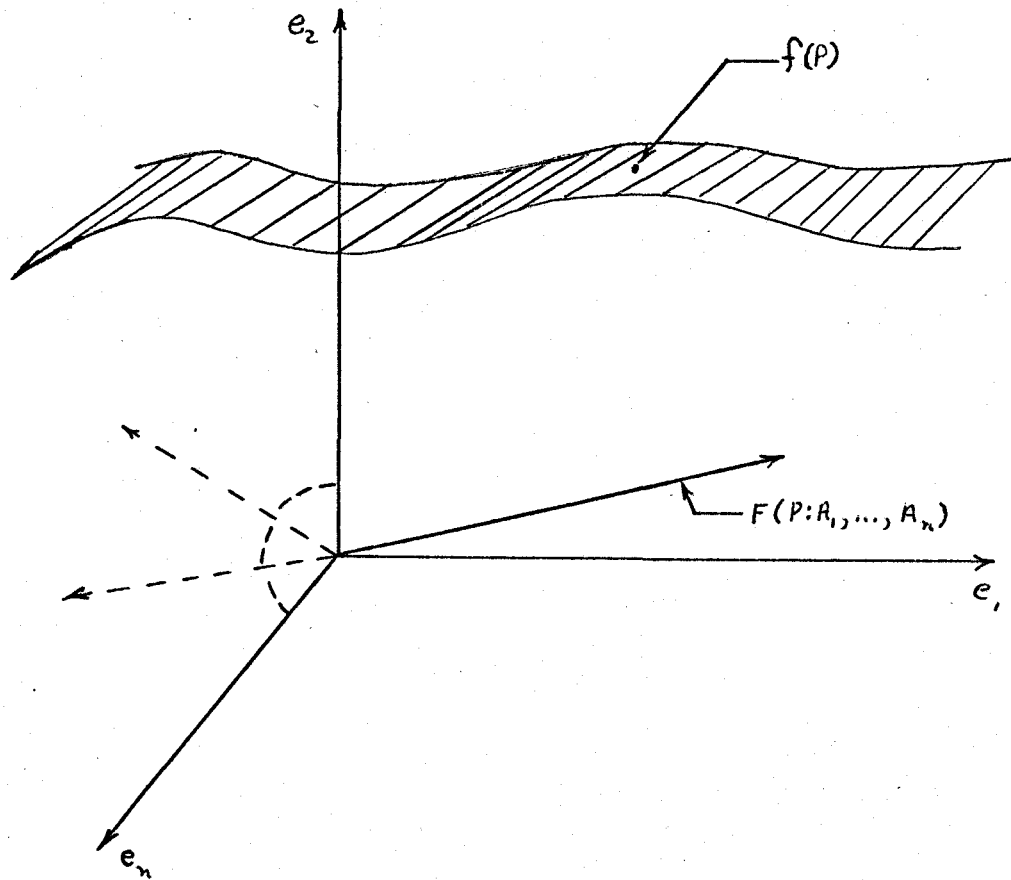


Figure 1.1 Schematic Illustration of Principal Problem

The statement of the problem is illustrated schematically in Figure 1.1. The unit vectors $e_1, e_2, \dots, e_n$ define a

space in the neighbourhood of the point set β . With P the function $f(P)$ is defined; it remains to estimate A_r , $r=1, \dots, n$ so that the function $F(P; A_1, \dots, A_n)$ coincides with the function $f(P)$ in this case. Here the A_r , $r=1, \dots, n$ are the scale factors of the vectors e_r , $r=1, \dots, n$.

It is the function of the parameters A_r , $r=1, \dots, n$ which is observable; however, there is a stationary random process $Z_r(t)$ (either time-continuous or time-discrete) which contaminates the observations. Stochastic approximation can be used in a situation of this nature.

1.3 Scope and Advantages of Stochastic Approximation Methods

Stochastic approximation methods are applicable to any problem that can be formulated as some form of regression problem in which repeated observations are made. To be specific, the use of these methods is particularly appropriate and advantageous when either one or both of two conditions occur. One condition is that the observation interval is so long that conventional methods of estimation are impractical because of the computational problems associated with processing long intervals of the observed data. The other condition is that there is no detailed knowledge available concerning the statistics of the processes $Z_r(t)$. The limitation on stochastic approximation methods is that there must be a unique solution to the regression problem of interest.

Conventional estimation methods usually proceed in two steps:

- (1) the observed data are used to estimate intermediate statistics; and
- (2) a set of (possibly nonlinear) equations relating the parameters of interest to these statistics are solved.

Stochastic approximation methods differ from this approach in two respects. First, the observation interval is divided into short subintervals of convenient fixed length (it can be of unit sample length also). Only the observations from a single subinterval are handled at a time, and after the data from a subinterval have been processed they are discarded and not used again. Second, the two separate algorithms for estimating statistics and solving equations are combined into a single algorithm.

In those situations in which stochastic approximation methods are applicable, their usage yields the following advantages:

- (1) Only a small interval of data needs to be processed.
- (2) Only simple computations are required, even when the actual functional dependence of the regression function on the parameters of interest is nonlinear.

- (3) The method may be employed in the absence of a priori knowledge of process statistics and in the absence of detailed knowledge of the relationship between the desired parameters and the observed data. In particular, the only requirement is that the regression function satisfy certain regularity conditions and that the regression problem have a unique solution.

If sufficient a priori knowledge concerning the statistics of $Z_r(t)$ and the functional relationship between the parameters and observed data is available, the third advantage can be replaced by the following desirable property: the methods can be made asymptotically efficient.

1.4 Major Contributions

To the area of stochastic approximation there have been three major contributions. These have been made by H. Robbins & S. Monro¹, J. Kiefer & J. Wolfowitz² and A. Dvoretzky³. Essentially the Robbins - Monro procedure is a method for finding the root of regression function whose form is unknown but that can be observed and sampled. The Kiefer - Wolfowitz procedure is similar in that it is a method for finding the extremum, maximum or minimum, of a regression function given only pertinent random observations. However, the basic idea of stochastic approximation is that any sort of iterative solution algorithm that is convergent,

based on direct observations of a regression function, can be adapted successfully to the case in which the observations are random. Dvoretzky formulates the problem in this light and proves several theorems to this effect showing both convergence with probability one and in the mean square sense. As a result, the original results concerning the Robbins - Monro and Kiefer - Wolfowitz methods follow as special cases of his results.

Although Dvoretzky's work represents a major contribution to the mathematical structure of stochastic approximation theory, it is of less practical importance. The primary reason is that in applications, one is usually concerned with the rate of convergence and its dependence on the parameters of the recursive solution algorithm. These factors are best handled by focusing attention on the particular algorithm being used.

1.5 Algorithms

In the area of control theory and specifically as pertains to the area of heuristic reinforced learning, two stochastic approximation algorithms have been developed and investigated by K. S. Fu⁴. The experience gained by studying the results of the first algorithm provided an insight into its advantages and disadvantages. The specific advantage desired in the algorithm was a faster rate of

convergence. This was the reason that led to the development of the second algorithm. It has all the advantages of a stochastic approximation algorithm as well as the feature of an accelerated rate of convergence.

It is evident from Fu's paper that there is room for improvement of the rate of convergence. In many situations a faster result is of great advantage. It means that less sampling is required, the number of computations is reduced and the overall time to get a result with a certain confidence level is less.

The objective of this work is to show the development of a new algorithm with a faster rate of convergence than the two already developed. Essentially it is required to show that the new algorithm converges regardless of the starting point and that it converges to the true value. The conditions and limitations on the convergence will be made explicit in the proof. Having proven convergence, comparison will be made between the two existing algorithms and the new algorithm developed here. The feature of most interest is the rate of convergence. It is this then, which will be highlighted in the comparisons.

1.6 Preview

Since the area of stochastic approximation is a relatively new field of study in mathematics and has been

applied to engineering only recently, it is essential that the theory as a whole be set in the proper prospective. In Chapter 2, a detailed outline of the historical sequence of contributions will be given along with a comprehensive review of the major blocks of theory presented in both the area of mathematical statistics and recently in electrical engineering. Chapter 3 contains an introduction to the relationship between the area of learning control systems and the area of stochastic approximation. In addition, the fundamental form of the stochastic approximation algorithm is given along with an outline of the previous algorithms in the area of concern. The basic proof of the convergence of the algorithm under study is given in detail. Chapter 4 contains an outline of the discrete maximum principle, an identification of the optimal control problem with the problem of optimization of the convergence rate of the algorithm, and a formulation and discussion of the optimization problem. Chapter 5 contains the results of computer simulations made using the new algorithm. A comparison of the sum of squared errors and sum of sample product squared errors for the new algorithm is made with similar error criteria for two previous algorithms introduced for comparison purposes. Chapter 6 contains an alternate proof for the new stochastic approximation algorithm using Dvoretzky's theorem. Chapter 7 follows with the conclusions

reached in this work and a brief review of the contents of this thesis.

CHAPTER II

Major Contributions to Stochastic Approximation: A Historical Review

2.1 Founding Contributors

The area of stochastic approximation is a relatively new field of study in mathematics. It is essentially the fusion of two major areas in mathematics: the area of random or stochastic processes and the area of deterministic approximation theory. The first major contribution which put forth a block of theory suggesting such a union of subjects was made by H. Robbins & S. Monroe¹. Their work was monumental not only from the point of view of being of major significance but that it was a pioneering work. Never before had a theory for solving a regression function stochastically been put forth. Within a year of their publication, there appeared the work of J. Kiefer & J. Wolfowitz² who extended the work of Robbins and Monroe and had applied it to the stochastic solution for the extremum, maximum or minimum, of a regression function. Both contributions were a new approach to an existing problem. It was not for some time afterwards that a generalized approach was taken and formulated. A. Dvoretzky³ is credited with just such a contribution. His work was the

major unifying theory to appear and to generalize the stochastic approximation problem. The Robbins - Monro technique and the Kiefer - Wolfowitz method are both special cases of Dvoretzky's Theorem. In addition to his theorem, Dvoretzky also provided a number of extensions and five generalizations thereby providing an all encompassing theory.

It is intended to present an outline of these three major contributions to the theory of stochastic approximation. In addition, a review of application and extensions of this work will be given, in particular, those areas pertaining to electrical engineering.

2.2 The Robbins - Monro Technique

Consider only the one dimensional Robbins - Monro technique for a scalar valued parameter α . Now, given a sequence of random entities $Z_1, Z_2, \dots$, and a scalar valued function of Z and α , $f(Z, \alpha)$. Each of the quantities Z_n , $n=1, 2, \dots$, may represent one or more random variables or a random process of given duration observed at sequential time intervals. It is required to find the value of α for which

$$m(\alpha) = m_0 \quad (2.2.1)$$

where m_0 is a nominal value of the function $m(\alpha)$ and where the function $m(\alpha)$ is defined by the equation

$$m(\alpha) = E\{f(Z, \alpha)\} \quad (2.2.2)$$

where the α which corresponds to this solution is called θ . In the proof of their technique, Robbins and Monro make a number of assumptions. These assumptions will be given along with the explanation rather than listing them arbitrarily at the conclusion of the description. The first assumption is that there exist constants k_0 and k'_0

$$0 < k_0 \leq k'_0 \leq \infty \quad (2.2.3)$$

such that,

$$k_0(\alpha - \theta)^2 \leq \{m(\alpha) - m_0\}(\alpha - \theta) \leq k'_0(\alpha - \theta)^2 \quad (2.2.4)$$

This simply says that $m(\alpha)$ must lie between two straight lines, one of positive slope k_0 , and the second of finite positive slope k'_0 . A schematic illustration is given in Figure (2.2.1).

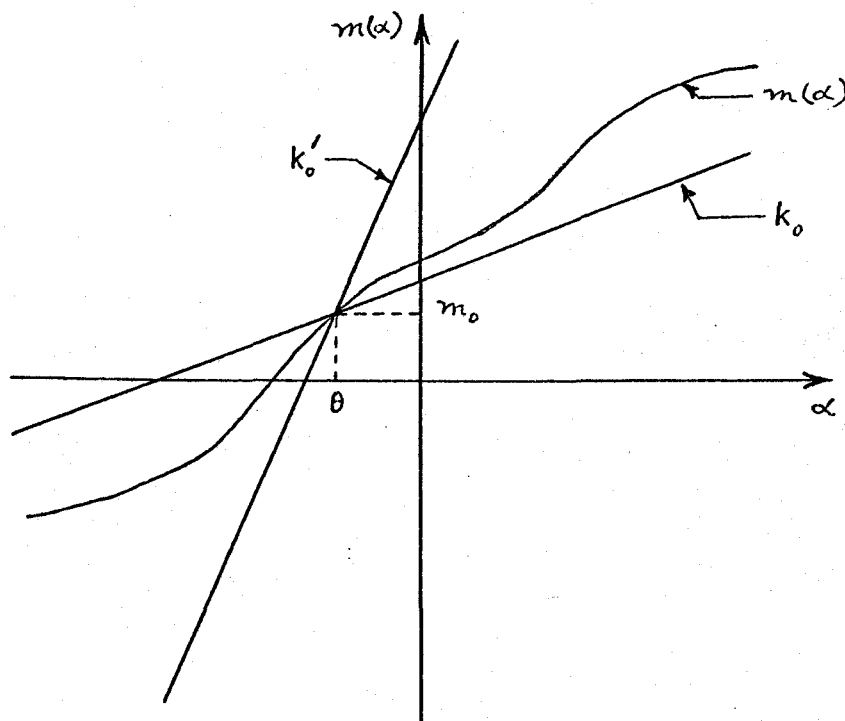


Figure (2.2.1) Schematic illustration of first assumption for Robbins - Monro Technique

Now suppose that the function $m(\alpha)$ could be observed directly for choices of α . Then for a function $m(\alpha)$ satisfying this assumption, the following method could be used to find θ . First choose α_1 arbitrarily, and observe $m(\alpha_1)$. If it is not equal to m_0 , then add a correction to α_1 of the form $-a_1\{m(\alpha_1)-m_0\}$. Then make an observation at

$$\alpha_2 = \alpha_1 - a_1\{m(\alpha_1) - m_0\} \quad (2.2.5)$$

and again make a similar correction but this time weighted by a_2 . This scheme could be continued until one approaches arbitrarily close to θ .

Now consider a modification of this method to the situation where the function $m(\alpha)$ cannot be observed but only the random variable

$$Y_n(\alpha) = f(Z_n, \alpha) \quad (2.2.6)$$

whose expected value is $m(\alpha)$. In similar fashion, again select an α_1 arbitrarily, and generate a sequence of estimates α_n by the recursion relation,

$$\alpha_{n+1} = \alpha_n + a_n\{Y_n(\alpha_n) - m_0\} \quad n=1,2,\dots \quad (2.2.7)$$

alternately written as

$$\alpha_{n+1} = \alpha_n + a_n\{f(Z_n, \alpha) - m_0\} \quad n=1,2,\dots \quad (2.2.8)$$

Here α_n is a sequence of non-stationary random variables converging to θ under certain assumptions.

More assumptions regarding the observable entities Z_n and the sequence α_n are given as follows:

Assumption 2: The random elements Z_n , $n=1,2,\dots$, are identically distributed and statistically independent.

Assumption 3: For all values of α the variance must be finite, that is,

$$\text{var } \{Y_n(\alpha)\} = \text{var } \{f(Z_n, \alpha)\} \leq \sigma^2 < \infty \quad (2.2.9)$$

Assumption 4: The sequence of weights a_n are positive monotone decreasing with

$$\sum_{n=1}^{\infty} a_n = \infty \quad (2.2.10)$$

and

$$\sum_{n=1}^{\infty} a_n^2 < \infty \quad (2.2.11)$$

The Robbins - Monro theorem states; based on the assumptions given, the sequence of estimates α_n approaches the true value θ in the mean square sense.

$$\lim_{n \rightarrow \infty} E\{(\alpha_n - \theta)^2\} = 0 \quad (2.2.12)$$

Even though the results of this theorem appear simple the fact remains that set in the proper context, the work of Robbins and Monro is a near fundamental achievement. The

criticism of J. Wolfowitz⁵ tends to bear this out.

2.3 The Kiefer - Wolfowitz Method

Consider the same formulation as in the Robbins - Monro formulation in the previous case, except that it is now desired to find the value of the scalar α which extremizes, minimizes or maximizes, the scalar valued function $m(\alpha)$. Denote this value θ . In order for the recursive search procedure to be successful in this case, it is required that $m(\alpha)$ have only a single extremum and no flat spots where $m'(\alpha)$, the derivative of $m(\alpha)$ with respect to α , is zero other than at the extremum. In short, it is required that $m'(\alpha)$, if it exists, be restricted as $m(\alpha)$ in the Robbins - Monro technique, that is, there exist constants k_0 and k'_0 where

$$0 < k_0 \leq k'_0 < \infty \quad (2.3.1)$$

such that

$$k_0(\alpha - \theta)^2 \leq \{m'(\alpha) - m'_0\}(\alpha - \theta) \leq k'_0(\alpha - \theta)^2 \Big|_{\theta=0} \quad (2.3.2)$$

Essentially, $m'(\alpha)$ must lie between two straight lines, one of positive slope and the second of finite positive slope. (See Figure 2.3.1 for illustration of the derivative of $m(\alpha)$).

This problem could be reformulated as a Robbins - Monro problem and search for the solution of

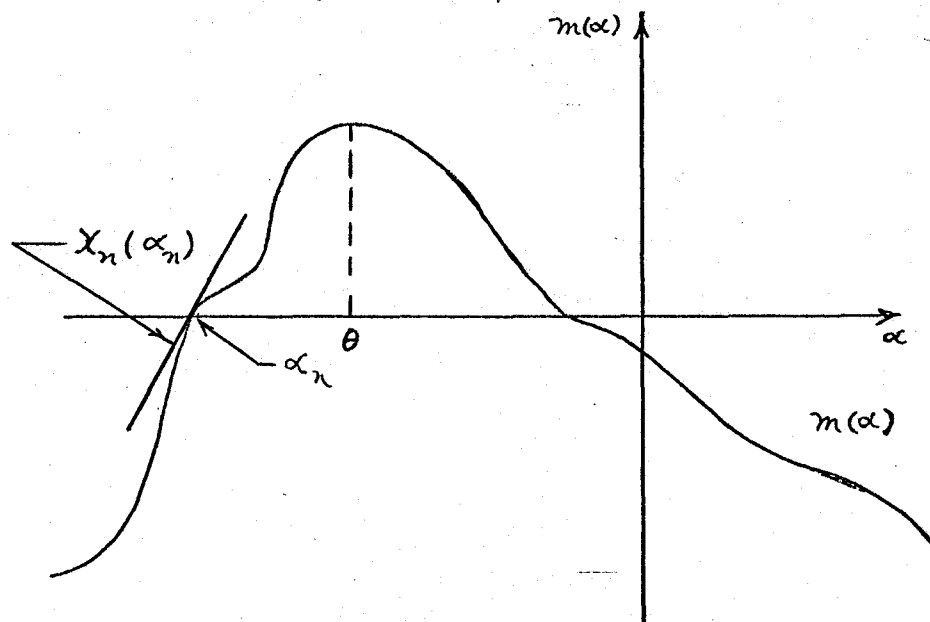


Figure 2.3.1 Illustration of requirement on derivative of $m(\alpha)$ in the Kiefer - Wolfowitz Method

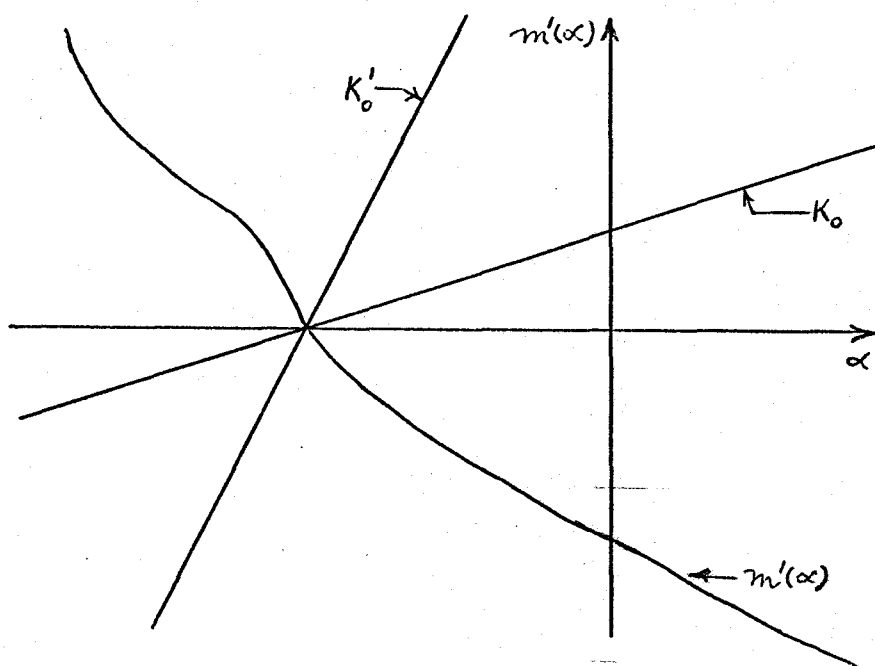


Figure 2.3.2 Locus of $m'(\alpha)$

$$m'(\alpha) = 0 \quad (2.3.3)$$

When possible, by all means, this is the best procedure. Often, however, this is neither possible nor feasible. The following two conditions state why:

- (1) It may not be possible to assume that the function $m(\alpha)$ is everywhere differentiable. Furthermore, the random variable

$$\frac{\partial}{\partial \alpha} f(Z_n, \alpha)$$

may not be well behaved; particularly it may be impossible to guarantee that

$$m'(\alpha) = \frac{\partial}{\partial \alpha} E\{f(Z_k, \alpha)\} = E\left\{\frac{\partial}{\partial \alpha} f(Z_n, \alpha)\right\} \quad (2.3.4)$$

let alone generate it in practice.

- (2) Although it may be quite simple to calculate or observe Z , and compute $f(Z, \alpha)$ the computation of the quantity

$$\frac{\partial}{\partial \alpha} f(Z_n, \alpha)$$

may be very difficult.

If either reason is valid, the Kiefer - Wolfowitz method can be applied.

Now consider at the n^{th} step of the search procedure, two observations of Z , Z_{2n-1} and Z_{2n} that are made and the two quantities

$$Y_{2n-1} = f(Z_{2n-1}, \alpha_n - c_n) \quad (2.3.5)$$

and

$$Y_{2n} = f(Z_{2n}, \alpha_n + c_n) \quad (2.3.6)$$

are calculated. The quantity

$$x_n(\alpha_n) = \frac{(Y_{2n} - Y_{2n-1})}{2c_n} \quad (2.3.7)$$

is taken as an estimate of the two-sided difference approximation to $m'(\alpha)$, namely,

$$\frac{m(\alpha_n + c_n) - m(\alpha_n - c_n)}{2c_n} \quad (2.3.8)$$

The sequence of estimates α_n is then generated by picking α_1 arbitrarily and using the recursion equation

$$\alpha_{n+1} = \alpha_n - a_n x_n(\alpha_n) \quad (2.3.9)$$

Here a_n has the same properties as in the Robbins - Monro procedure, that is, it is a sequence of positive monotone

decreasing weights with

$$\sum_{n=1}^{\infty} a_n = \infty \quad (2.3.10)$$

and

$$\sum_{n=1}^{\infty} a_n^2 < \infty \quad (2.3.11)$$

The condition on c_n in expression (2.3.8) is that it approach zero as n becomes very large so that the sequence of differences approximate a derivative more closely.

In addition, an assumption similar to Assumption 2 in the Robbins - Monro technique is made; that is, the random entities Z_n , $n=1,2,\dots$, are identically distributed and statistically independent.

Furthermore, on the function $m(\alpha)$, the following conditions are required:

Condition 1: there exist positive β and B such that

$$|\alpha' - \theta| + |\alpha'' - \theta| < \beta \text{ implies } |m(\alpha') - m(\alpha'')| < B |\alpha' - \alpha''| \quad (2.3.12)$$

Condition 2:^{*} there exist ρ and R such that

$$|\alpha' - \alpha''| < \rho \text{ implies } |m(\alpha') - m(\alpha'')| < R \quad (2.3.13)$$

Condition 3: for every $\delta > 0$, there exists a positive $\pi(\delta)$ such that

*Note condition 1 implies condition 2 but not the converse.

$$|\alpha - \theta| < \delta \text{ implies } \inf_{\frac{1}{2}\delta > \varepsilon > 0} \frac{|m(\alpha + \varepsilon) - m(\alpha - \varepsilon)|}{\varepsilon} > \pi(\delta) \quad (2.3.14)$$

Under these conditions it was shown by Kiefer and Wolfowitz that for a weighting sequence a_n and difference sequence c_n satisfying equation (2.3.10), equation (2.3.11) and the following set of conditions, namely,

$$\sum_{n=1}^{\infty} a_n c_n < \infty \quad (2.3.15)$$

and

$$\sum_{n=1}^{\infty} a_n^2 c_n^{-2} < \infty \quad (2.3.16)$$

that the sequence α_n converges to θ in probability.

2.4 The Work of A. Dvoretzky

Having seen the two initial and particular examples of stochastic approximation, namely the Robbins - Monro technique for approximating the point where a regression function assumes a given value and the Kiefer - Wolfowitz Method which finds the extremum of a regression function, it is apparent that the need for a more general theory existed. Dvoretzky³ formulates the problem in this more general sense, that is, he considers a random element such as noise superimposed on a convergent deterministic scheme. In this light he formulated and proved the following theorem:

Theorem: Let $\alpha_n(\rho_1, \dots, \rho_n)$, $\beta_n(\rho_1, \dots, \rho_n)$ and $\gamma_n(\rho_1, \dots, \rho_n)$ be non-negative measurable functions of real variables $\rho_1, \rho_2, \dots, \rho_n$, satisfying the condition that $\alpha_n(\rho_1, \dots, \rho_n)$ are bounded and that

$$\lim_{n \rightarrow \infty} \alpha_n(\rho_1, \dots, \rho_n) \rightarrow 0 \quad (2.4.1)$$

for a sequence $\rho_1, \rho_2, \dots$

The sum of the series

$$\sum_{n=1}^{\infty} \beta_n(\rho_1, \dots, \rho_n) < \infty \quad (2.4.2)$$

is bounded and converges for any sequence $\rho_1, \rho_2, \dots$

The series

$$\sum_{n=1}^{\infty} \gamma_n(\rho_1, \dots, \rho_n) = \infty \quad (2.4.3)$$

diverges for any sequence $\rho_1, \rho_2, \dots$, bounded in absolute value, that is, for any sequence $\rho_1, \rho_2, \dots$, such that

$$\sup_{n=1,2,\dots} |\rho_n| < c \quad (2.4.4)$$

c being an arbitrary finite number. Let θ be a real number, and $T_1, T_2, \dots$, be measurable transformations, satisfying the inequality

$$|T_n(\rho_1, \rho_2, \dots, \rho_n) - \theta| \leq \max \{ \alpha_n, (1 + \beta_n) |\rho_n - \theta| - \gamma_n \} \quad (2.4.5)$$

for any real sequence $\rho_1, \rho_2, \dots$. Further let X_1 and $Y_1, Y_2, \dots$, be random variables, and for $n \geq 1$ let

$$X_{n+1} = T_n (X_1, \dots, X_n) + Y_n + g_n (X_1, \dots, X_n) \quad (2.4.6)$$

where $g_n (r_1, \dots, r_n)$ are measurable functions such that the series

$$\sum |g_n (r_1, r_2, \dots, r_n)| \quad (2.4.7)$$

uniformly converges and its sum is uniformly bounded for any $r_1, r_2, \dots$.

Let

$$E \{Y_n | X_1, X_2, \dots, X_n\} = 0 \quad (2.4.8)$$

with probability 1. Let the series

$$\sum_{n=1}^{\infty} E \{Y_n^2\} < \infty \quad (2.4.9)$$

and let

$$E \{X_1^2\} < \infty \quad (2.4.10)$$

Then as $n \rightarrow \infty$

$$P \{ \lim_{n \rightarrow \infty} X_n = \theta \} = 1 \quad (2.4.11)$$

and

$$\lim_{n \rightarrow \infty} E \{ (X_n - \theta)^2 \} = 0 \quad (2.4.12)$$

Along with this basic theorem, Dvoretzky also proved an extension and five generalizations.

2.5 Generalizations and Applications

Having developed the theory to this level, the work begun by the three groups of researchers, namely Robbins and Monro, Kiefer and Wolfowitz, and Dvoretzky, was extended and modified by a number of people. An attempt will be made to give a survey of first, the generalizations which stemmed from the work outlined to this point and also some of the applications where this theory has been used.

The first extension of the Robbins - Monro technique was made by Wolfowitz⁵. He showed that under weaker assumptions than required by Robbins and Monro there was still convergence in probability to the root. Further, Blum⁶ showed that under still weaker conditions there was convergence in probability and even convergence with probability 1. In the same paper Blum showed that for weakened conditions in the Kiefer - Wolfowitz method convergence could be strengthened to convergence with probability 1. In a concurrent publication, Blum⁷ extended the Robbins - Monro and Kiefer - Wolfowitz techniques to the multidimensional case. He dealt with vector valued parameters α where the function of the vector could now be interpreted as planes in a hyperspace.

Another area of exploration has been the rates of convergence and their dependence upon the sequence a_n in the Robbins - Monro technique, and the sequence a_n and c_n in the Kiefer - Wolfowitz method. In addition to this, regularity properties of $m(\alpha)$ were investigated. The key figures in this area were Chung⁸, Derman⁹, Burkholder¹⁰, Sacks¹¹ and Dupac¹². The work of these people has been diverse in nature and of major importance to the applications of Robbins - Monro and Kiefer - Wolfowitz methods. In particular Sacks¹¹ indicates a method for selecting a weighting sequence a_n for the Robbins - Monro Technique and Dupac¹² similarly suggests sequences for the Kiefer - Wolfowitz method.

The generalizations of Dvoretzky which were also extended to non-independent observations were further extended by Sakrison¹³. Essentially Sakrison's¹³ work was an attempt to formulate conditions that are more suitable for practical work. Driml and Nedoma¹⁴ worked in the same vein but tried to extend the one-dimensional scalar case of the Robbins - Monro Technique to the continuous time case. The most extensive attempts to bring these methods to practical applications have been made by Alberts and Gardner¹⁵. They have attempted to make practical choices of the weighting sequences a_n in the Robbins - Monro technique.

As it became evident that stochastic approximation could provide a useful tool to the engineer, attempts were

made to apply the techniques to physical problems. In the realm of electrical engineering, numerous applications were tried and documented. The Robbins - Monro Technique was applied to the problem of parameter estimation in radar and radio astronomy where signals involved bandwidths of from 100 HZ. on up and the total observation time was quite long. As such, the amount of data to be processed was quite large and hence the stochastic approximation technique chosen was quite appropriate. Sackrison^{16, 17} outlines the theoretical and practical details in two papers. Sackrison^{18, 19} also looked at the optimization of filters and detectors. He applied the Kiefer - Wolfowitz method in order to ascertain two advantages over conventional methods;

- (1) The error weight could be more general than square error.
- (2) The method combines the processes of measuring statistics and solving filter equations into a single compact algorithm.

Similar work has been done by Kushner²⁰ who used the Robbins - Monro Technique in filter design considering mean square error, additive signal and noise.

It is thus evident, that even though the theory of stochastic approximation is quite restrictive, requiring regularity conditions as well as bounds on weighting sequences, its implementation to some areas of electrical

engineering has already begun. As the theory expands the future augurs well for the application of stochastic approximation techniques to engineering problems.

CHAPTER III

Development of a Stochastic Approximation Algorithm

3.1 Background

The problem of stochastic approximation has been mentioned by Sklansky²¹ within the framework of *learning* control systems. He has associated the term *learning* control with the hierarchic arrangement of three feedback loops. The first loop contains a controller or "compensator" in a simple feedback configuration. The second or "adaptive" loop contains a "system identifier" or pattern recognizer that adjusts the compensator in response to changes in the estimated dynamic parameters of the plant. The third loop or "learning" loop contains a teacher--a type of controller--which "trains" the "pattern recognizer" to make optimum or near optimum recognitions. Based on a stored set of past controls used in conjunction with a given recognition, a control is initiated by the teacher as either of two forms. If there is no previous record or "experience" of the given situation, then the adaptive loop performs an adaptive control procedure and stores the situation, control and results. If there had been record of a similar condition in its performance history, the teacher would have selected the corresponding control policy and initiated execution.

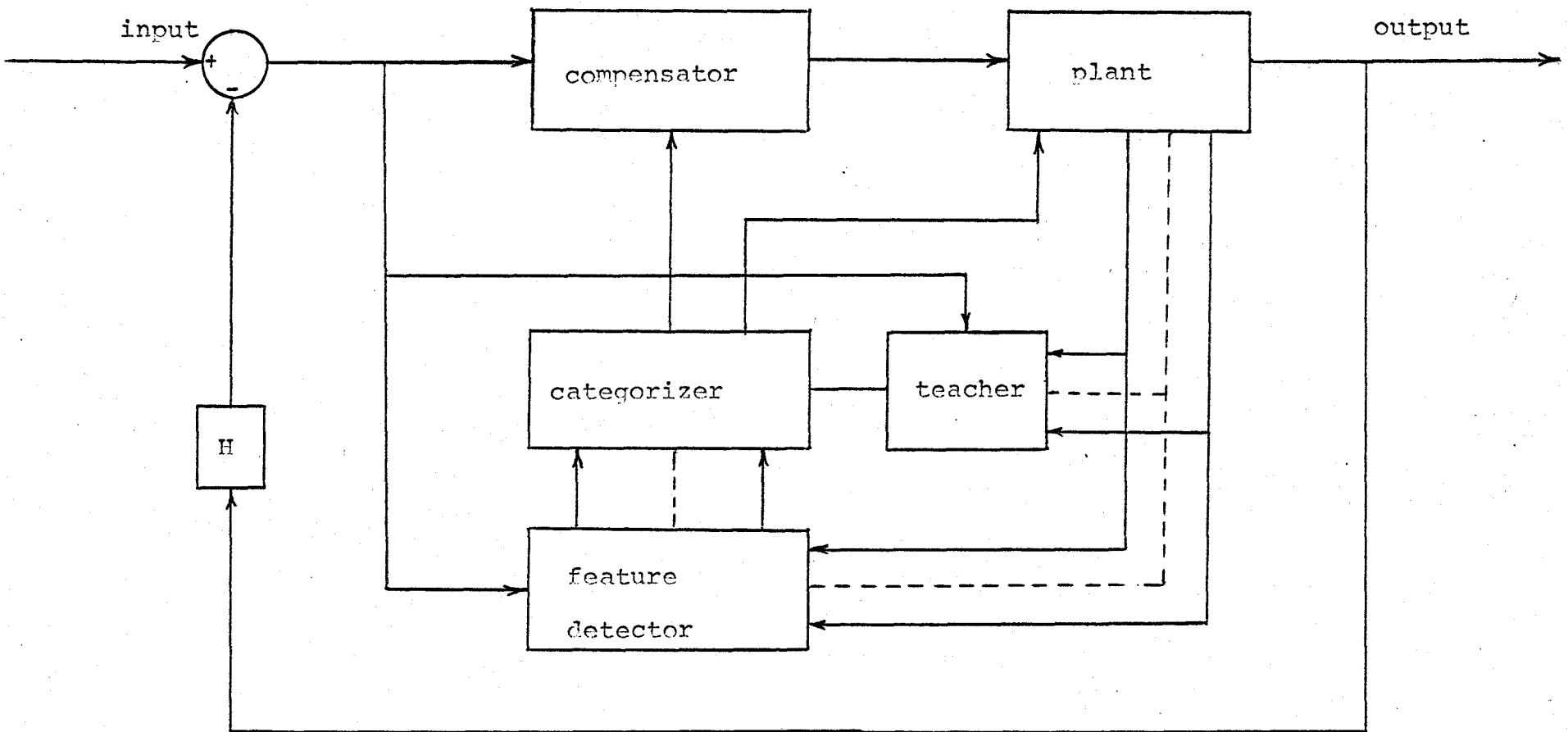


Figure 3.1 Schematic illustration of Learning Control System

This would have alleviated the need for an adaptive control. Essentially, once an adaption in a given situation has been performed, there is no need to recompute it if the adaptor is a part of a learning system; the results of the first adaption, which have been stored, are applied to the same control situation when it re-occurs.

The method by which the "teacher" associated the recognized situation of plant parameters and input signal--called control situation--with a given control successfully adapted at the previous occurrence of the same situation is via reinforcement probabilities. These are probabilities that are assigned to all possible control possibilities for a given control situation and are up-dated ie. reinforced or penalized, based on the outcome of an executed control policy.

There is an intimate relationship between stochastic approximation methods and the reinforcement learning technique just outlined. In fact, it can be shown* that the two are essentially the same and that the reinforcement learning process is an example of an application of stochastic approximation theory. In addition, stochastic approximation techniques can be used to estimate a cost function, plant parameters, input signals and optimal control policies as pertain to the particular area of

*See Appendix A for proof of relationship of reinforced learning and stochastic approximation

control cited here. Further applications will be outlined later in this work.

3.2 Convergence

Direct reference has been made to the two algorithms developed by K. S. Fu et al. By way of introduction, the two algorithms will be presented and categorized. They are of the Dvoretzky type being a special class as concerns their proof**.

Consider an algorithm of the form

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{f(r_{n+1}) - \hat{x}_n\} \quad n=1,2,\dots \quad (3.2.1)$$

to be used in the presence of an ergodic process ξ_n ,

where $\hat{x}_n$ is the n^{th} estimate of x ,

where x is the true value of the parameter being estimated,

r_n is the n^{th} sample taken and used to calculate,

$f(r_n)$ a function of the samples,

and γ_n is a gain sequence.

It is required of the function $f(r_n)$ that

$$E \{f(r_n)\} = x \quad (3.2.2)$$

The form (3.2.1) has been used in two specific ways. The first algorithm of Fu uses

$$f(r_n) = r_n \quad n=1,2,\dots \quad (3.2.3)$$

where the samples

**see Appendix B for proof of first two stochastic approximation algorithms

$$r_n = x + \xi_n, \quad \xi_n \text{ being a random component of zero mean noise,} \quad (3.2.4)$$

and

$$\gamma_n = \frac{1}{n+\alpha} \quad n=1,2,\dots \quad (3.2.5)$$

where α is a constant. The selection of the constant α is arbitrary; but, if *a priori* statistics are known and the process is known to be normally distributed then

$$\alpha = \frac{v_0^2}{\sigma^2} \quad (3.2.6)$$

where σ^2 is the variance of the distribution

and v_0^2 is the initial value of expected mean square error.

Selecting γ_n and α as in (3.2.5) and (3.2.6) respectively, gives the form (3.2.1) the best convergence when using the function relationship (3.2.3)

The second algorithm of Fu uses

$$f(r_n) = \frac{1}{n} \sum_{i=1}^n r_i \quad n=1,2,\dots \quad (3.2.7)$$

where the samples

$$r_n = x + \xi_n, \quad \xi_n \text{ again being a random component of zero mean noise,} \quad (3.2.8)$$

and

$$\gamma_n = \frac{n}{n(n+1) + \alpha} \quad n=1,2,\dots \quad (3.2.9)$$

where α is a constant. The selection of the constant α is arbitrary; but, when *a priori* statistics are known and the process is known to be normally distributed then

$$\alpha = \frac{v_0^2}{\sigma^2} \quad (3.2.10)$$

where σ^2 is the variance of the distribution and v_0^2 is the initial value of the expected mean square error.

Selecting γ_n and α as in (3.2.9) and (3.2.10) respectively gives the form (3.2.1) the best convergence when using the functional relationship (3.2.7).

Now the first algorithm of Fu

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{r_{n+1} - \hat{x}_n\} \quad n=1,2,\dots \quad (3.2.11)$$

with γ_{n+1} as in (3.2.5) and the second algorithm of Fu

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \left\{ \frac{1}{n+1} \sum_{i=1}^{n+1} r_i - \hat{x}_n \right\} \quad n=1,2,\dots \quad (3.2.12)$$

with γ_{n+1} as in (3.2.9) will be compared with the algorithm to be developed here.

3.3 Development

Consider an algorithm of the Dvoretzky type of the form

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{f(r_{n+1}) - \hat{x}_n\} \quad n=1,2,\dots \quad (3.3.1)$$

similar to (3.2.1). The algorithm (3.3.1) is to be used in the presence of an ergodic process ξ_n ,

where $\hat{x}_n$ is the n^{th} estimate of x ,

where x is the true value of the parameter being estimated,

r_n is the n^{th} sample taken and used to calculate $f(r_n)$ a function of the samples,

and γ_n is a gain sequence.

It should be noted that the process of taking the samples r_n gives

$$r_n = x + \xi_n \quad (3.3.2)$$

where the true parameter x that is being sought is contaminated by the zero mean ergodic process ξ_n .

It is suggested that the function f be chosen as

$$f(r_{n+1}) = [\hat{R}_{r_{n+1}}(\ell)]^{\frac{1}{2}} \quad (3.3.3)$$

where $\hat{R}_{r_{n+1}}(\ell)$ is an estimate of the sample autocorrelation function of the samples $r_1, r_2, \dots, r_{n+1}$ with $\ell-1 \leq n$.

From the definition of autocorrelation

$$R(\ell) = E \{ r_n r_{n+\ell} \} \quad (3.3.4)$$

and if (3.3.2) is recalled

$$r_n = x + \xi_n$$

then combining these two equations gives

$$R(\ell) = E \{ [x + \xi_n] [x + \xi_{n+\ell}] \} \quad (3.3.5)$$

Expanding the product of binomials and rearranging the expectation operator gives,

$$R(\ell) = x^2 + E \{ x \xi_n \} + E \{ x \xi_{n+\ell} \} + E \{ \xi_n \xi_{n+\ell} \} \quad (3.3.6)$$

Now since x is a constant and since ξ_n and $\xi_{n+\ell}$ are from a zero mean ergodic process, then

$$E \{ x \xi_n \} = x E \{ \xi_n \} \quad (3.3.7)$$

and

$$E \{ x \xi_{n+\ell} \} = x E \{ \xi_{n+\ell} \} \quad (3.3.8)$$

and these two terms vanish. Hence, if an estimate of the sample autocorrelation function is taken it gives

$$\hat{R}_{r_{n+1}}(\ell) = x^2 + \hat{R}_{\xi_{n+1}}(\ell) \quad (3.3.9)$$

where $\hat{R}_{\xi_{n+1}}(\ell)$ is an estimate of the sample autocorrelation function of the random noise elements only.

Hence (3.3.1) can be written as

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{ [x^2 + \hat{R}_{\xi_{n+1}}(\ell)]^{\frac{1}{2}} - \hat{x}_n \} \quad n=1,2,\dots \quad (3.3.10)$$

Now subtracting x from both sides and rearranging terms gives

$$(\hat{x}_{n+1} - x) = (1 - \gamma_{n+1}) (\hat{x}_n - x) + \gamma_{n+1} \psi_{n+1} \quad (3.3.11)$$

where

$$\psi_{n+1} = x \left\{ \left(1 + \frac{\hat{R}_{\xi_{n+1}}^{(l)}}{x^2} \right)^{\frac{1}{2}-1} \right\} \quad (3.3.12)$$

If (3.3.11) is iterated then an expansion can be developed in terms of the initial mean square error. After the first iteration, the equation (3.3.11) is of the form

$$\hat{(x_{n+1}-x)} = (1-\gamma_{n+1})(1-\gamma_n)\hat{(x_{n-1}-x)} + (1-\gamma_{n+1})\gamma_n\psi_n + \gamma_{n+1}\psi_{n+1} \quad (3.3.13)$$

After the second iteration (3.3.13) will appear as

$$\begin{aligned} \hat{(x_{n+1}-x)} &= (1-\gamma_{n+1})(1-\gamma_n)(1-\gamma_{n-1})\hat{(x_{n-2}-x)} \\ &\quad + (1-\gamma_{n+1})(1-\gamma_n)\gamma_{n-1}\psi_{n-1} \\ &\quad + (1-\gamma_{n+1})\gamma_n\psi_n + \gamma_{n+1}\psi_{n+1} \end{aligned} \quad (3.3.14)$$

Repeating this procedure $n+1$ times gives the following general form:

$$\begin{aligned} \hat{(x_{n+1}-x)} &= (1-\gamma_{n+1})(1-\gamma_n)\dots(1-\gamma_2)(1-\gamma_1)\hat{(x_0-x)} \\ &\quad + (1-\gamma_{n+1})(1-\gamma_n)\dots(1-\gamma_2)\gamma_1\psi_1 \\ &\quad + (1-\gamma_{n+1})(1-\gamma_n)\dots(1-\gamma_3)\gamma_2\psi_2 \\ &\quad + \dots + (1-\gamma_{n+1})(1-\gamma_n)\gamma_{n-1}\psi_{n-1} \\ &\quad + (1-\gamma_{n+1})\gamma_n\psi_n \\ &\quad + \gamma_{n+1}\psi_{n+1} \end{aligned} \quad (3.3.15)$$

Now (3.3.15) can be rewritten into a closed form;

$$(\hat{x}_{n+1}-x) = \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\} (\hat{x}_0-x) + \sum_{i=n+1}^1 \gamma_i \psi_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \quad (3.3.16)$$

$n=1,2,\dots$

where all void products are taken as unity.

Changing the order of the first term and the limits of summation in the second term gives,

$$(\hat{x}_{n+1}-x) = (\hat{x}_0-x) \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\} + \sum_{i=1}^{n+1} \gamma_i \psi_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \quad (3.3.17)$$

$n=1,2,\dots$

Squaring (3.3.17) gives;

$$\begin{aligned} (\hat{x}_{n+1}-x)^2 &= (\hat{x}_0-x)^2 \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\}^2 + \left[\sum_{i=1}^{n+1} \gamma_i \psi_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \right]^2 \\ &\quad + 2(\hat{x}_0-x) \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\} \cdot \left[\sum_{i=1}^{n+1} \gamma_i \psi_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \right] \end{aligned} \quad (3.3.18)$$

$n=1,2,\dots$

Substituting from (3.3.12) for ψ_i gives, after some rearrangement;

$$\begin{aligned} (\hat{x}_{n+1}-x)^2 &= (\hat{x}_0-x)^2 \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\}^2 + x^2 \left[\sum_{i=1}^{n+1} \gamma_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \right]^2 \\ &\quad + \left[\sum_{i=1}^{n+1} \gamma_i (x^2 + \hat{R}_{\xi_i}(\ell))^{1/2} \prod_{j=i+1}^{n+1} (1-\gamma_j) \right]^2 \\ &\quad + 2(x_0-x) \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\} \left[\sum_{i=1}^{n+1} \gamma_i (x^2 + \hat{R}_{\xi_i}(\ell))^{1/2} \prod_{j=i+1}^{n+1} (1-\gamma_j) \right] \end{aligned}$$

$$\begin{aligned}
& - 2x(\hat{x}_0 - x) \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\} \left[\sum_{i=1}^{n+1} \gamma_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \right] \\
& - 2x \left\{ \sum_{i=1}^{n+1} \gamma_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \right\} \left[\sum_{i=1}^{n+1} \gamma_i (x^2 + \hat{R}_{\xi_i}(\ell))^{\frac{1}{2}} \prod_{j=i+1}^{n+1} (1-\gamma_j) \right] \\
& n=1, 2, \dots \quad (3.3.19)
\end{aligned}$$

In general

$$\hat{R}_{\xi_{n+1}}(\ell) = E \{ \xi_k \xi_{k+\ell} \} = \sigma^2 e^{-\alpha |\ell|} \quad (3.3.20)$$

for coloured Gaussian noise. The wider the frequency spectrum the smaller the value of α and hence the smaller the value of $\hat{R}_{\xi_{n+1}}(\ell)$ for the same non zero value of ℓ . As the spectrum becomes wider--the noise approaches white-- α approaches zero and $\hat{R}_{\xi_{n+1}}(\ell)$ becomes an impulse,

$$\hat{R}_{\xi_{n+1}}(\ell) = \sigma^2 \delta(\ell) \quad (3.3.21)$$

and for $\ell \neq 0$ & $\ell > 1$,

$$\hat{R}_{\xi_{n+1}}(\ell) \rightarrow 0 \quad (3.3.22)$$

Now, if in (3.3.19)

$$(x^2 + \hat{R}_{\xi_i}(\ell))^{\frac{1}{2}} = x \left(1 + \frac{\hat{R}_{\xi_i}(\ell)}{x^2} \right)^{\frac{1}{2}} \quad x \neq 0 \quad (3.3.23)$$

is used as a substitution, (3.3.19) can be rewritten as

$$\begin{aligned}
(\hat{x}_{n+1}-x)^2 &= (\hat{x}_0-x)^2 \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\}^2 + x^2 \left[\sum_{i=1}^{n+1} \gamma_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \right]^2 \\
&+ x^2 \left[\sum_{i=1}^{n+1} \gamma_i \left(1 + \frac{\hat{R}_{\xi i}(\ell)}{x^2} \right)^{\frac{1}{2}} \prod_{j=i+1}^{n+1} (1-\gamma_j) \right]^2 \\
&+ 2x(\hat{x}_0-x) \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\} \left[\sum_{i=1}^{n+1} \gamma_i \left(1 + \frac{\hat{R}_{\xi i}(\ell)}{x^2} \right)^{\frac{1}{2}} \prod_{j=i+1}^{n+1} (1-\gamma_j) \right] \\
&- 2x(\hat{x}_0-x) \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\} \left[\sum_{i=1}^{n+1} \gamma_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \right] \\
&- 2x^2 \left\{ \sum_{i=1}^{n+1} \gamma_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \right\} \left[\sum_{i=1}^{n+1} \gamma_i \left(1 + \frac{\hat{R}_{\xi i}(\ell)}{x^2} \right)^{\frac{1}{2}} \prod_{j=i+1}^{n+1} (1-\gamma_j) \right]
\end{aligned}$$

n=1,2,... (3.3.24)

For the case of white noise or even slightly coloured Gaussian noise, substituting from (3.3.22) for $\hat{R}_{\xi i}(\ell)$ in (3.3.24) reduces it to the following form:

$$(\hat{x}_{n+1}-x)^2 = (\hat{x}_0-x)^2 \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\}^2 \quad (3.3.25)$$

n=1,2,...

Therefore, for the mean square error to become zero, regardless of the starting value, and for the algorithm to converge, requires only that the following limit exist:

$$\lim_{n \rightarrow \infty} \prod_{i=1}^{n+1} (1-\gamma_i) \rightarrow 0 \quad (3.3.26)$$

The selection of a gamma sequence that satisfies (3.3.26) is arbitrary within the restrictions of the limit given above.

The type of Γ -sequences used in the thesis are of two

basic types. These will be introduced, discussed and their effect on convergence illustrated.

Hence, form (3.3.1) can be written as

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{ (\hat{R}_{r_{n+1}}(\ell))^{\frac{1}{2}} - \hat{x}_n \} \quad n=1,2,\dots \quad (3.3.27)$$

and will converge to the true value being sought if the measurements, r_n , are of the form (3.3.2) and the Γ -sequence satisfies (3.3.26)

3.4 Gamma Sequence

The selection of the Γ -sequence for (3.3.27) essentially regulates the rate of convergence of the algorithm. The optimization of this rate will be discussed in the next chapter. The present form of the Γ -sequences selected will be given and it will be shown how they satisfy (3.3.26) and hence (3.3.27) converges.

The first Γ -sequence is given by

$$\Gamma = \{ \gamma_i \in \Gamma \mid \gamma_i = \frac{1}{i+\alpha} \quad \forall \quad i = 1,2,\dots \} \quad (3.4.1)$$

for any non-negative arbitrary constant α . Hence,

$$\begin{aligned} (1-\gamma_i) &= 1 - \frac{1}{i+\alpha} \\ &= \frac{i+\alpha-1}{i+\alpha} \end{aligned} \quad (3.4.2)$$

Taking the continued product of (3.4.2) gives,

$$\prod_{i=r}^n (1-\gamma_i) = \prod_{i=r}^n \frac{i+\alpha-1}{i+\alpha}$$

This expands to,

$$\prod_{i=r}^n (1-\gamma_i) = \left(\frac{r+\alpha-1}{r+\alpha}\right) \left(\frac{r+\alpha}{r+\alpha+1}\right) \cdots \left(\frac{n+\alpha-2}{n+\alpha-1}\right) \left(\frac{n+\alpha-1}{n+\alpha}\right)$$

which reduces to,

$$\prod_{i=r}^n (1-\gamma_i) = \frac{r+\alpha-1}{n+\alpha} \quad (3.4.3)$$

Now if a limit is taken

$$\lim_{n \rightarrow \infty} \prod_{i=r}^n (1-\gamma_i) = \lim_{n \rightarrow \infty} \frac{r+\alpha-1}{n+\alpha}$$

giving

$$\lim_{n \rightarrow \infty} \prod_{i=r}^n (1-\gamma_i) \rightarrow 0 \quad (3.4.4)$$

Hence for a Γ -sequence defined by (3.4.1) the statement of the limit (3.4.4) satisfies (3.3.26) and hence the algorithm (3.3.27) using said definition converges to the true value being sought regardless of the starting value-- provided that the starting value is finite.

The second Γ -sequence is given by

$$\Gamma = \{\gamma_i \in \Gamma \mid \gamma_i = 1 - \frac{1}{i^p} \forall i=1,2,\dots; p=1,2,\dots\} \quad (3.4.5)$$

From this Γ -sequence

$$(1-\gamma_i) = \frac{1}{i^p} \quad (3.4.6)$$

Taking the continued product of (3.4.6) gives,

$$\prod_{i=r}^n (1-\gamma_i) = \prod_{i=r}^n \frac{1}{i^p} \quad (3.4.7)$$

Now if the limit is taken

$$\lim_{n \rightarrow \infty} \prod_{i=r}^n (1 - \gamma_i) = \lim_{n \rightarrow \infty} \prod_{i=r}^n \frac{1}{ip}$$

giving

$$\lim_{n \rightarrow \infty} \prod_{i=r}^n (1 - \gamma_i) \rightarrow 0 \quad (3.4.8)$$

Hence for a Γ -sequence defined by (3.4.5) the statement of the limit (3.4.8) satisfies (3.3.26) and hence the algorithm (3.3.27) using said definition converges to the true value being sought regardless of the starting value--provided that the starting value is finite.

3.5 Consistency and Bias

An estimator should not be considered bad simply because it can assume a value that deviates considerably from the true value being sought. But if the bulk of the values of the estimator deviate considerably from the true value, the estimator can be considered bad, particularly, if a large sampling has been taken. Hence, a desirable property is that there be a high probability that the estimator be near the parameter it is intended to estimate for large sample sizes.

By definition, an estimator $\hat{x}_n$ of x is said to be a *consistent* estimator, if for any positive numbers δ and ϵ there exists an integer N such that the probability that

$$|\hat{x}_n - x| < \epsilon$$

is greater than $1-\delta$ for all $n > N$; that is

$$P\{|\hat{x}_n - x| < \epsilon\} > 1-\delta \quad \forall n > N \quad (3.5.1)$$

The definition is similar to the definition of convergence in the mathematical sense, except that here it is said that, given any small ϵ , a sample size can be found large enough so that, for all larger sample sizes, the probability that $\hat{x}_n$ differs from the true value x more than ϵ is as small as desired. In such a case $\hat{x}_n$ converges in probability to x . So convergence in probability means that $\hat{x}_n$ is a consistent estimator of x .

The criterion of consistency is not very practical sometimes, since it has to do with a limiting property. There are two fundamental facts which pertain to this fact. First, samples have a finite number of observations while the definition of consistency requires an infinite number. Second, when there is one consistent estimator $\hat{x}_n$ of θ , it is possible to have infinitely many. For example, if $\hat{x}_n$ is consistent, so is

$$\frac{n+a}{n+b} \cdot \hat{x}_n$$

for all fixed numbers $a \neq -n$ and $b \neq -n$.

Now if the restriction "for large n " is removed, a selection from among all consistent estimators can be made

resulting in a much smaller class by applying the definition of an unbiased estimator.

By definition, an estimator $\hat{x}_n$ of x is unbiased if

$$E \{ \hat{x}_n \} = x \quad (3.5.2)$$

This definition applies for all n and x , and requires that the mean of the sampling distribution of any statistic equal the parameter which the statistic is supposed to estimate.

Now for the class of algorithms presented here, consider for a moment only the sample autocorrelation function given by

$$\hat{R}_{r_n}(\ell) = \frac{1}{n} \sum_{k=1}^{n-\ell} (r_k r_{k+\ell}) \quad (3.5.3)$$

Recalling from (3.3.2) that

$$r_n = x + \xi_n$$

and substituting in (3.5.3) for r_k gives

$$\hat{R}_{r_n}(\ell) = \frac{1}{n} \sum_{k=1}^{n-\ell} (x^2 + \xi_k x + \xi_{k+\ell} x + \xi_k \xi_{k+\ell})$$

which can be rewritten as

$$\hat{R}_{r_n}(\ell) = x^2 + \frac{x}{n} \sum_{k=1}^{n-\ell} \xi_k + \frac{x}{n} \sum_{k=1}^{n-\ell} \xi_{k+\ell} + \frac{1}{n} \sum_{k=1}^{n-\ell} \xi_k \xi_{k+\ell} \quad (3.5.4)$$

Since ξ_i is an element of zero mean Gaussian white noise,

then for large n , the first and second summation in (3.5.4)

tends to zero; that is,

$$\frac{x}{n} \sum_{k=1}^{n-\ell} \xi_k \rightarrow 0, \quad (3.5.5)$$

$$\frac{x}{n} \sum_{k=1}^{n-\ell} \xi_{k+\ell} \rightarrow 0, \quad (3.5.6)$$

$$\text{and} \quad \frac{1}{n} \sum_{k=1}^{n-\ell} \xi_k \xi_{k+\ell} \rightarrow \delta(\ell) \quad (3.5.7)$$

Now if $\ell \neq 0$, then (3.5.4) reduces to

$$\hat{R}_{r_n}(\ell) = x^2, \quad n \gg 1 \quad (3.5.8)$$

Now recall (3.3.17)

$$(\hat{x}_{n+1} - x) = (\hat{x}_0 - x) \left\{ \prod_{i=1}^{n+1} (1 - \gamma_i) \right\} + \sum_{i=1}^{n+1} \gamma_i \psi_i \prod_{j=i+1}^{n+1} (1 - \gamma_j) \quad n=1, 2, \dots$$

and (3.3.12) rewritten as

$$\psi_{n+1} = [x^2 + R_{\xi_{n+1}}(\ell)]^{\frac{1}{2} - x} \quad (3.5.9)$$

If, to the algorithm

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{ (\hat{R}_{r_{n+1}}(\ell))^{\frac{1}{2} - \hat{x}_n} \} \quad (3.5.10)$$

written in the form (3.3.17) is applied the transformation

$$\hat{\theta}_n = (\hat{x}_n - x) \quad (3.5.11)$$

then (3.3.17) can be rewritten as

$$\hat{\theta}_{n+1} = \hat{\theta}_0 \left\{ \prod_{i=1}^{n+1} (1 - \gamma_i) \right\} + \sum_{i=1}^{n+1} \gamma_i \psi_i \prod_{j=i+1}^{n+1} (1 - \gamma_j) \quad n=1, 2, \dots \quad (3.5.12)$$

which is a zero seeking algorithm. Substituting from (3.5.8) for $\hat{R}_n(\ell)$ into (3.5.9) and substituting the result in (3.5.10) for ψ_n gives, after taking expectations of both sides,

$$\lim_{n \rightarrow \infty} \left[E \{ \hat{\theta}_{n+1} \} \right] = 0 \quad (3.5.13)$$

Now applying the transformation of equation (3.5.10) in reverse, gives

$$E \{ \hat{x}_{n+1} \} = x \quad (3.5.14)$$

Hence, it can now be said the algorithm

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{ (\hat{R}_{n+1}(\ell))^{1/2} - \hat{x}_n \} \quad n=1,2,\dots$$

is unbiased for large n .

CHAPTER IV

Optimization of Convergence

4.1 Preamble for Optimization

The convergence of the algorithm of the form

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{f(r_{n+1}) - \hat{x}_n\} \quad (4.1.1)$$

with

$$f(r_{n+1}) = (\hat{R}_{r_{n+1}}(\ell))^{\frac{1}{2}} \quad (4.1.2)$$

has been proven in the previous chapter. The most important factor to be considered next is the rate of convergence of this algorithm. In this algorithm $f(r_{n+1})$ has been chosen. Also, the Γ -sequence has been chosen to satisfy condition (3.3.26), a necessary condition. Beyond that the choice of the Γ -sequence is theoretically arbitrary. The theme, then, is to select the Γ -sequence so that the convergence of (4.1.1) is optimal in some sense and subject to the constraint (3.3.26) which is the requirement for convergence. In other words, choose a Γ -sequence to minimize a cost functional while still retaining a convergent algorithm by satisfying condition (3.3.26).

4.2 The Discrete Maximum Principle for Optimization

The discrete maximum principle formulated by Katz³¹ from the continuous maximum principle of Pontryagin, provides a method of obtaining an optimal solution for very general dynamical processes. It treats the optimization problem of minimizing or maximizing a functional subject to certain constraints. The beauty of this principle is that it is not restricted to dynamical processes. Any problem which can be formulated within the framework of state space and for which a cost function can be written in an analytic form can be approached with the maximum principle. For all problems a state space can be defined so that the equations describing it can be written in a standard form. It is often somewhat more difficult to write a cost function since this requires an intimate knowledge of the problem. In addition, a cost function must take into account the objectives to be achieved and the level of penalties to be given if the objectives are not pursued.

In general, an optimum control problem can be transformed into the problem of minimizing a function such as an inner product

$$M = \langle \underline{b}, \underline{x}(t_k) \rangle_f \quad (4.2.1)$$

subject to certain constraining functionals. The control

strategy which minimizes (or maximizes) this function is referred to as the optimum control strategy. In equation (4.2.1), $\underline{x}$ is a state vector of the n^{th} order process under consideration, and $\underline{b}$ is a column vector which depends upon the coordinates to be minimized (or maximized). It is interesting to note that this class of problems is contained within the framework of the Mayer problem in the calculus of variations. A simple geometrical interpretation of the maximum principle is that the control vector $\underline{u}$ is chosen in such a way that the state vector $\underline{x}(t_{k_f})$ moves "farthest" in the direction of $-\underline{b}$, and thus the scalar function M takes on a minimum value.

Suppose that a process under consideration can be characterized by equation (3.2.2)

$$\underline{x}_{k+1} = \underline{f}_k(\underline{x}_k, \underline{u}_k, k) \quad k=1, 2, \dots, k_f \quad (4.2.2)$$

It is required to determine the control strategy $\underline{u}$ so that the scalar function given in (4.2.1) is minimized (or maximized). Frequently, the extremization of the scalar function M is not easy to accomplish. If some simpler function can be found which is closely related to the scalar function and the process dynamics, and if it is easier to perform the optimization with respect to this simpler function, the solution to the optimization problem may then be obtained in a simpler manner. Intuitively, the

scalar function may be minimized by maximizing the energy or the power in the system. This physical intuition leads to the speculation that there may exist an energy function such that its maximization implies the minimization of the scalar function. This function is the Hamiltonian. It is defined as the sum of the kinetic energy and the potential energy and is expressed as the inner product of the momentum vector and the coordinate vector of the system. The simplicity of the Hamiltonian function and its very nature tends to lead one to suspect that maximization of the Hamiltonian function may imply minimization of the scalar function, and that the use of the Hamiltonian may lead to a simple method for solving optimization problems. Pontryagin first discovered this fact for the continuous case and formulated his findings as the celebrated maximum principle.

The maximum (or minimum) principle states that, if the control vector $\underline{u}$ is optimum, that is, if it minimizes (or maximizes) the scalar function M , then the Hamiltonian $H(\underline{x}_k, \underline{\lambda}_k, \underline{u}_k, k)$ is maximized (or minimized) with respect to $\underline{u}_k$ over the control interval. This statement indicates that maximum H implies minimum M and minimum H implies maximum M .

To optimize a process, then, the Hamiltonian must be optimized with respect to the control vector $\underline{u}$. This results in the derivation of difference (or differential)

equations which are solved as a boundary value problem.

By way of summary*, if a system is given as described by equation (4.2.2)

$$\underline{x}_{k+1} = \underline{f}(\underline{x}_k, \underline{u}_k, k)$$

and if for this process a cost functional can be written in the form

$$J = [\theta_k(\underline{x}_k, k)]_{k=k_0}^{k=k_f} + \sum_{k=k_0}^{k_f-1} \phi(\underline{x}_k, \underline{u}_k, k) \quad (4.2.3)$$

then the optimization of the trajectory $\underline{x}_k$ can be achieved by defining a Hamiltonian $H(\underline{x}_k, \lambda_k, \underline{u}_k, k)$, using it to derive difference equations and solving them as a boundary value problem. The definition of the Hamiltonian is given as

$$H_k = \phi(\underline{x}_k, \underline{u}_k, k) + \lambda_{k+1}^T \underline{f}(\underline{x}_k, \underline{u}_k, k) \quad k=k_0, \dots, k_f \quad (4.2.4)$$

Taking partial derivatives of the Hamiltonian with respect to $\lambda_k, \underline{u}_k$ and $\underline{x}_k$ gives the following difference conditions:

$$\underline{x}_{k+1} = \frac{\partial H_k}{\partial \lambda_{k+1}} \quad \text{or} \quad \underline{x}_{k+1} = \underline{f}(\underline{x}_k, \underline{u}_k, k), \quad (4.2.4)$$

$$\frac{\partial H_k}{\partial \underline{u}_k} = 0 \quad \text{or} \quad \frac{\partial \phi_k}{\partial \underline{u}_k} + \left(\frac{\partial \underline{f}^T}{\partial \underline{u}_k} \right) \lambda_{k+1} = 0, \quad (4.2.5)$$

and

$$\lambda_k = \frac{\partial H_k}{\partial \underline{x}_k} \quad \text{or} \quad \lambda_{k+1} = \left(\frac{\partial \underline{f}^T}{\partial \underline{x}_k} \right)^{-1} \left[\lambda_k - \frac{\partial \phi_k}{\partial \underline{x}_k} \right] \quad (4.2.6)$$

*See Appendix C for proof of Pontryagin Maximum Principle

The boundary conditions, if not specified, can be obtained by using the following transversality equations:

$$\eta_{k_0} \left[\lambda_{k_0} - \frac{\partial \theta_{k_0}}{\partial \underline{x}_{k_0}} \right] = 0 \quad (4.2.7)$$

and

$$\eta_{k_f} \left[\lambda_{k_f} - \frac{\partial \theta_{k_f}}{\partial \underline{x}_{k_f}} \right] = 0 \quad (4.2.8)$$

The application of the above conditions and the solution of the equations will yield an optimization.

4.3 Identification with Optimal Control Theory

The basic problem of optimization as stated in the first section can be reformulated as a problem in optimal control. This enables one to identify the problem with the state formulation of the optimal control problem. Once the identification is made, then the whole block of optimal control theory--Pontryagin's Maximum Principle as discretized by Katz--can be applied. Once applied, the resulting difference equations and transversality conditions can be used to solve the problem.

Essentially, the problem is to identify the stochastic approximation algorithm with the state equations of the discrete maximum formulation. Recall the algorithm from equation(4.1.1)

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{f(r_{n+1}) - \hat{x}_n\}$$

Rewritting this in the following form

$$\hat{x}_{n+1} = (1-\gamma_{n+1})\hat{x}_n + \gamma_{n+1}f(r_{n+1}) \quad n=1,2,\dots \quad (4.3.1)$$

Now identifying this equation with (4.2.2) gives the following correspondence:

$\hat{x}_{n+1}$ corresponds with $\underline{x}_{k+1}$,
 $\hat{x}_n$ corresponds with $\underline{x}_k$,
 γ_n corresponds with $\underline{u}_k$,
 and n corresponds with k .

The choice of cost function for the problem is the basic mean square error.

$$J_n = \frac{1}{n} \sum_{i=1}^n (\hat{x}_i - x)^2 \quad (4.3.2)$$

The problem can now be stated. It is desired to minimize the cost function J_n in equation (4.3.2) with respect to the Γ -sequence such that it satisfies (4.3.1) and subject to the convergence constraint stated in equation (3.3.26)

$$\lim_{n \rightarrow \infty} \prod_{i=1}^{n+1} (1-\gamma_i) \rightarrow 0$$

The first step is to define a Hamiltonian function so that the discrete maximum principle can be applied. Using the basic form of (4.2.4) gives the following equation:

$$H_n = (\hat{x}_n - x)^2 + \lambda_{n+1}^{(1)} [(1-\gamma_{n+1})\hat{x}_n + \gamma_{n+1}f(r_{n+1})] + \lambda_{n+1}^{(2)} \prod_{i=1}^{n+1} (1-\gamma_i) \quad n=1,2,\dots \quad (4.3.3)$$

where $\lambda_{n+1}^{(1)}$ and $\lambda_{n+1}^{(2)}$ are Lagrangian multipliers, referred to as the co-state or adjoint variables of the system.

Now applying condition (4.2.6) to the Hamiltonian (4.3.3) gives the first difference equation

$$\frac{\partial H_n}{\partial x_n} = \lambda_n^{(1)}$$

or

$$\lambda_n^{(1)} = 2(\hat{x}_n - x) + (1 - \gamma_{n+1})\lambda_{n+1}^{(1)} \quad (4.3.4)$$

Next, applying condition (4.2.5) to the Hamiltonian gives the second difference equation

$$\frac{\partial H_n}{\partial \gamma_n} = 0$$

or

$$-\lambda_{n+1}^{(2)} \left\{ \prod_{i=1}^{n-1} (1 - \gamma_i) \right\} (1 - \gamma_{n+1}) = 0 \quad (4.3.5)$$

For large values of n the Hamiltonian becomes very close to zero.

$$H_n = 0$$

or

$$(\hat{x}_n - x)^2 + \lambda_{n+1}^{(1)} [(1 - \gamma_{n+1})\hat{x}_n + \gamma_{n+1}f(x_{n+1})] + \lambda_{n+1}^{(2)} \prod_{i=1}^{n+1} (1 - \gamma_i) = 0 \quad (4.3.6)$$

From equation (4.3.5) substitute into the Hamiltonian (4.3.6).

This gives the equation

$$(\hat{x}_n - x)^{2+\lambda}_{n+1} \left[(1-\gamma_{n+1}) \hat{x}_n + \gamma_{n+1} f(r_{n+1}) \right] = 0 \quad (4.3.7)$$

Now equations (4.3.4) and (4.3.7) are the two difference equations to be solved for the optimum r -sequence and the optimum estimation trajectory. But first there is a need for boundary conditions and some auxiliary conditions to eliminate x from these equations.

Now first, recall the form of the algorithm given in equation (4.1.1) given here in a modified form

$$\hat{x}_n = \hat{x}_{n-1} + \gamma_n \{f(r_n) - \hat{x}_{n-1}\} \quad n=1,2,\dots \quad (4.3.8)$$

Iterating this form to $\hat{x}_0$ gives

$$\begin{aligned} \hat{x}_n &= (1-\gamma_n) \hat{x}_{n-1} + \gamma_n f(r_n) \\ &= (1-\gamma_n) (1-\gamma_{n-1}) \hat{x}_{n-2} + (1-\gamma_n) \gamma_{n-1} f(r_{n-1}) + \gamma_n f(r_n) \\ &= (1-\gamma_n) (1-\gamma_{n-1}) (1-\gamma_{n-2}) \hat{x}_{n-3} + (1-\gamma_n) (1-\gamma_{n-1}) \gamma_{n-2} f(r_{n-2}) \\ &\quad + (1-\gamma_n) \gamma_{n-1} f(r_{n-1}) + \gamma_n f(r_n) \end{aligned}$$

Iterating n times gives,

$$\begin{aligned} \hat{x}_n &= (1-\gamma_n) (1-\gamma_{n-1}) (1-\gamma_{n-2}) \dots (1-\gamma_1) \hat{x}_0 \\ &\quad + (1-\gamma_n) (1-\gamma_{n-1}) \dots (1-\gamma_2) \gamma_1 f(r_1) \\ &\quad + (1-\gamma_n) (1-\gamma_{n-1}) \dots (1-\gamma_3) \gamma_2 f(r_2) + \dots \\ &\quad + (1-\gamma_n) (1-\gamma_{n-1}) \gamma_{n-2} f(r_{n-2}) \\ &\quad + (1-\gamma_n) \gamma_{n-1} f(r_{n-1}) \\ &\quad + \gamma_n f(r_n) \end{aligned} \quad (4.3.9)$$

Writing this equation in closed form gives,

$$\hat{x}_n = \hat{x}_0 \prod_{i=1}^n (1-\gamma_i) + \sum_{i=1}^n \gamma_i f(r_i) \prod_{j=i+1}^n (1-\gamma_j) \quad (4.3.10)$$

$n=1,2,\dots$

where the void product is taken as one.

Also recalling the iterated form of the mean error from (3.3.17)

$$(\hat{x}_{n+1}-x) = (\hat{x}_0-x) \left\{ \prod_{i=1}^{n+1} (1-\gamma_i) \right\} + \sum_{i=1}^{n+1} \gamma_i \psi_i \prod_{j=i+1}^{n+1} (1-\gamma_j) \quad n=1,2,\dots$$

and substituting from (3.3.17), and (4.3.10) for $\hat{x}_n$ and $(\hat{x}_{n+1}-x)$ in (4.3.4) and (4.3.7) gives the difference equations to be solved simultaneously for optimal Γ -sequence and optimal estimation trajectory. Equation (4.3.4) yields upon the substitution

$$(1-\gamma_{n+1}) \lambda_{n+1} = \lambda_n - 2 \left\{ (\hat{x}_0-x) \prod_{i=1}^n (1-\gamma_i) + \sum_{i=1}^n \gamma_i \psi_i \prod_{j=i+1}^n (1-\gamma_j) \right\} \quad (4.3.11)$$

Equation (4.3.7) yields upon similar substitution

$$\begin{aligned} & (\hat{x}_0-x)^2 \left\{ \prod_{i=1}^n (1-\gamma_i) \right\}^2 + \left[\sum_{i=1}^n \gamma_i \psi_i \prod_{j=i+1}^n (1-\gamma_j) \right]^2 \\ & + 2 (\hat{x}_0-x) \prod_{i=1}^n (1-\gamma_i) \left\{ \sum_{i=1}^n \gamma_i \psi_i \prod_{j=i+1}^n (1-\gamma_j) \right\} \\ & = \lambda_{n+1}^{(1)} \left\{ \hat{x}_0 \prod_{i=1}^n (1-\gamma_i) + \sum_{i=1}^n \gamma_i f(r_i) \prod_{j=i+1}^n (1-\gamma_j) \right\} \end{aligned} \quad (4.3.12)$$

In addition to equations (4.3.11) and (4.3.12), transversality equations are needed to solve these difference equations. Applying the transversality condition of equations (4.2.7) and (4.2.8) to the cost function in equation (4.3.2) gives

$$\lambda_{n_f}^{(1)} = 0 \quad (4.3.13)$$

$$\lambda_{n_0}^{(1)} = 0 \quad (4.3.14)$$

Now using the transversality conditions (4.3.13) and (4.3.14) with equations (4.3.11) and (4.3.12), the optimization of the rate of convergence of the algorithm (4.1.1)

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{f(x_{n+1}) - \hat{x}_n\} \quad n=1,2,\dots$$

having selected the function f as in (4.1.2), with respect to an optimum Γ -sequence, can be obtained by solving the boundary value problem formulated in this section.

4.4 Boundary Value Problem

It is evident from the non-linear nature of equations (4.3.11) and (4.3.12) that the boundary value problem formulated above is a discrete two point boundary value problem which in the general case cited here cannot be solved. Even if the mechanics of the mathematics would be tractable, the resulting solutions for the optimum Γ -sequence and estimation trajectory are both dependent on the initial error

$$E_0 = \hat{x}_0 - x \quad (4.4.1)$$

considered as a constant in the boundary value problem. Since the optimality depends on the value of this constant it is not to one's advantage to use this approach to the problem. It is suggested that since the very essence of stochastic approximation is the basic simplicity of computations involved in the approximation, it would grossly hamper the efficiency of the algorithm if the boundary value problem had to be solved for a value of gamma for every step and between every iterated step. From the present state of knowledge, it would appear to be wise to accept the penalty of having to take a few extra iterations of the algorithm rather than incurring the cost of trying to obtain a closed form optimal Γ -sequence. It would appear that only when applications would require such a rapid convergence, based on a time limit allowed for a parameter approximation or a cost saving based on repeated application of this algorithm, that the effort be expended to solve the boundary value problem once and store the result for repeated usage. An alternative would be to solve a reduced problem of the same type for a closed form solution and then make use of this suboptimal Γ -sequence. This will probably have to be left to the engineering judgment of the individuals making the particular application.

4.5 Uniqueness of the Optimum Gamma Sequence

For the stochastic approximation algorithm given in equation (4.1.1) to converge optimally to solution requires a Γ -sequence dependent on the starting point and on the function in (4.1.2). Hence, it can be said that no single Γ -sequence can give an optimal convergence from different starting points, and as such a single unique Γ -sequence does not exist which can make the algorithm converge optimally every time. The best that can be expected is a suboptimal convergence that results in solutions somewhat more quickly than previous algorithms. This, in fact, will be shown in the simulation results.

CHAPTER V

Simulation Results

5.1 Structure of the Simulations

It has been proven that the stochastic approximation algorithm making use of sample autocorrelation as the sample information converges to the true value sought. This is achieved even in the presence of a random contaminating environment which interferes with the sampling. The numerical simulations and comparison of results have been designed to test this basic property, that is, the ability of the algorithm to converge to the sought value. In addition, it is desirable to investigate the rate of convergence of the algorithm. This latter section of the investigation develops quite easily into a manifold investigation. Since it has been shown that there does not exist a unique universal optimal r -sequence, the relative merits of these sequences as compared to the state of the art algorithms and their corresponding sequences is left to be determined and evaluated through experiment and/or simulation.

The basic concept behind the simulation has been the parameter identification made on samples which are a linear combination, that is, the sum of the parameter and noise contaminant. The samples were constructed thus:

$$r_n = x + \xi_n \quad (5.1.1)$$

where as before

r_n is the n^{th} sample

x is the true value sought

and ξ_n is an element of zero mean normally distributed noise*

From the autocorrelation function and power spectrum of the random process, it can be seen that

$$R_{r_n}(\ell) = \sigma^2 \delta(\ell) \quad (5.1.2)$$

holds to a good approximation and that the power spectrum is crudely uniform.

Recalling then the two existing algorithms, that is, the first and second algorithms (of Fu) from equations (3.2.11) and (3.2.12)

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{r_{n+1} - \hat{x}_n\} \quad n=1,2,\dots \quad (5.1.3)$$

and

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \left\{ \frac{1}{n+1} \sum_{i=1}^{n+1} r_i - \hat{x}_n \right\} \quad n=1,2,\dots \quad (5.1.4)$$

and comparing these with the new algorithm developed

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{ (R_{r_{n+1}}(\ell))^{\frac{1}{2}} - \hat{x}_n \} \quad n=1,2,\dots \quad (5.1.5)$$

where

*See Appendix F for complete details of the random process.

$$\hat{R}_{n+1}^{(l)} = \frac{1}{n-l} \sum_{i=1}^{n-l} r_i r_{i+l} \quad (5.1.6)$$

for the Γ -sequences as selected in equations (3.4.1) and (3.4.5) with $p=1$. Basically, then, the Γ -sequences are

$$\gamma_n = 1/n \quad n=1,2,\dots \quad (5.1.7)$$

or

$$\gamma_n = 1 - 1/n \quad n=1,2,\dots \quad (5.1.8)$$

Now in order to be able to make some comparison of the relative merits of the three algorithms, some criteria had to be selected. In fact, two measures of error were used to establish not only merit but consistent merit. They were the sample square error (S.S.E.) which is the discrete equivalent of the integral square error, and the time or interval sample square error (T.S.S.E.) which is the discrete equivalent of the time integral square error. More precisely,

$$\text{S.S.E.} = \frac{1}{n} \sum_{i=1}^n (\hat{x}_i - x)^2 \quad (5.1.9)$$

and

$$\text{T.S.S.E.} = \frac{1}{n} \sum_{i=1}^n \frac{i}{k} (\hat{x}_i - x)^2 \quad (5.1.10)$$

where k is simply a scaling constant which is the same through this work.

It was felt that two criteria would be of definite benefit to this study. The S.S.E. is particularly sensitive to large errors particularly those which occur at the

beginning, that is, for small values of n . The T.S.S.E. is sensitive to small errors which may persist after some time, that is, for large values of n . Hence, by using both indicators, a measure of initial or transient deviation as well as residual error after some time can be achieved for absolute or comparative purposes by using S.S.E. and T.S.S.E. respectively.

5.2 Illustration of Simulation Results

The method employed for the simulations was simple. A sequence of r_n 's $n=1,2,\dots, 200$ was generated and stored for a given value of standard deviation, σ , of the random component $\xi_n \in N(0, \sigma^2)$. Each of the three algorithms based its sample information on these r_n at every n and calculated an estimate $\hat{x}_n$. Along with this was computed the S.S.E. and T.S.S.E. at every n for all three algorithms.

In an actual simulation, the first estimate for each algorithm is taken to be the first sample r_n . From this starting point, the approximation trajectory (A.P.), that is, the successive estimates or approximations $\hat{x}_n$, for each algorithm is computed and normalized. This is done for a given Λ where

$$\Lambda \equiv \frac{\sigma}{\bar{x}} \quad (5.2.1)$$

is a measure of the noise content contaminant in the make-up of the sample r_n as compared to the parameter being sought. For example, if the variance, σ^2 , of the noise is 16 and the

parameter, x , being sought is 2, then $\Lambda=2$. This is approximately equal to -6db signal to noise ratio.

Now for a given Λ and ℓ , the approximation trajectory (A.P.), sum square errors (S.S.E.) and the time sum square errors (T.S.S.E.) are plotted as a function of n . Each of these contain three loci, one for each algorithm according to the following key:

Algorithm 1 — — — as in equation (5.1.3)

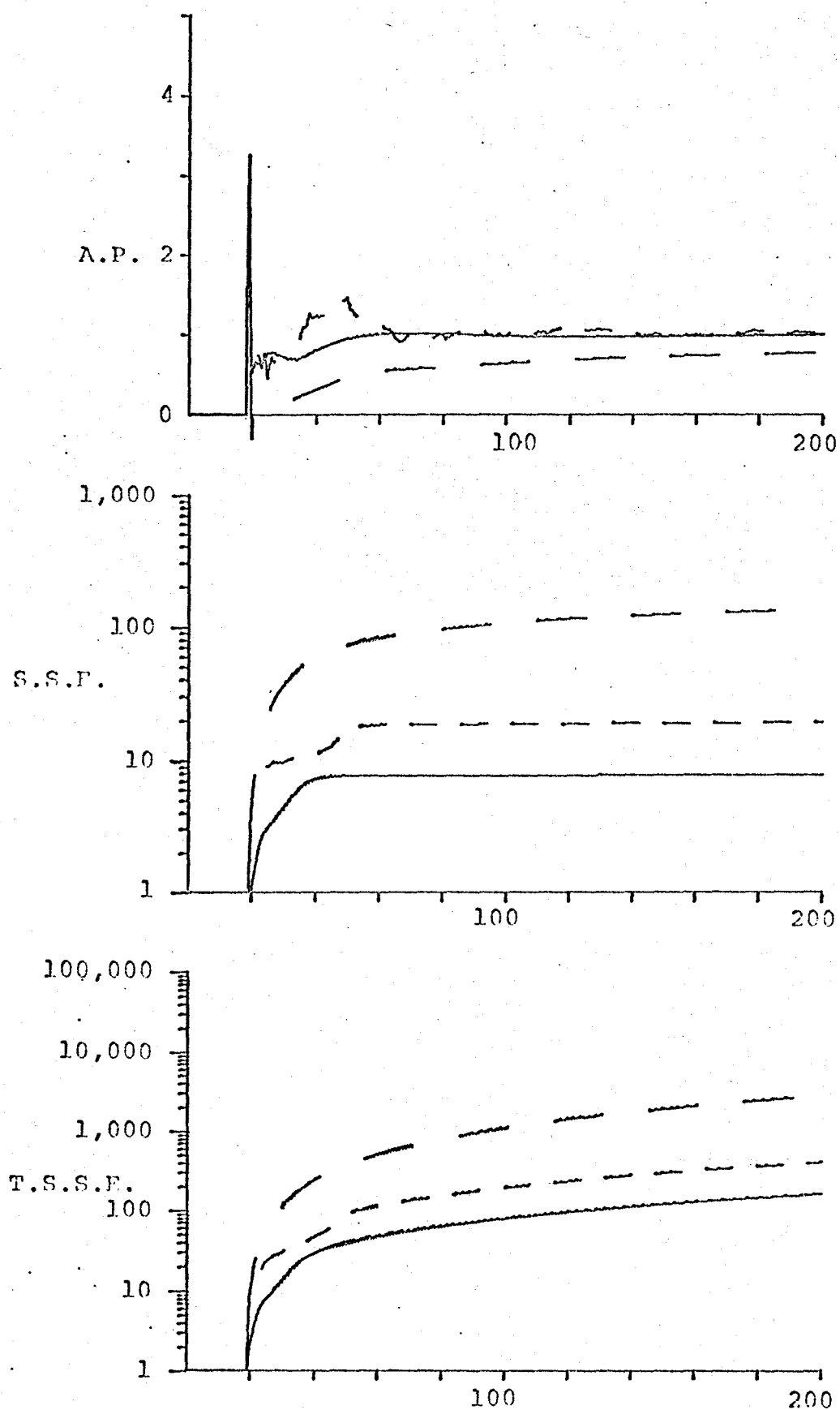
Algorithm 2 — — — as in equation (5.1.4)

Algorithm 3 — — — as in equation (5.1.5)

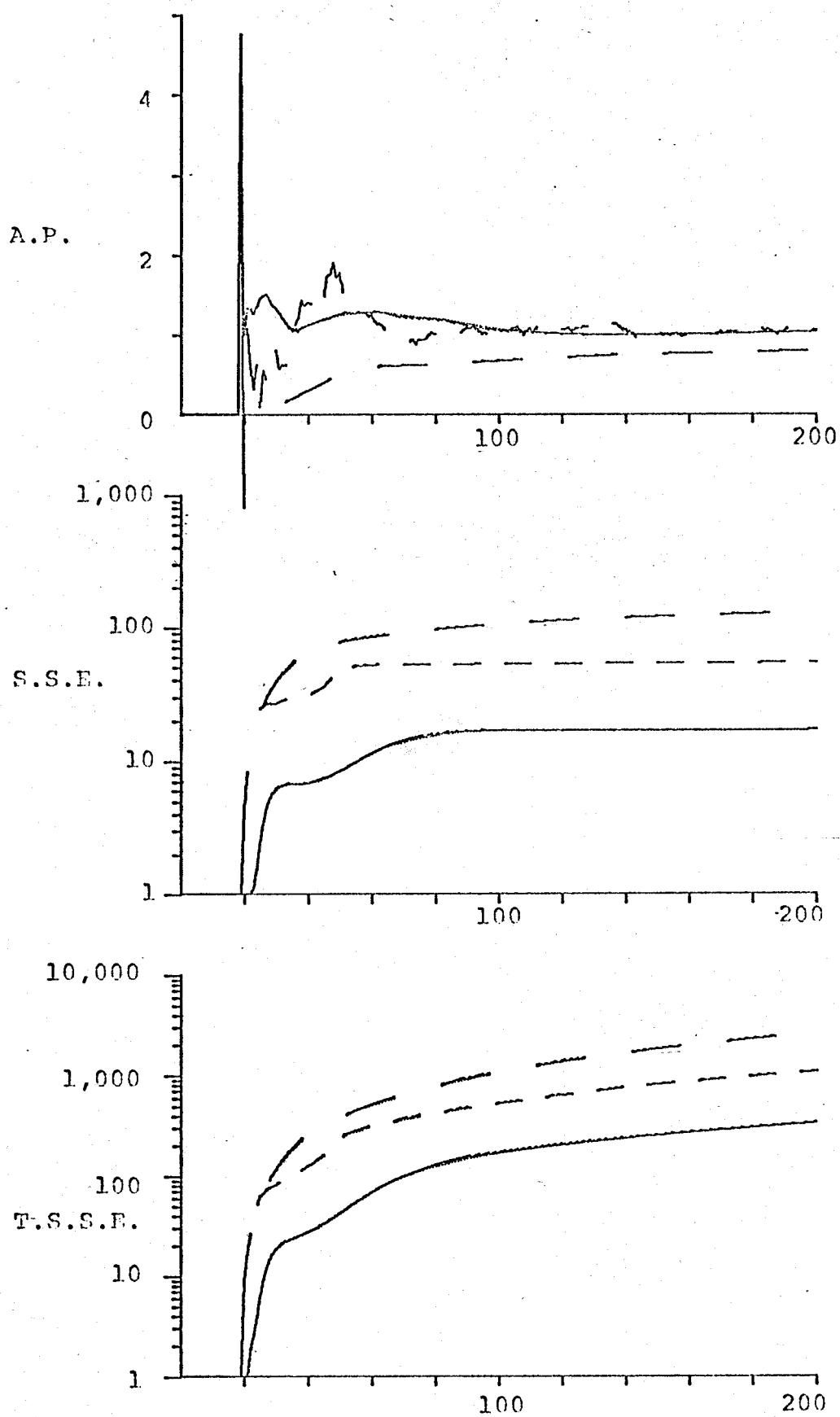
Eight sample simulations will be illustrated first for various values of Λ and for $\ell=20$ and $\gamma_n=1/n$ as shown in Table (5.2.1)

TABLE (5.2.1)

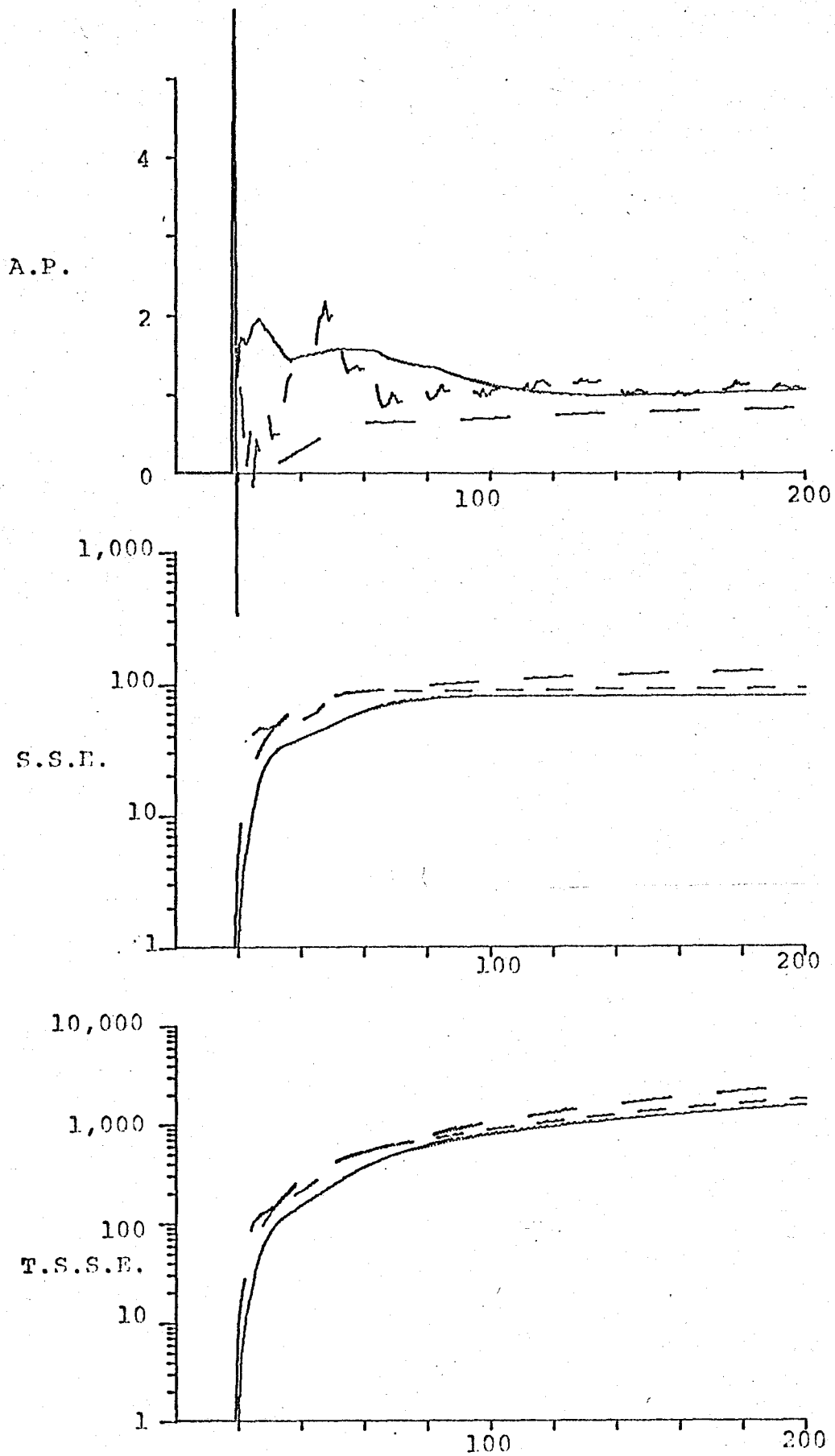
<u>Graph No.</u>	<u>Λ</u>
(5.2.1)	1.50
(5.2.2)	1.75
(5.2.3)	2.00
(5.2.4)	2.25
(5.2.5)	2.50
(5.2.6)	2.75
(5.2.7)	3.00
(5.2.8)	3.25



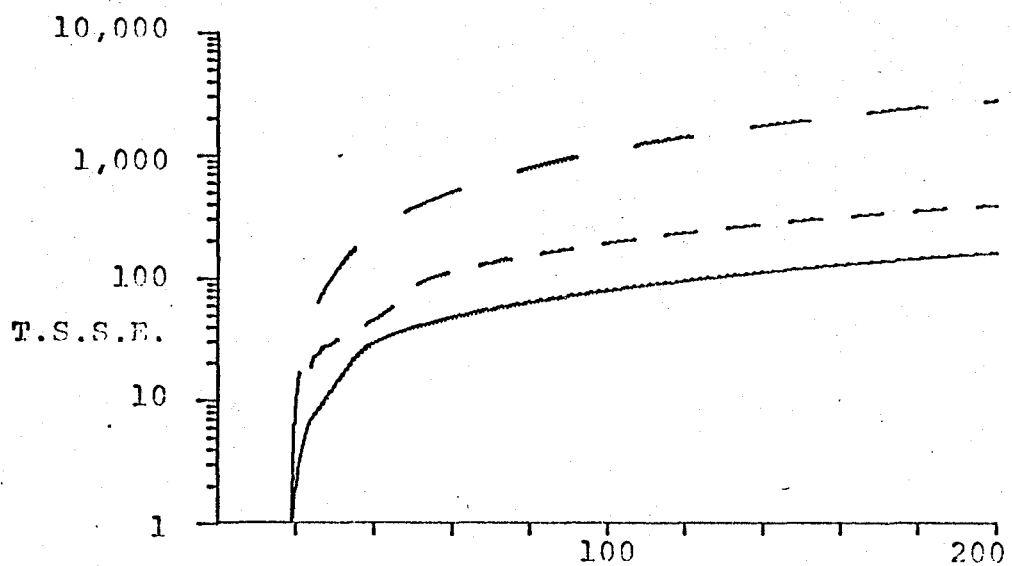
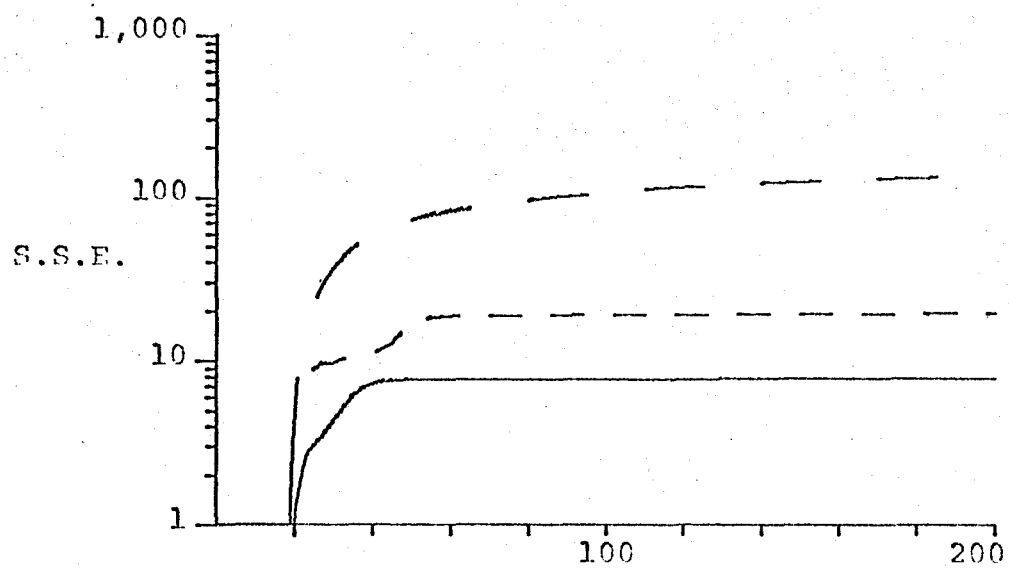
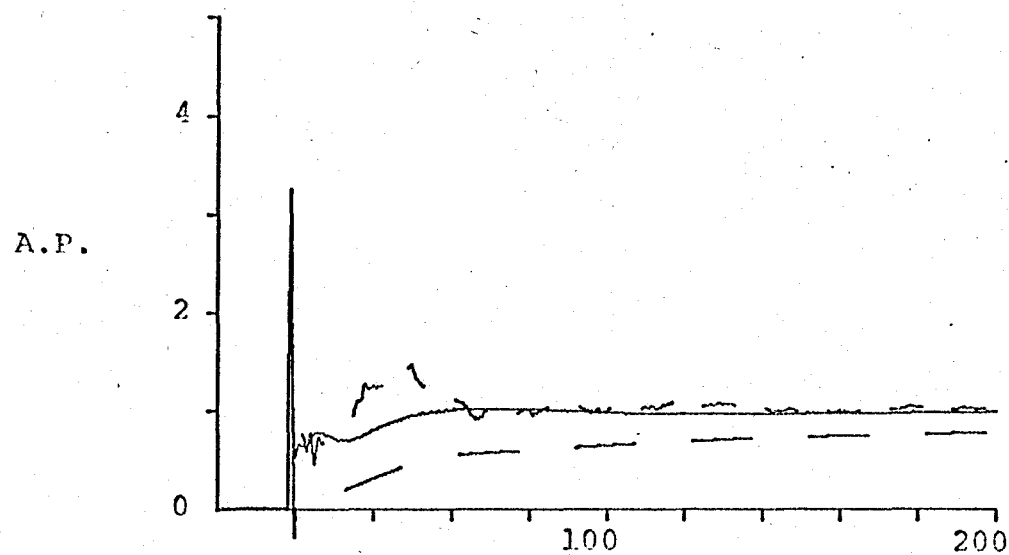
Graph (5.2.1) A.P., S.S.F., and T.S.S.F. as a function of n .
 $\Lambda = 1.50$



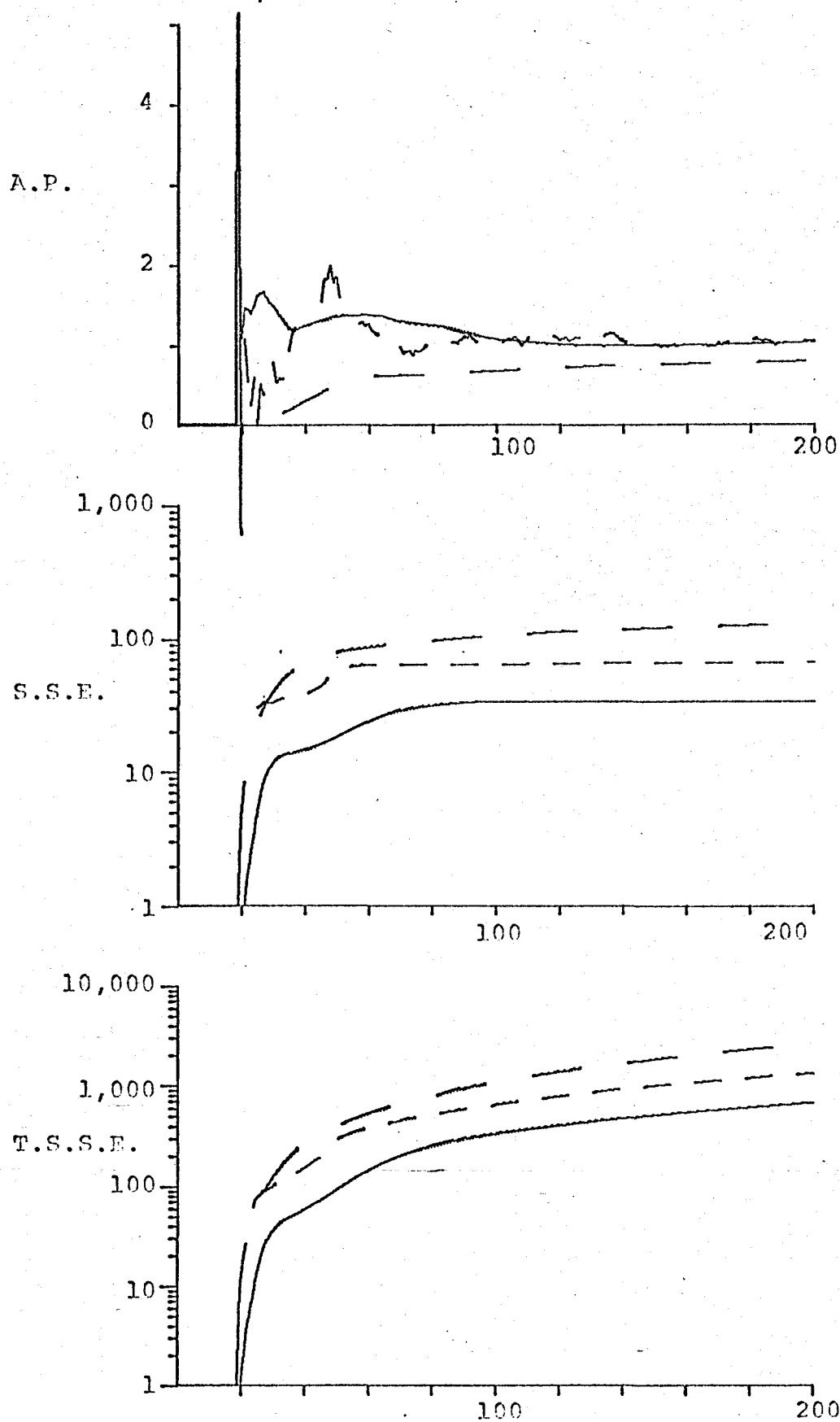
Graph (5.2.2) A.P., S.S.E., and T.S.S.E. as a function of n
 $\Lambda = 1.75$



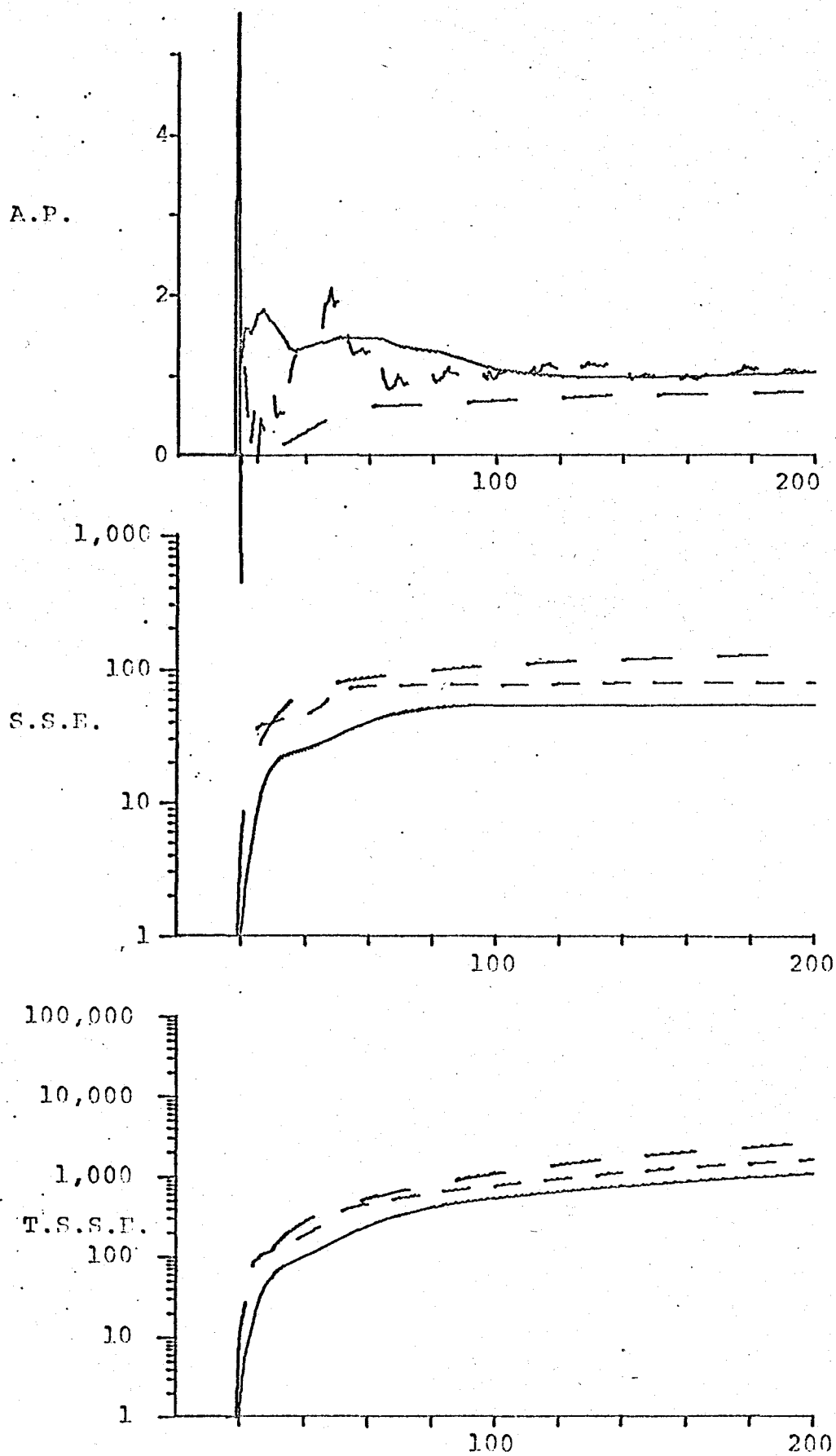
Graph (5.2.3) A.P., S.S.E., and T.S.S.E. as a function of n
 $\Lambda = 2.00$



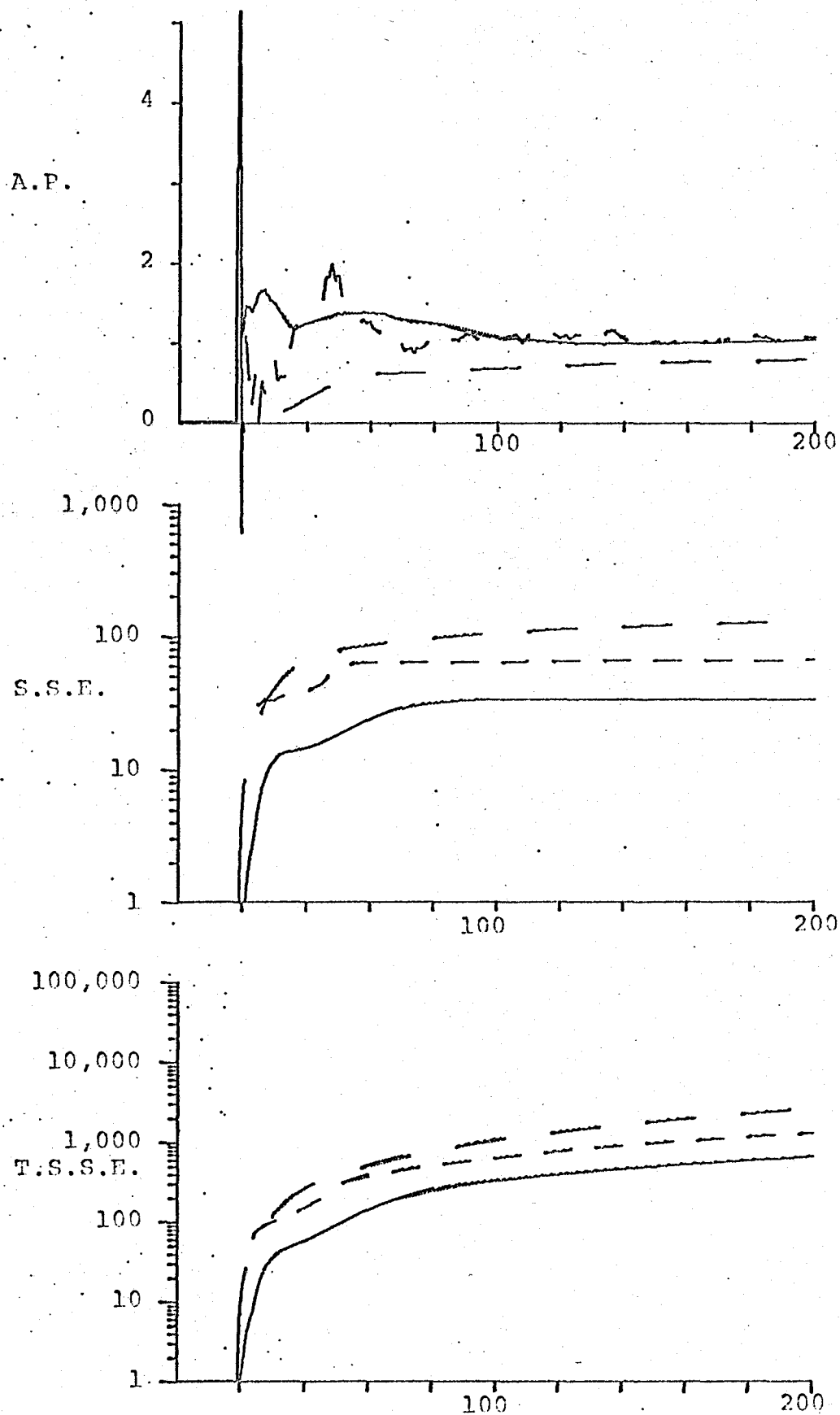
Graph (5.2.4) A.P., S.S.E., and T.S.S.E. as a function of n
 $\Lambda = 2.25$



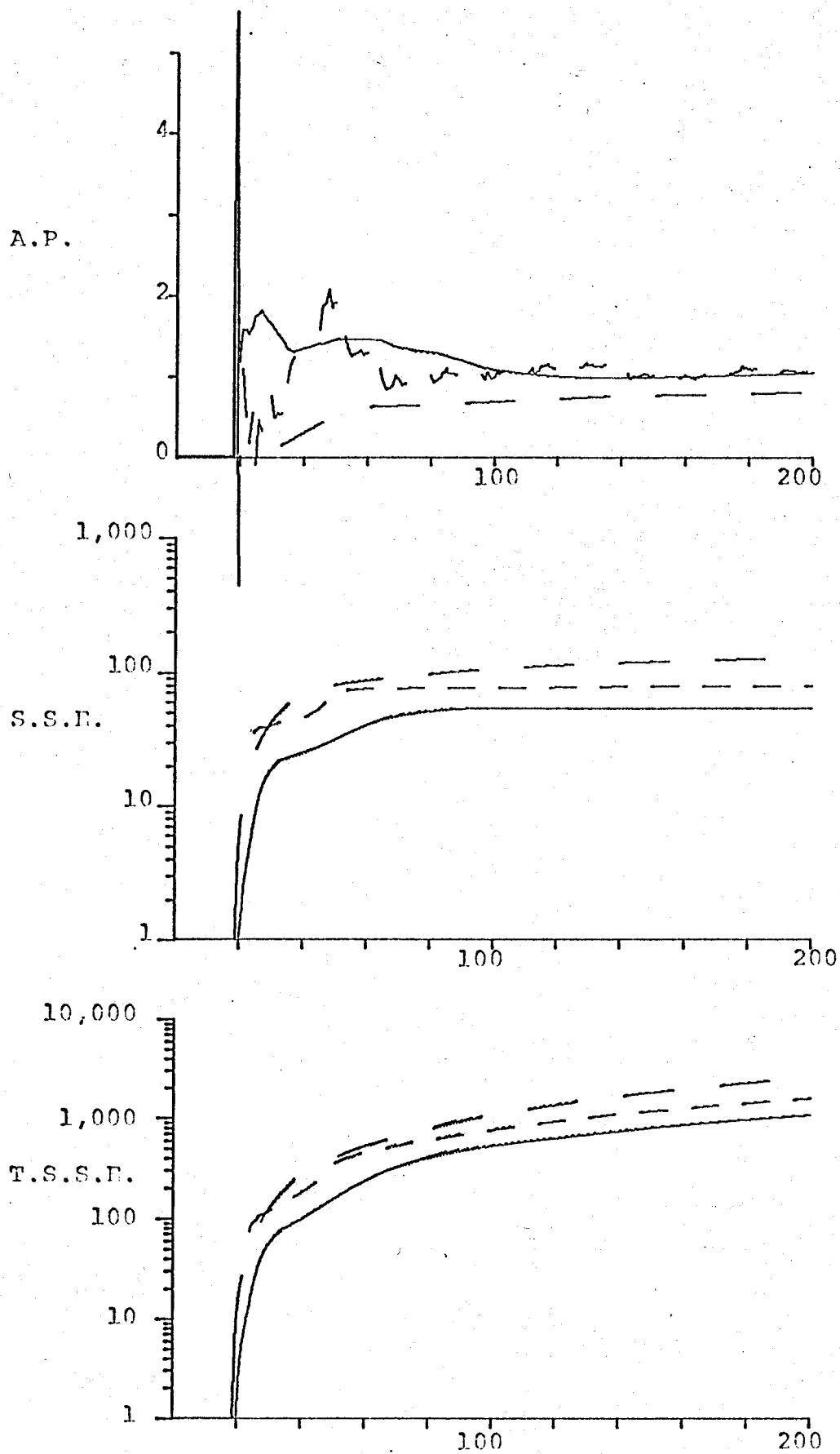
Graph (5.2.5) A.P., S.S.E., AND T.S.S.E. as a function of n
 $\Lambda = 2.50$



Graph (5.2.6) A.P., S.S.E., AND T.S.S.E. as a function of n
 $\Lambda = 2.75$



Graph (5.2.7) A.P., S.S.E., and T.S.S.E. as a function of n .
 $\Lambda = 3.00$

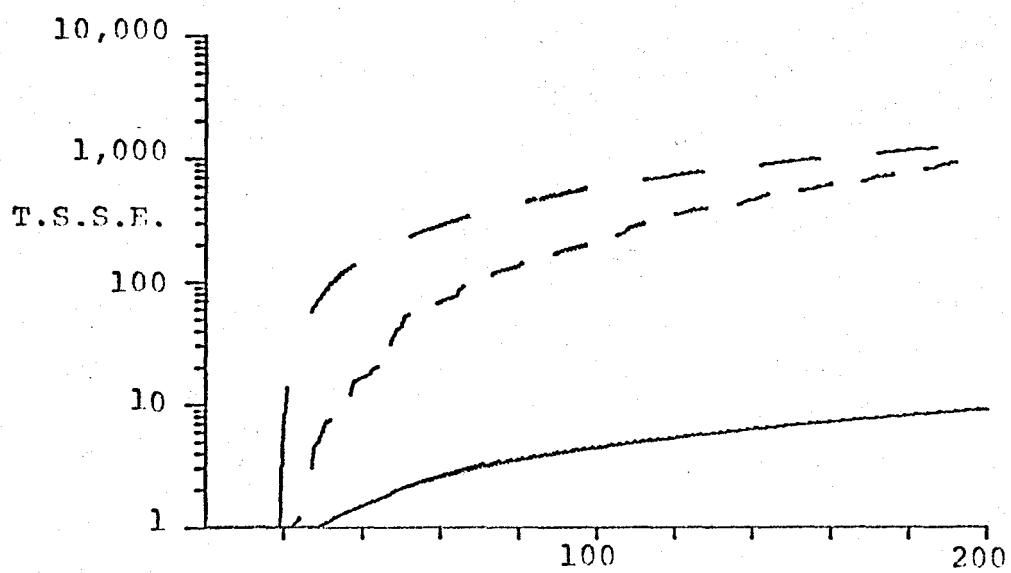
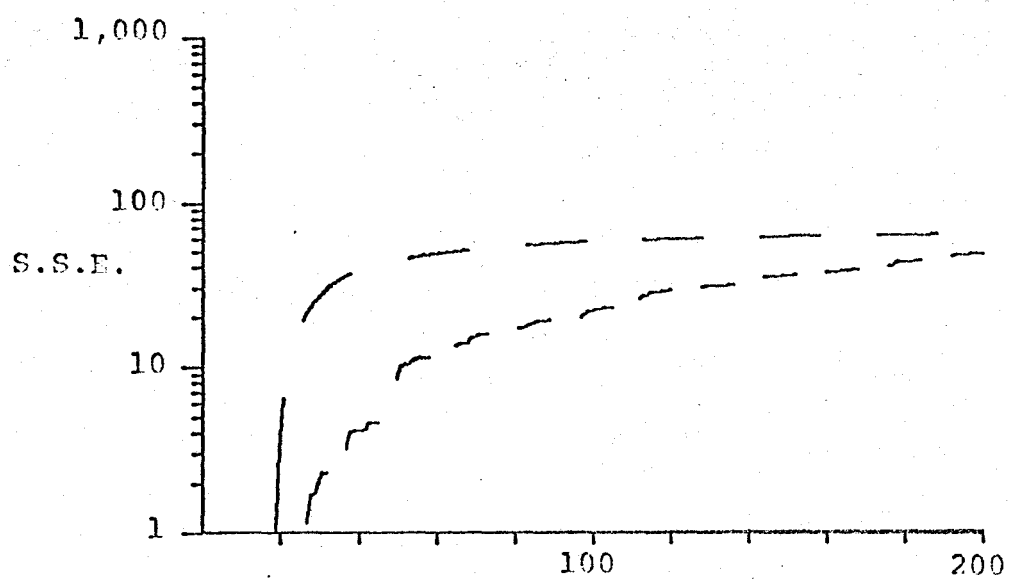
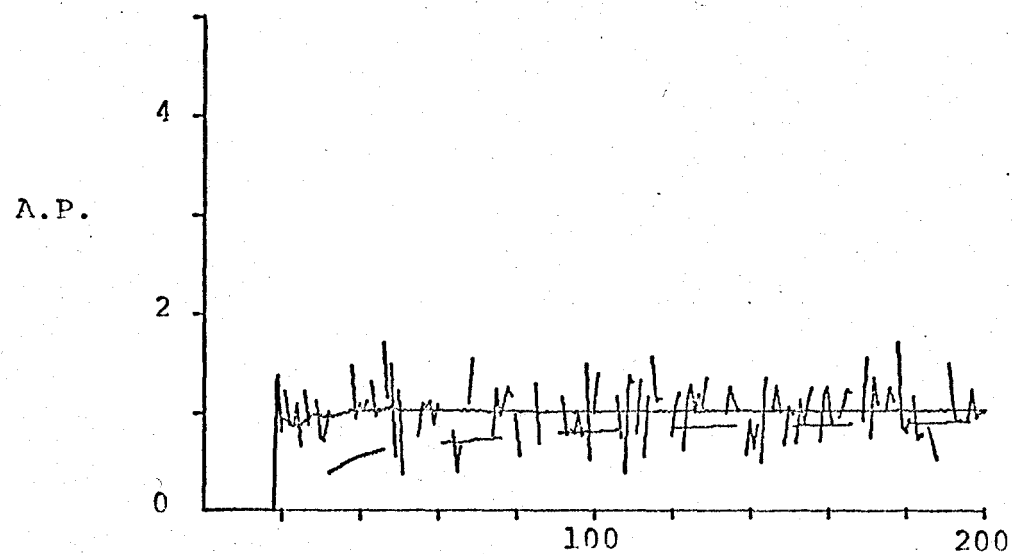


Graph (5.2.8) A.P., S.S.E., and T.S.S.E. as a function of n .
 $\Lambda = 3.75$

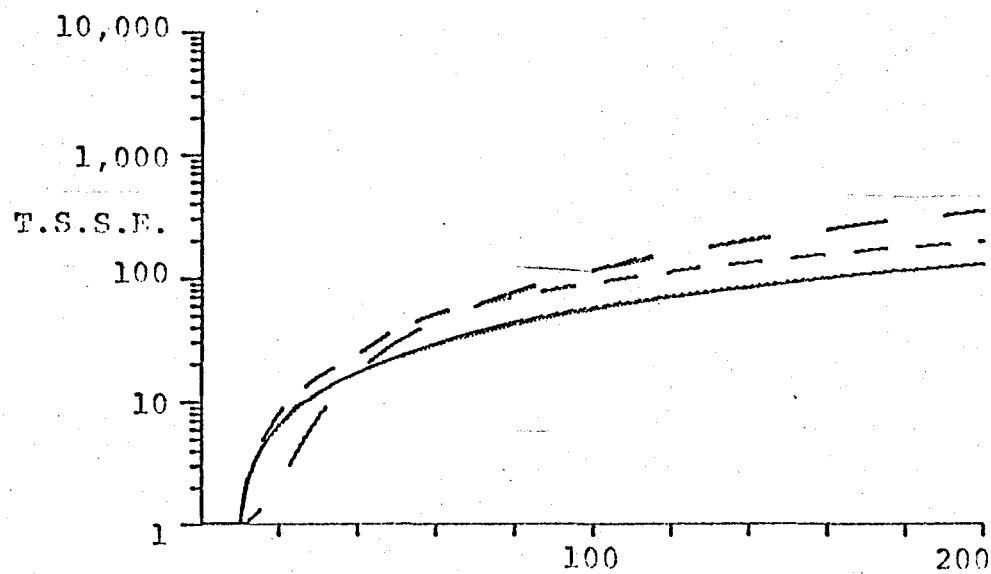
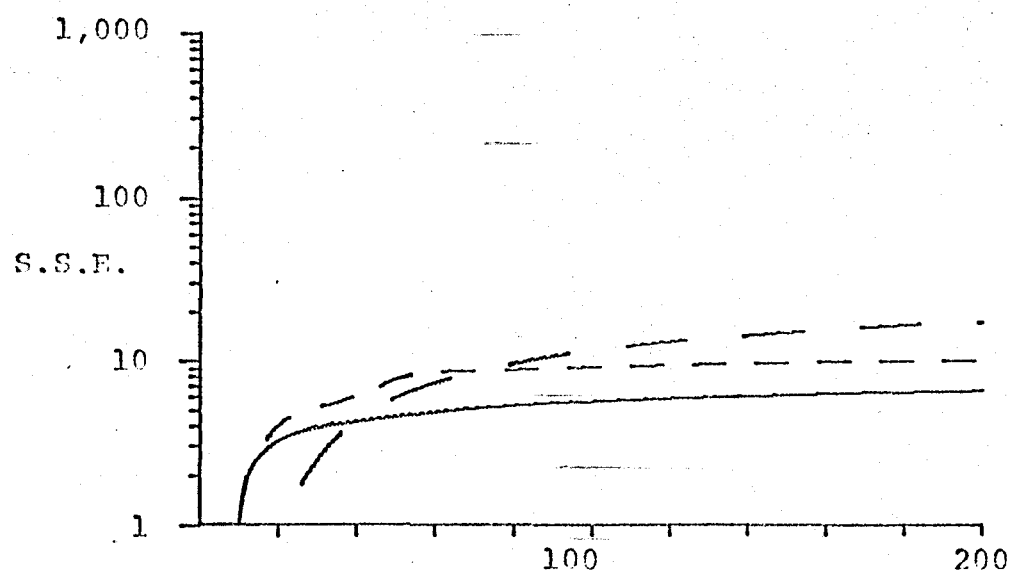
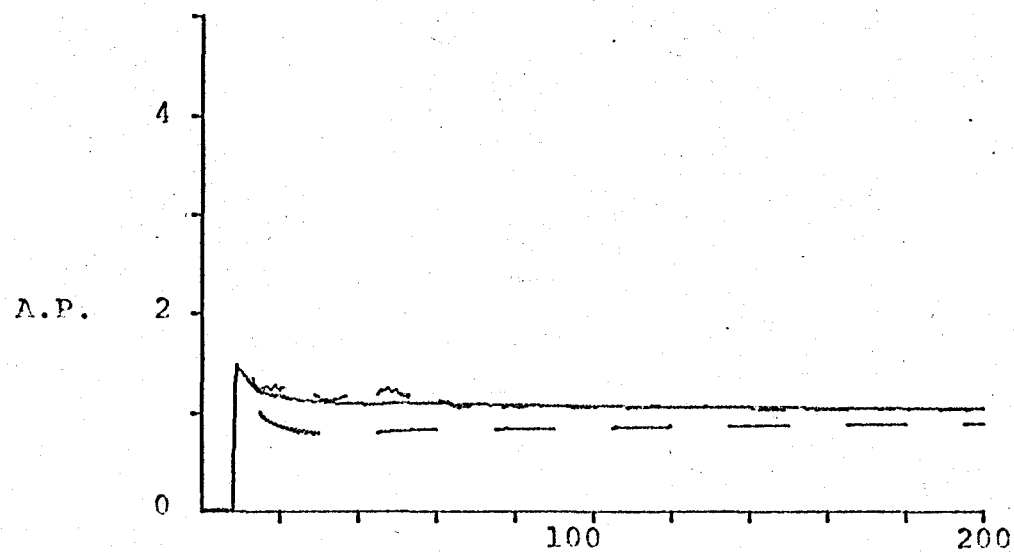
In addition to the simulations that have been shown in Graphs (5.2.1) through to (5.2.8), a number of sample simulations using similar conditions but with the Γ -sequence as in equation (5.1.8). The format remains the same, with A.P., S.S.E., and T.S.S.E. plotted on the same graph for all three algorithms. The actual noise conditions and, in particular, the value of ℓ for the third algorithm is shown with the other information on Table (5.2.2).

Table (5.2.2)

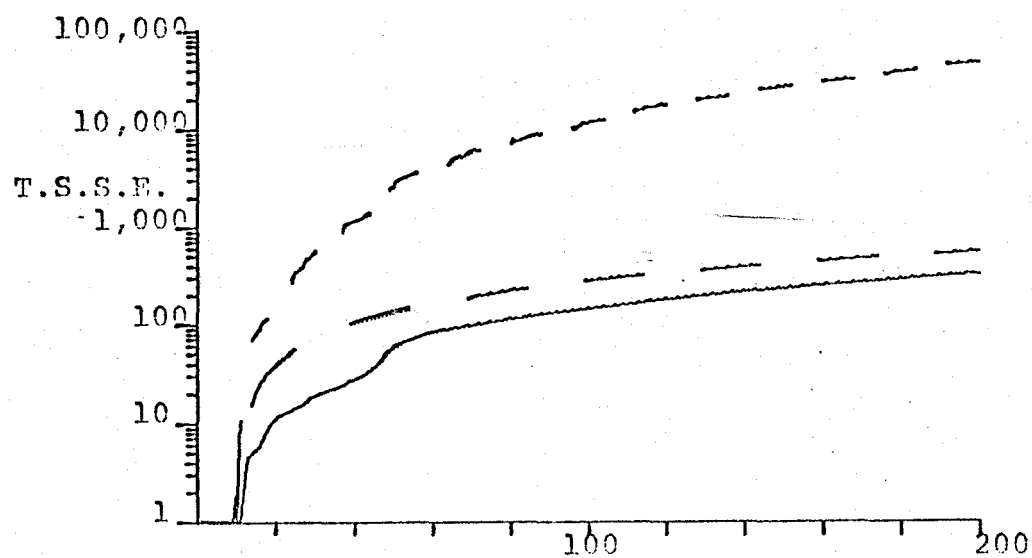
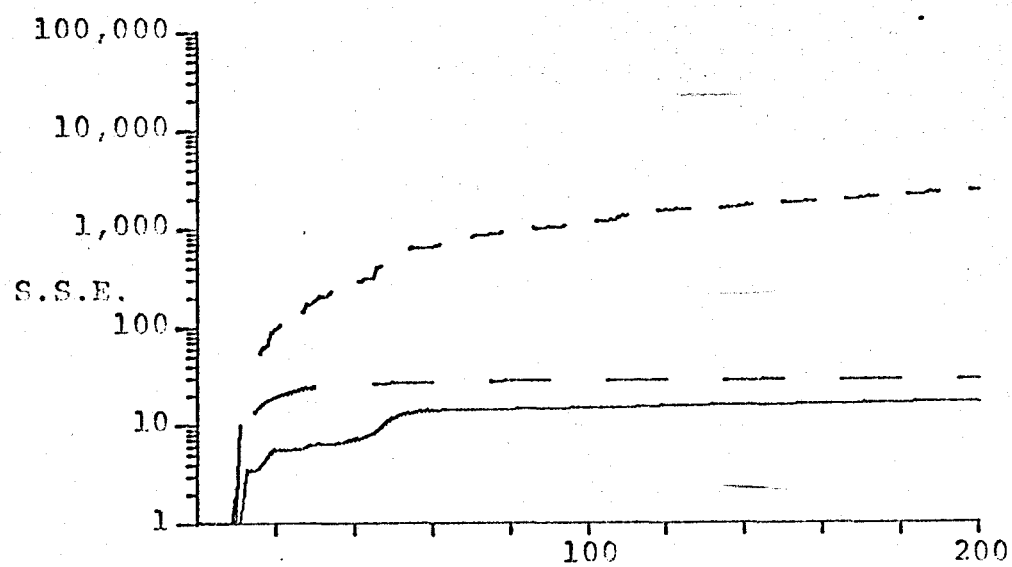
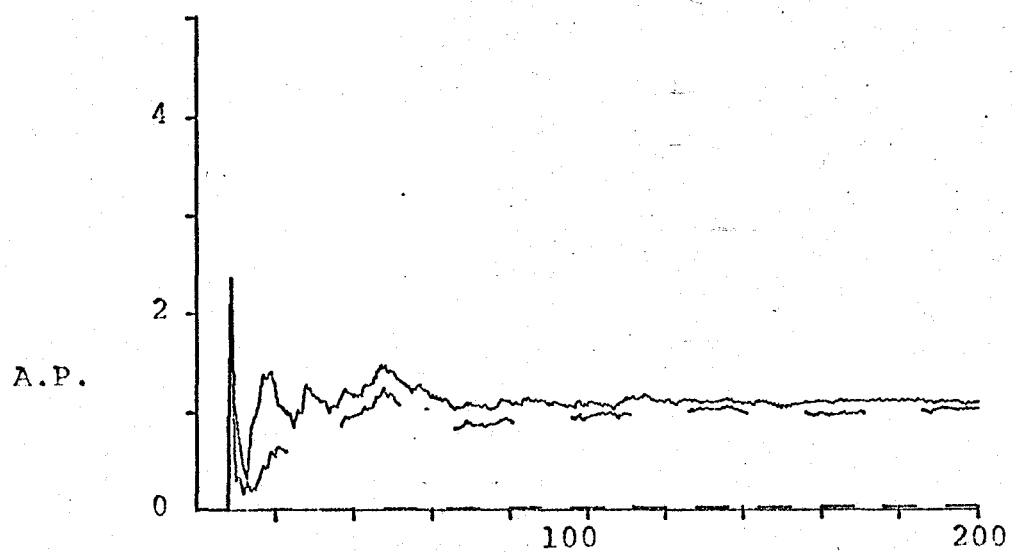
<u>Graph No.</u>	<u>Λ</u>	<u>ℓ</u>
(5.2.9)	0.50	10
(5.2.10)	0.75	10
(5.2.11)	1.75	10
(5.2.12)	2.00	10
(5.2.13)	2.25	10
(5.2.14)	0.25	20
(5.2.15)	0.50	20
(5.2.16)	0.75	20



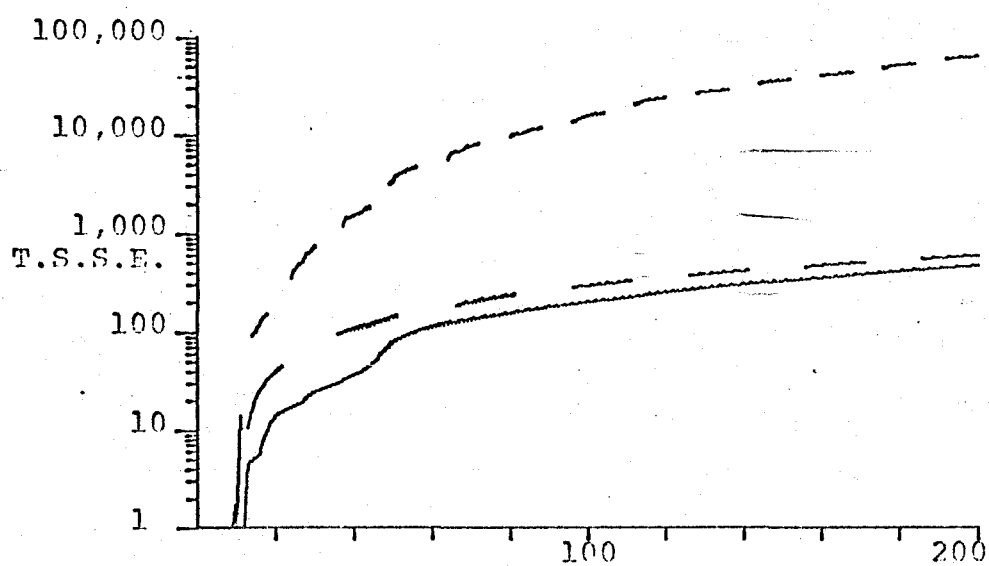
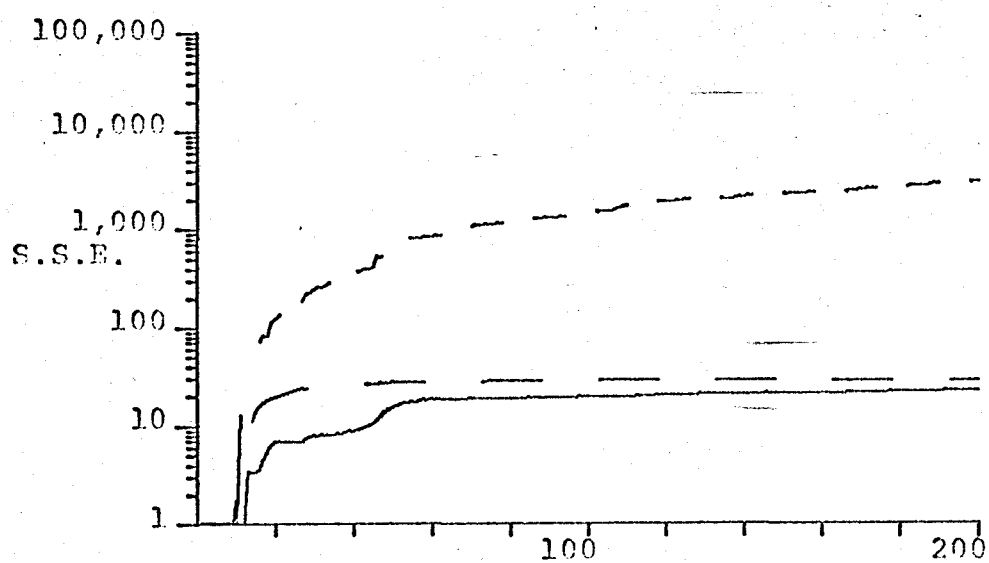
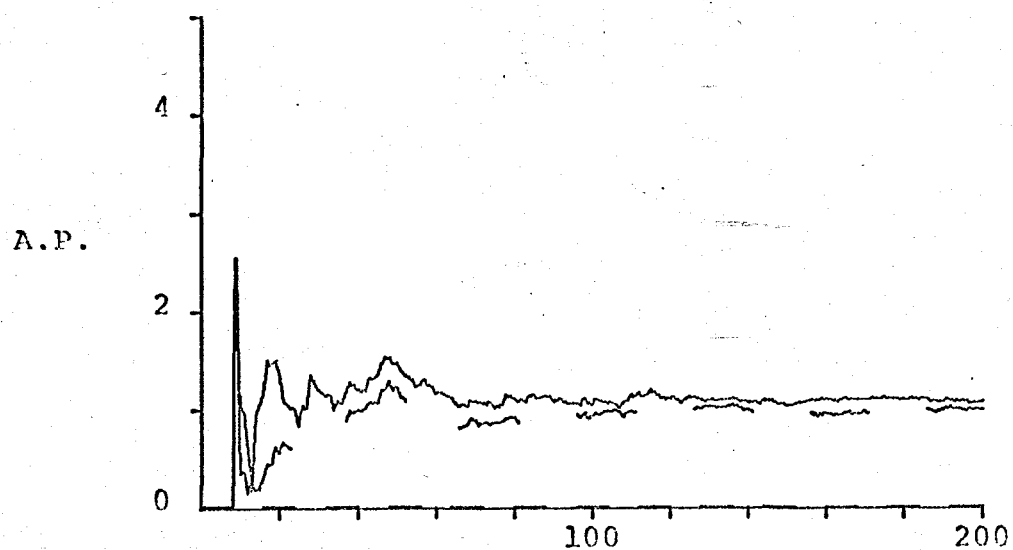
Graph (5.2.9) A.P., S.S.E. and T.S.S.E. as a function of n
 $\Lambda = 0.50$



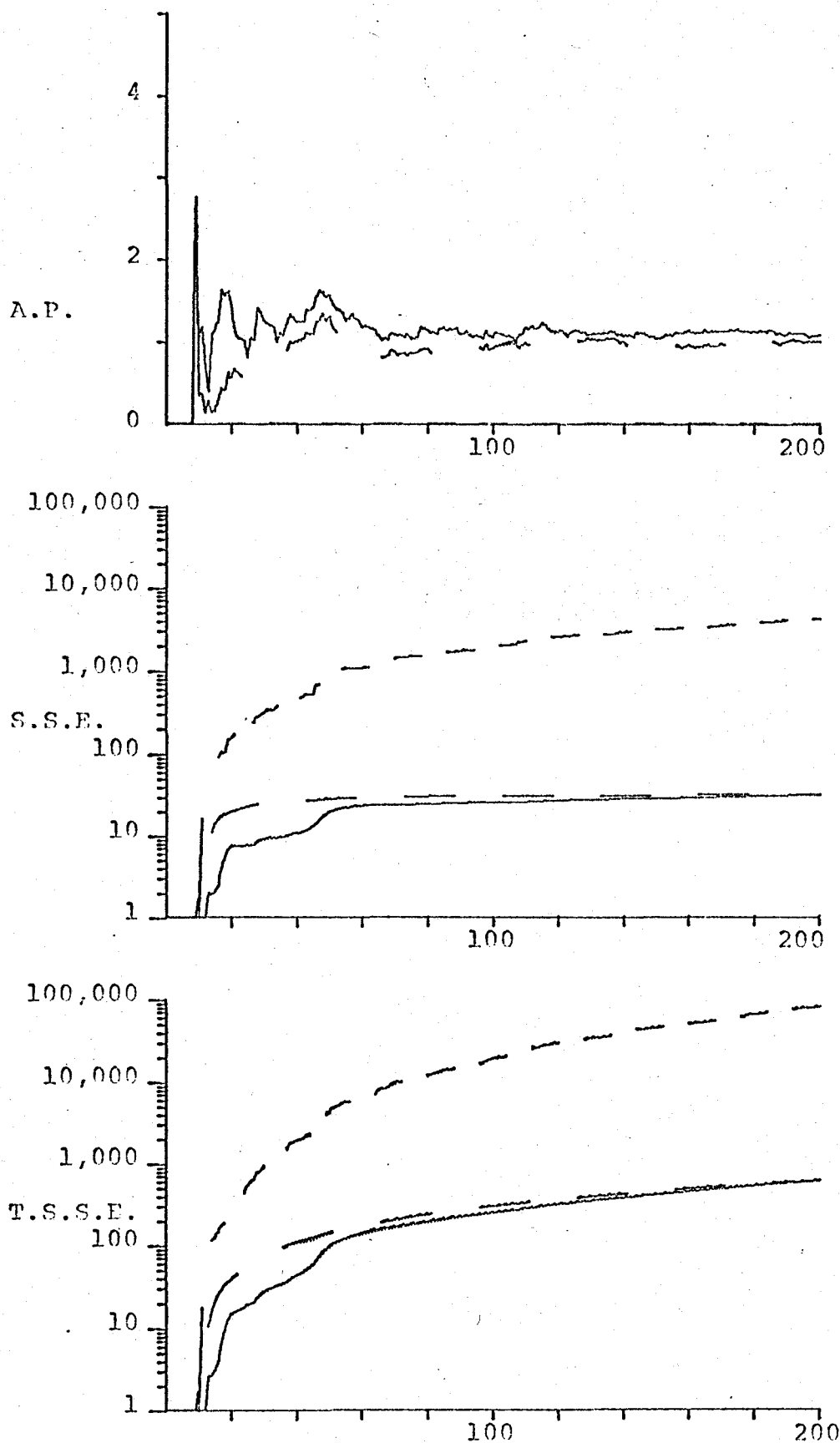
Graph (5.2.10) A.P., S.S.E. and T.S.S.E. as a function of n
 $\Lambda = 0.75$



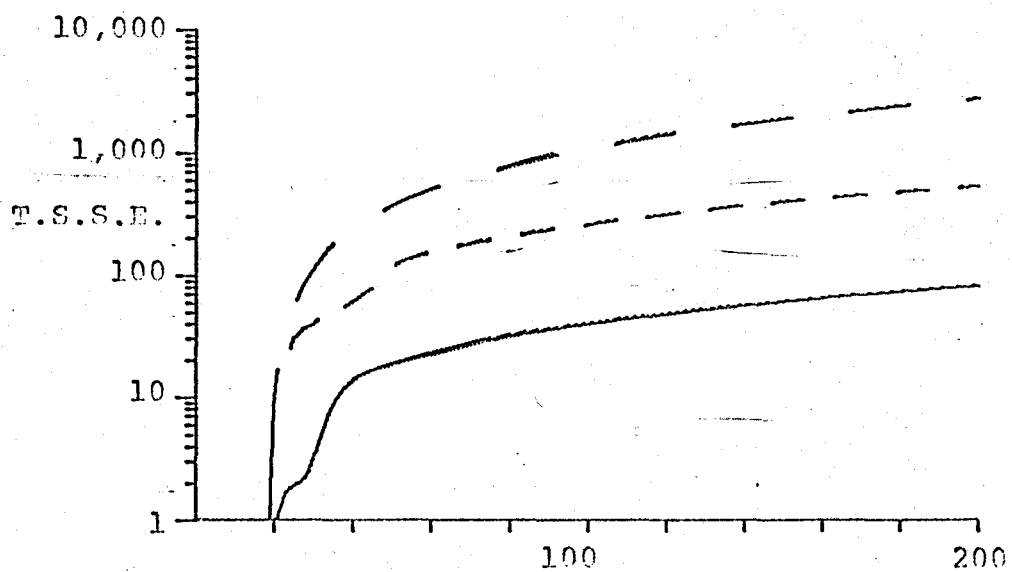
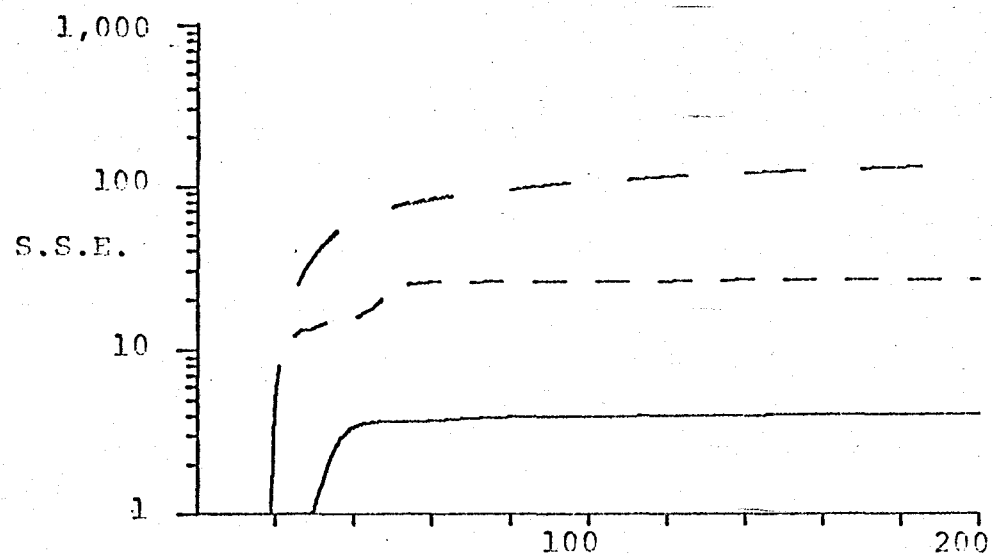
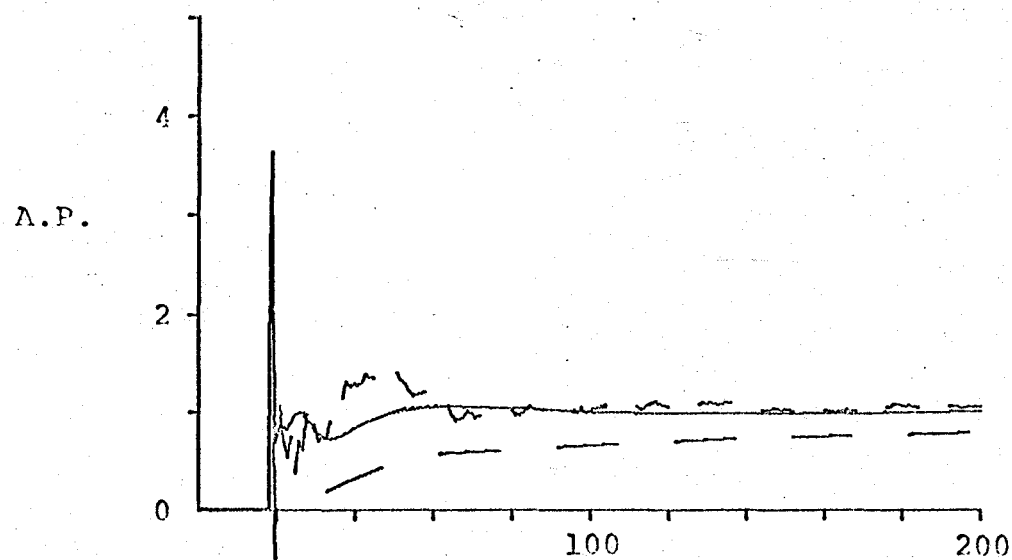
Graph (5.2.11) A.P., S.S.E. and T.S.S.E. as a function of n
 $\Lambda = 1.75$



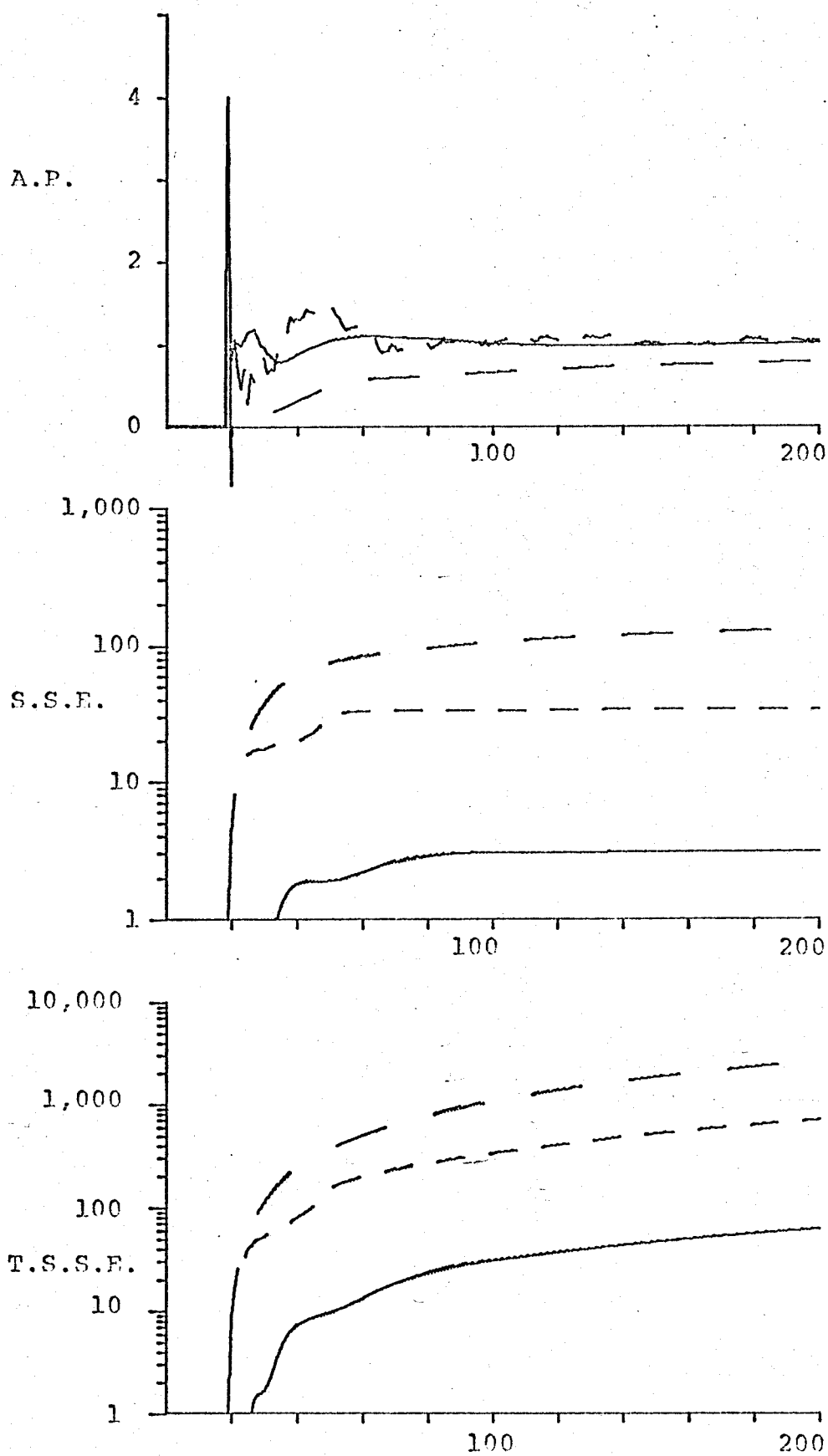
Graph (5.2.12) A.P., S.S.E. and T.S.S.E. as a function of n
 $\Lambda = 2.00$



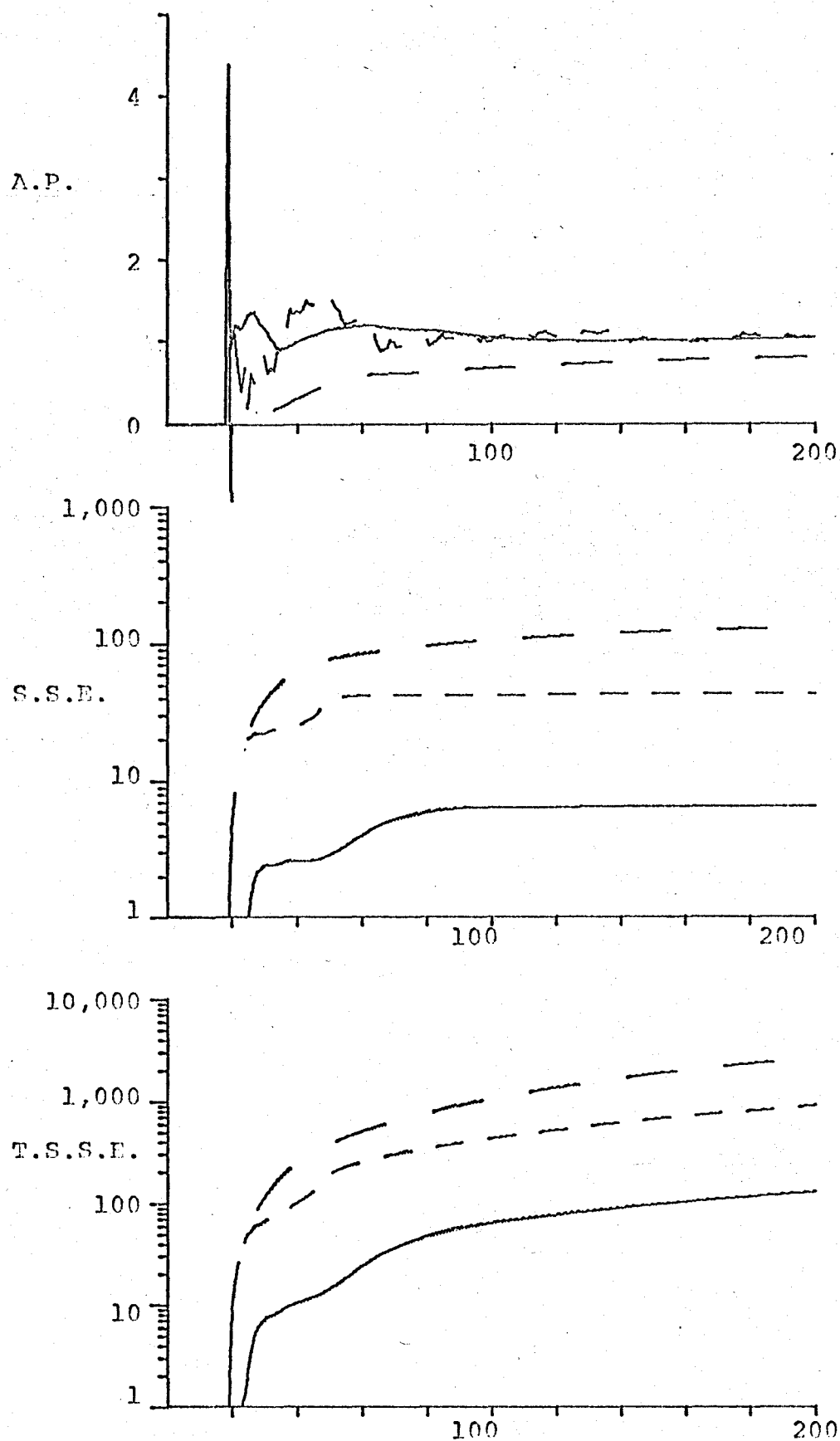
Graph (5.2.13) A.P., S.S.E. and T.S.S.E. as a function of n
 $\Lambda = 2.25$



Graph (5.2.14) A.P., S.S.E. and T.S.S.E. as a function of n
 $\Lambda = 0.25$



Graph (5.2.15) A.P., S.S.E. and T.S.S.E. as a function of n
 $\Lambda = 0.50$

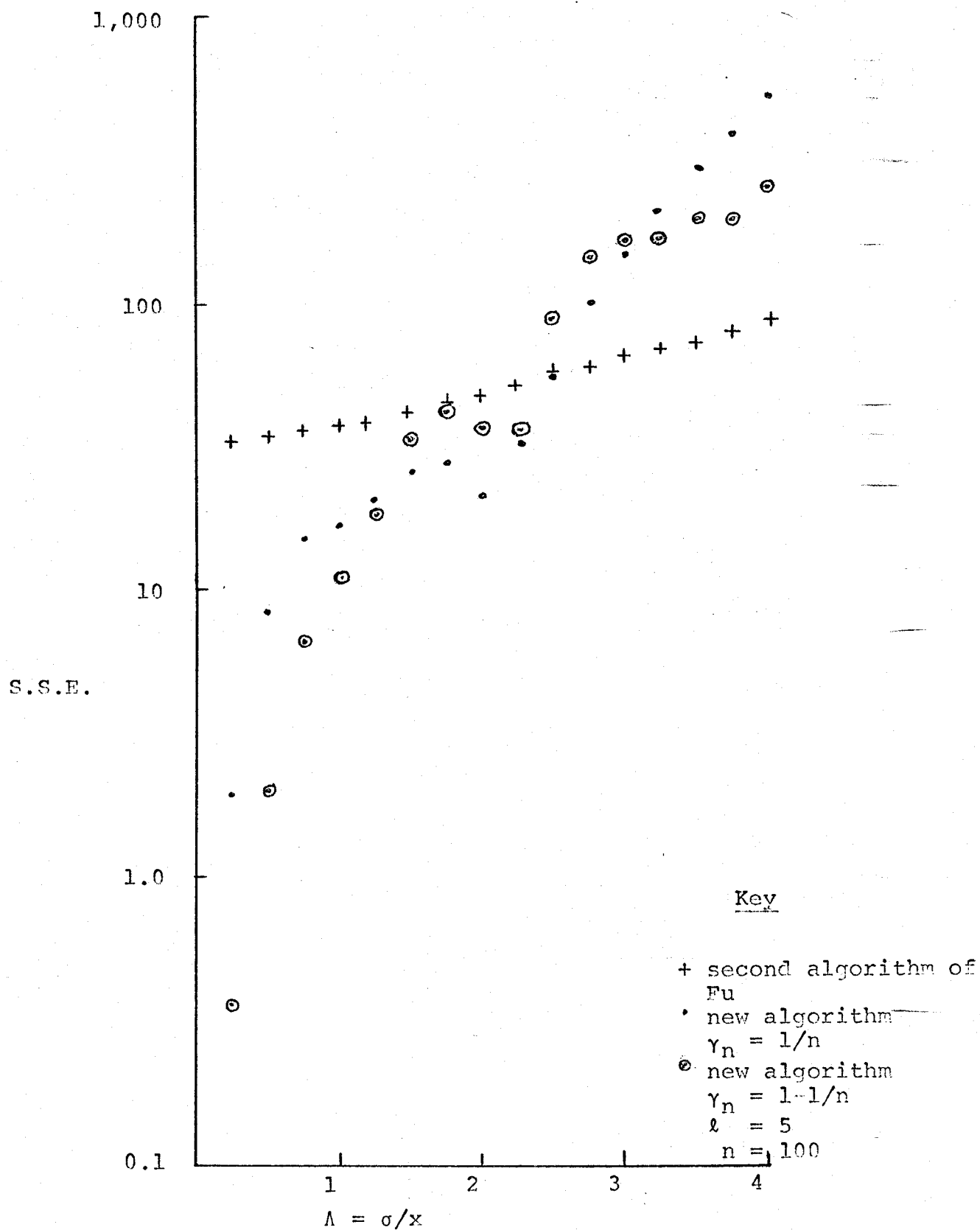


Graph (5.2.16) A.P., S.S.E. and T.S.S.E. as a function of n
 $\Lambda = 0.75$

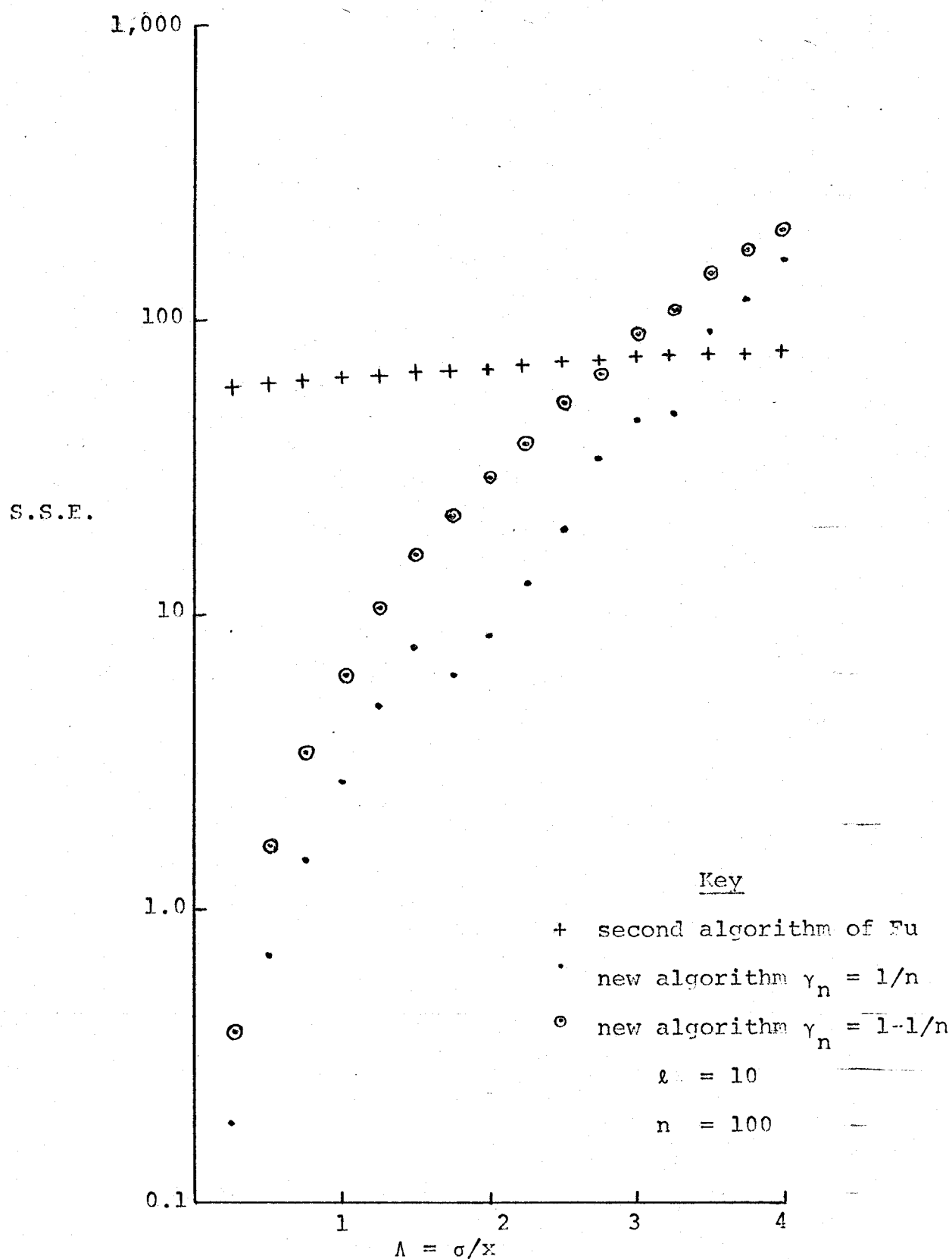
Now a more convenient evaluation of the simulation results can be made to give a more complete picture and a truer comparison of performance. What is done in the first instance is to plot the value of the sum of square errors (S.S.E.) after 100 iterations over the range of the Λ ratio as used in the simulation (Λ is defined in equation (5.2.1)). Only the second algorithm of Fu is used since it is the better of the two to use for comparison purposes. The new algorithm is thus evaluated for both Γ -sequences as defined in equations (5.1.7) and (5.1.8) in the following five graphs: (5.2.17), (5.2.18), (5.2.19), (5.2.20) and (5.2.21).

It should be noted that for both Γ -sequences there is a noise level, that is, a value of the Λ ratio at which the new algorithm has the same sum of square errors as the second algorithm of Fu. This value of the Λ ratio is called the value of equi-utility for the two algorithms being compared. As it turns out, for all values of the Λ ratio below the value of equi-utility, the new algorithm has a small sum of squared errors and also converges faster than Fu's second algorithm. For values of the Λ ratio above the value of equi-utility the new algorithm is not better than the existing techniques.

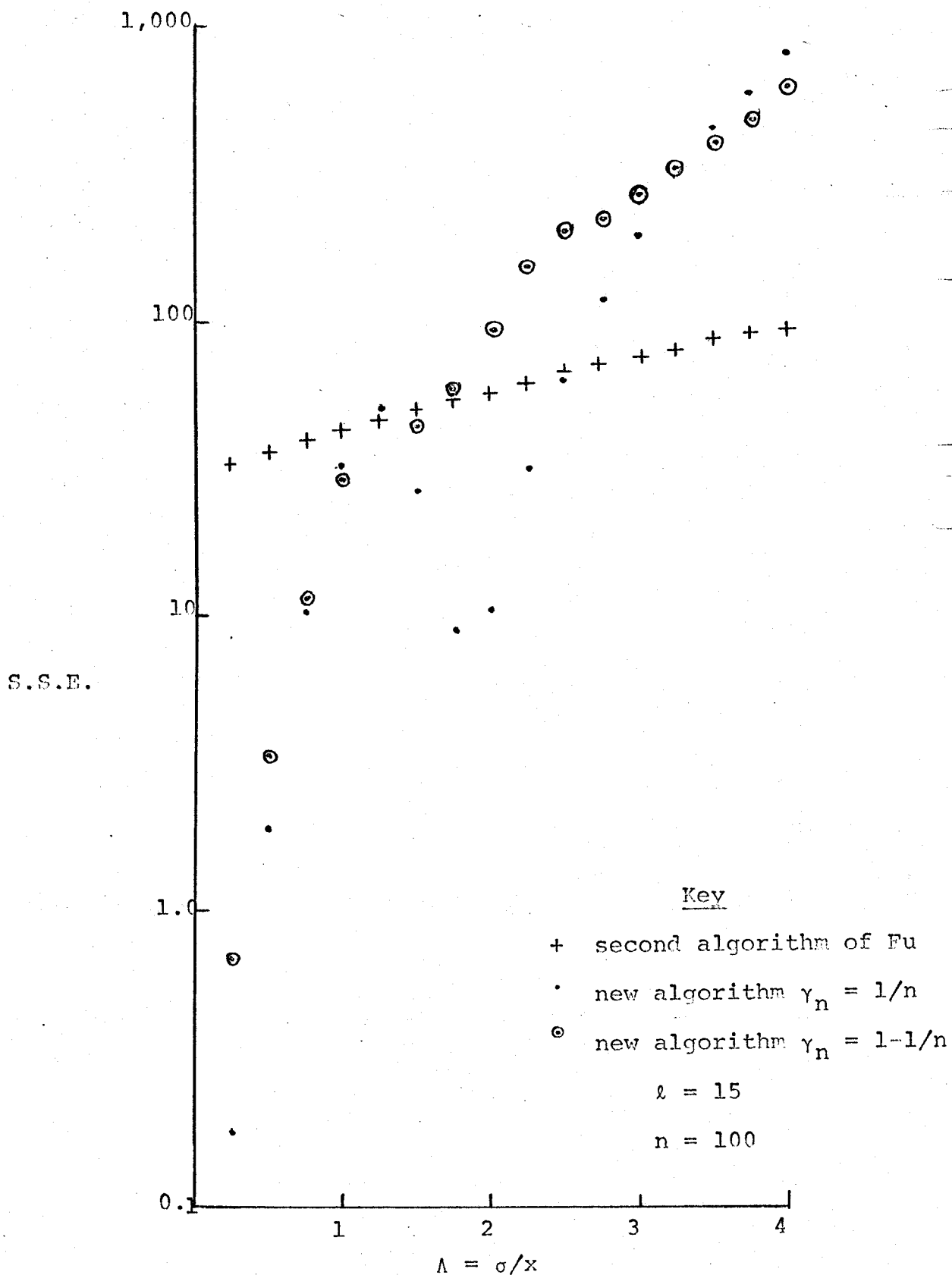
To make an evaluation of the point of equi-utility for both Γ -sequences and both error measures, that is, S.S.E. and T.S.S.E., the value of the point of equi-utility



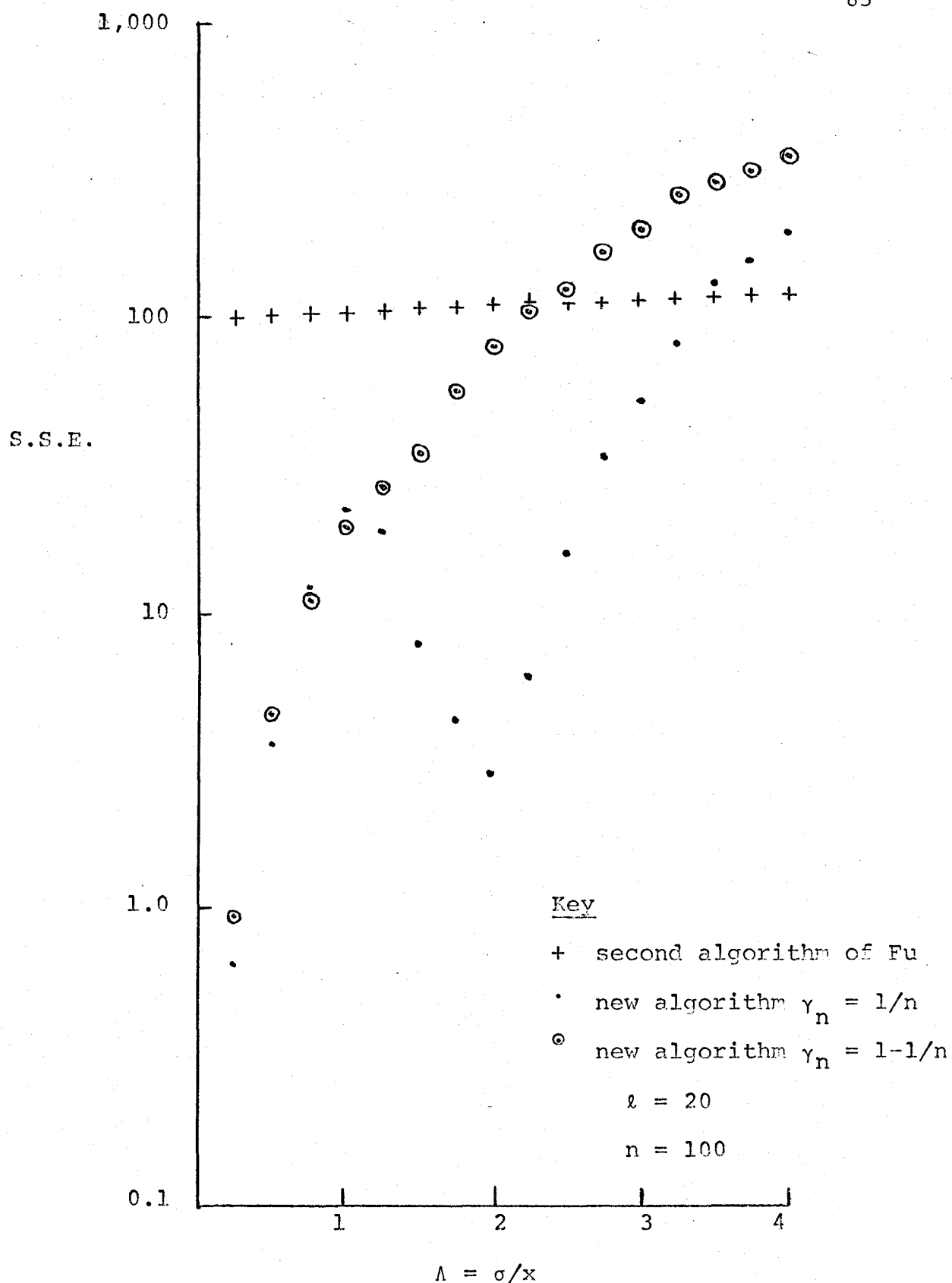
Graph (5.2.17) Sum of square errors (S.S.E.) as a function of Λ ratio



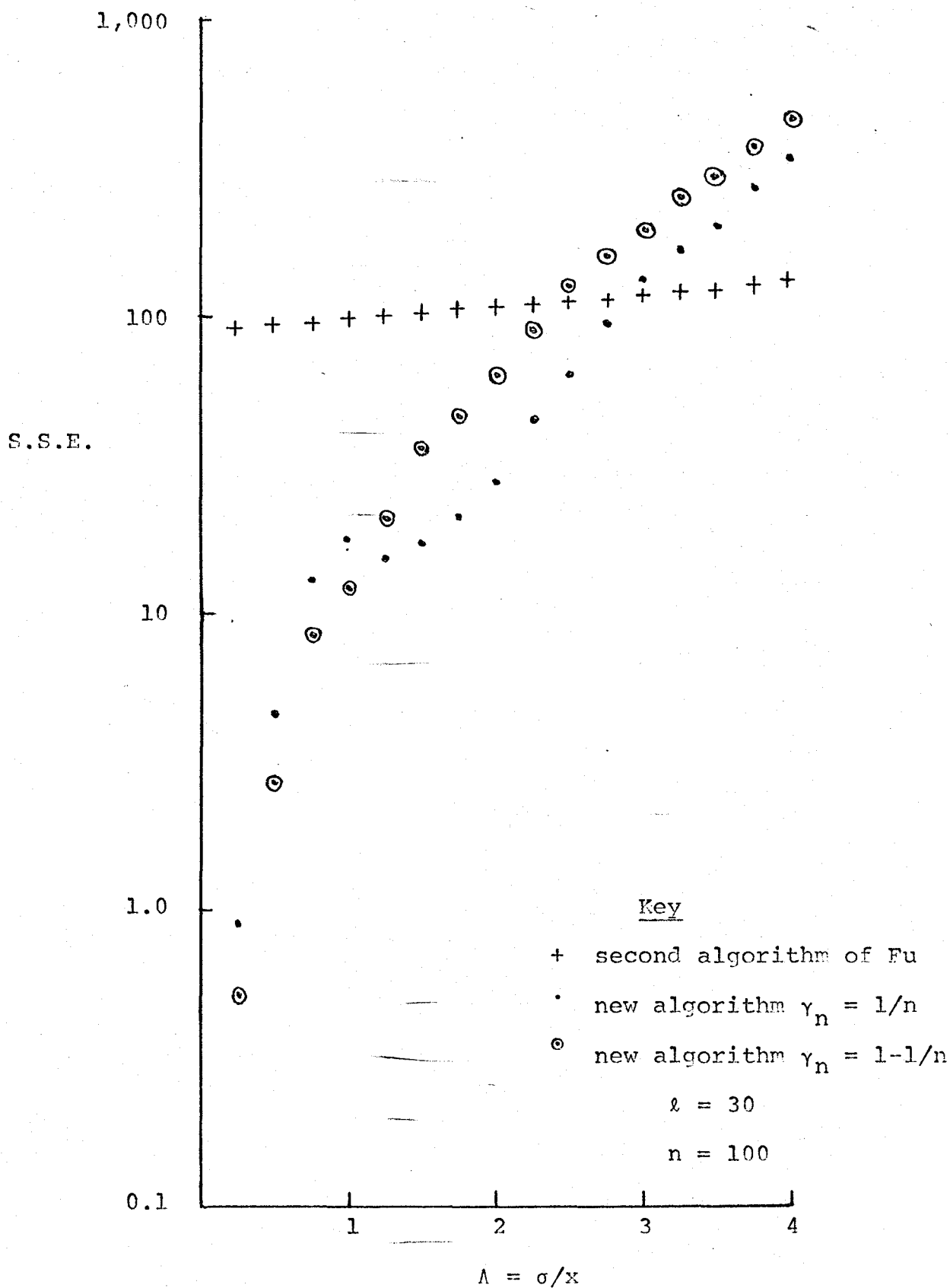
Graph (5.2.18) Sum of square errors (S.S.E.) as a function of Λ ratio



Graph (5.2.19) Sum of square errors (S.S.E.) as a function of Λ ratio



Graph (5.2.20) Sum of square errors (S.S.E.) as a function of Λ ratio



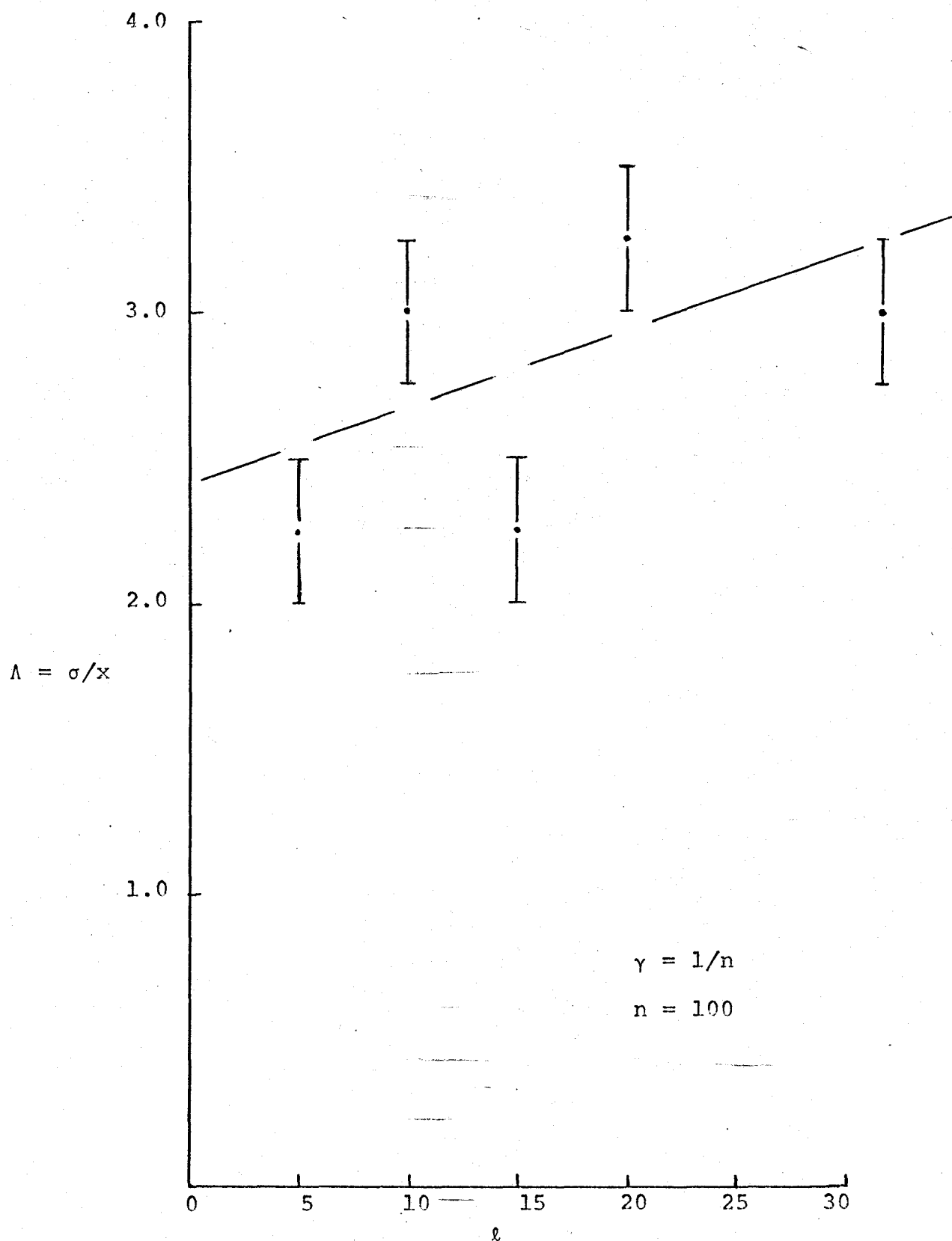
Graph (5.2.21) Sum of square errors (S.S.E.) as a function of Λ ratio

was plotted for a given Γ -sequence and error indicator after 100 iterations as a function of the sample autocorrelation delay ℓ . This is shown on the four graphs (5.2.22), (5.2.23), (5.2.24) and (5.2.25). It can be seen that if a longer sample autocorrelation delay is taken, then the range of Λ ratio over which the new algorithm is effective is increased. Both error indicators S.S.E. and T.S.S.E. bear out these facts.

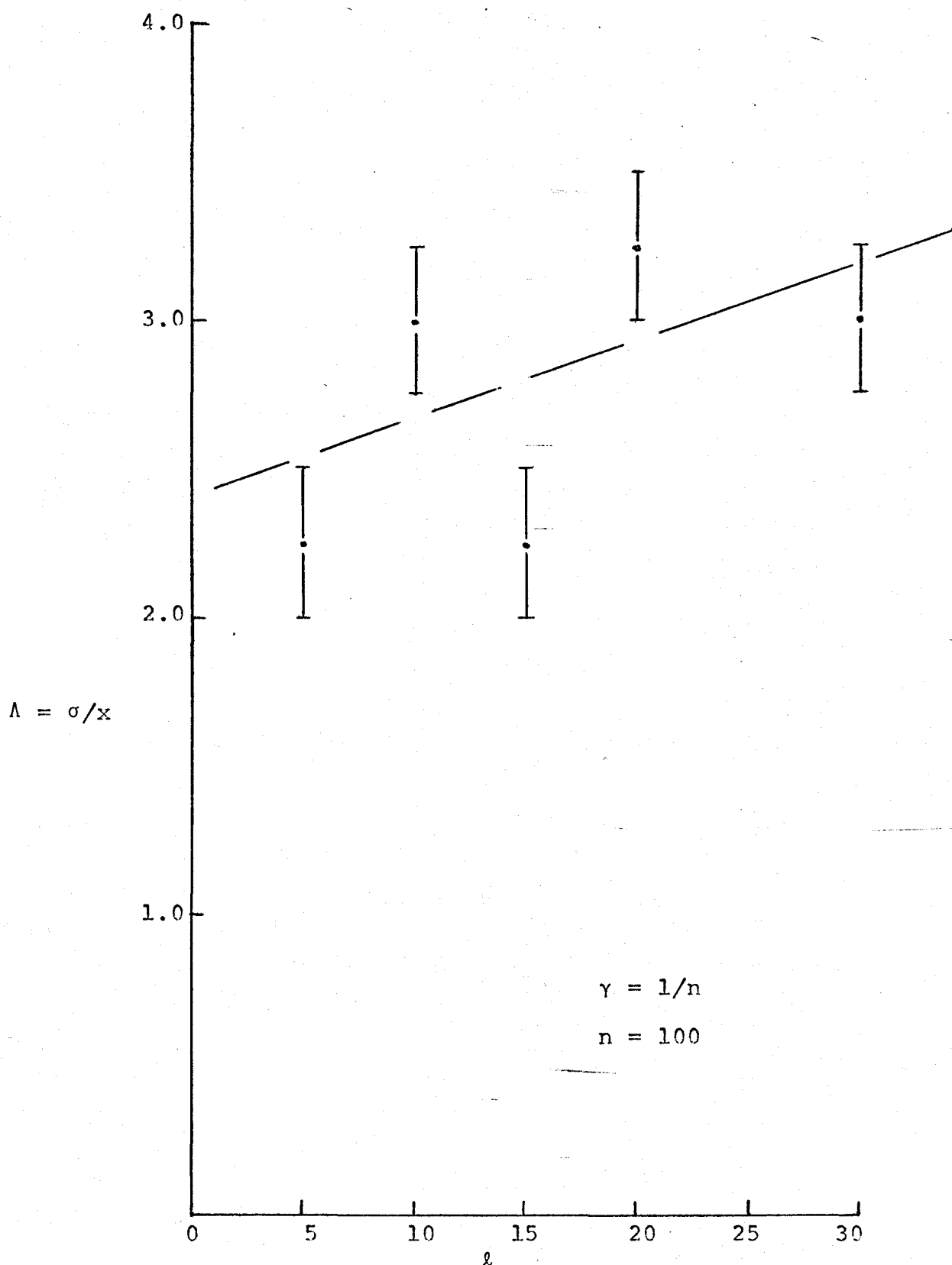
5.3 Summary of Results

The simulation of the algorithms has been so arranged as to test their relative merits and to establish that these merits are consistent. Two error criteria, the sum of square errors S.S.E. and the time sum of square errors T.S.S.E. has been used. Independently, each error measure evaluates the relative merits of each algorithm. Together they evaluate the consistent merits of the algorithms tested.

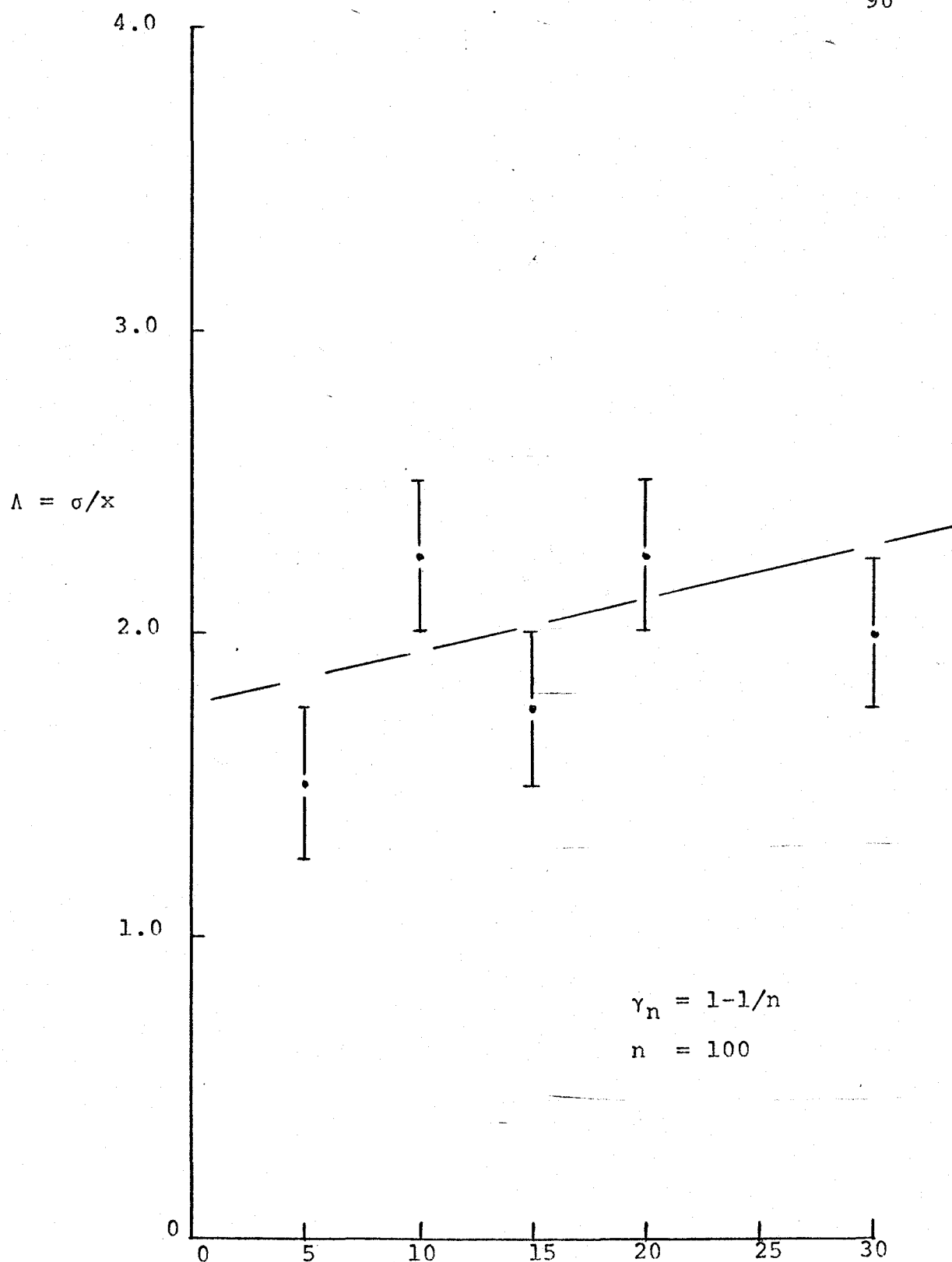
It has been shown that over a range of the Λ ratio, that is, $\Lambda \lesssim 2.5$, the new algorithm is of decided value. It is consistently of more value for this range as verified by the concurrence of both error measures. It is to this end that the particular choice of error measures made here has been taken for evaluation purposes.



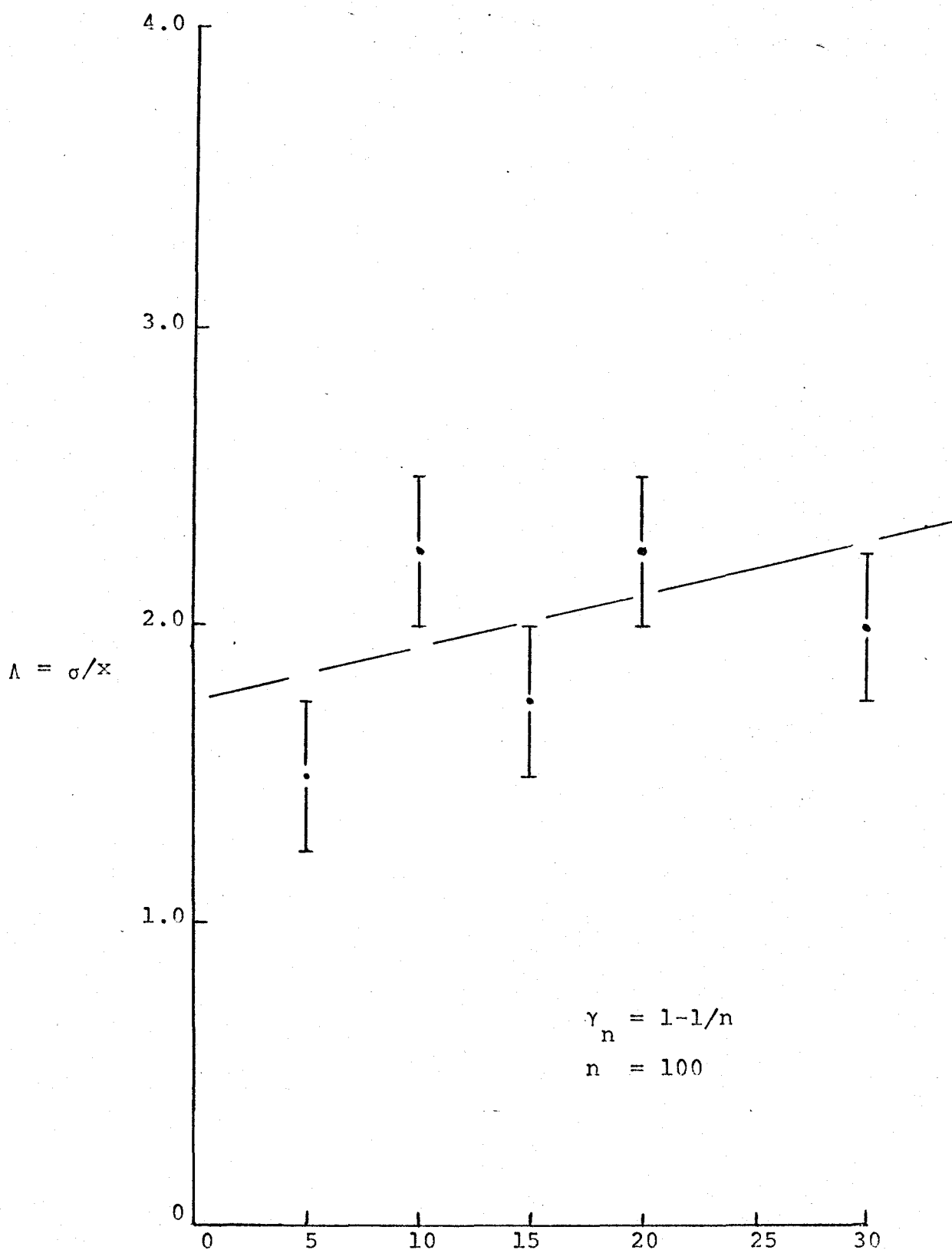
Graph (5.2.22) The Λ ratio of equi-utility for algorithms (5.1.4) and (5.1.5) based on S.S.F. evaluation as a function of sample autocorrelation delay



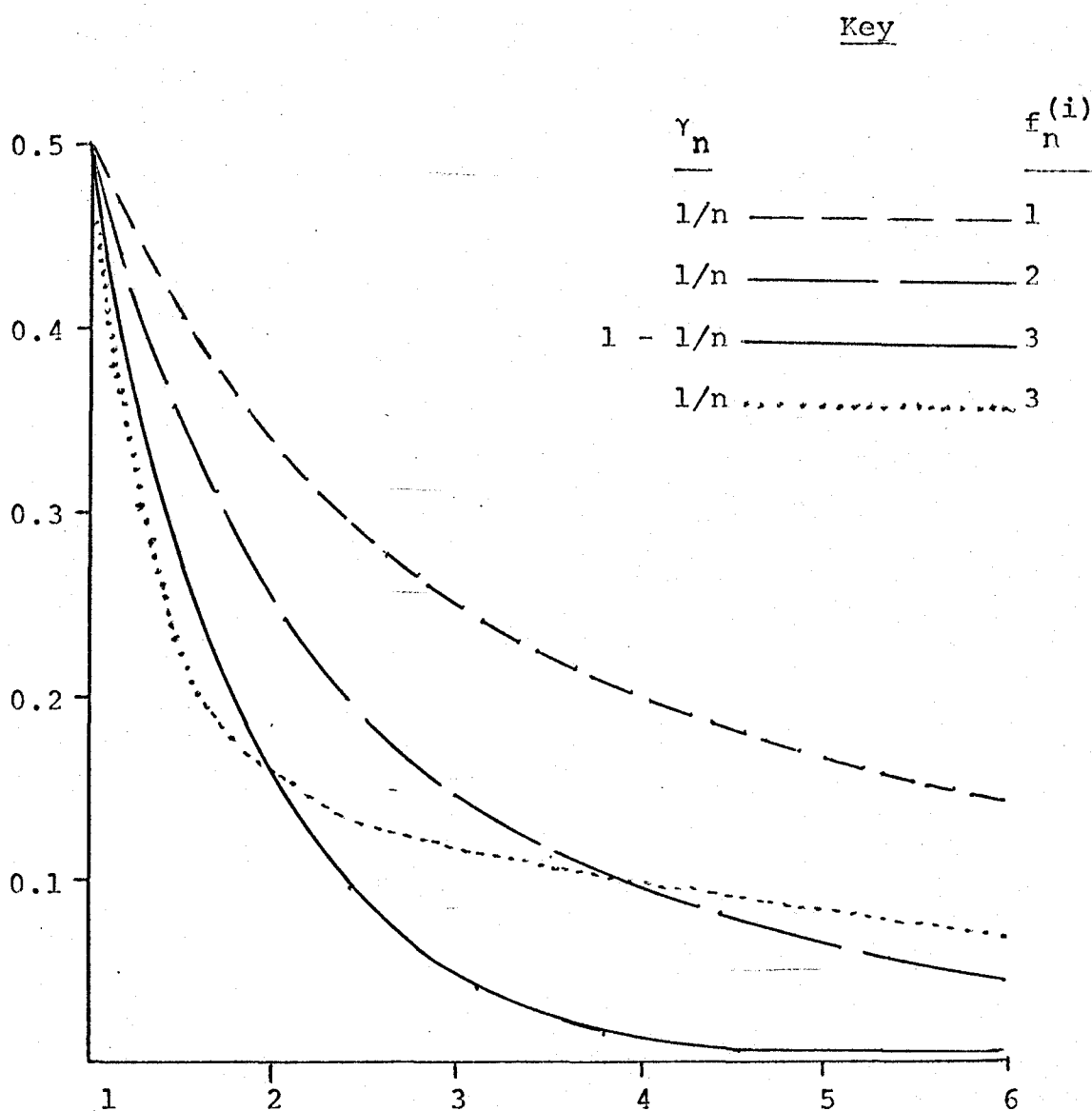
Graph (5.2.23) The ratio of equi-utility for algorithms (5.1.4) and (5.1.5) based on T.S.S.E. evaluation as a function of sample autocorrelation delay



Graph (5.2.24) The Λ ratio of equi-utility for algorithms (5.1.4) and (5.1.5) based on S.S.E. evaluation as a function of sample autocorrelation delay



Graph (5.2.25) The Λ ratio of equi-utility for algorithms (5.1.4) and (5.1.5) based on T.S.S.E. evaluation as a function of sample autocorrelation delay



Graph (5.2.26) Normalized Expected Mean Square Error as a Function of n .

$$(\sigma^2 = 1)$$

CHAPTER VI

Alternate Proof of Convergence

6.1 Alternate Approach

The essence of the work to this point has been the proof of the convergence of an algorithm and the comparison of its basic convergence property with two existing algorithms. During the course of the actual development work on the proof for convergence, it was desired by the author to have a different check on the theory developed. It is to this end that pursuit of another direction of proof was tried.

To be a valid confirmation of the first proof of convergence as given in Chapter III the alternate proof had to take a different approach as well as make use of different techniques. It is with these basic concepts that the author set out to develop an alternate or check proof for the basic convergence property of the algorithm.

It was thought that since the algorithm presented was of the Dvoretzky type, it would be natural to use his theorem to illustrate the property of convergence. It would also follow, that if the algorithm presented here does satisfy Dvoretzky's theorem, then all the properties associated with this class of algorithms is also true.

First a statement of Dvoretzky's theorem will be made outlining the conditions and results of his theorem. Then a stage by stage proof will be given showing that the algorithm presented here satisfies all the requirements of Dvoretzky's theorem, and hence, the generality of his proof and the resultant properties apply.

6.2 Statement of Dvoretzky's Theorem

Dvoretzky's theorem on stochastic approximation pertains to methods for successive approximations of a sought value, when, because of the stochastic nature of the problem, the observations or measurements have certain errors. The essential idea is to think of the random element as noise superimposed on a convergent deterministic scheme. Then the approximation procedure appears as an estimation scheme in a noisy environment.

Dvoretzky considered a general stochastic approximation procedure and proved a theorem, the statement of which follows below.

Theorem: Let $\alpha_n(\rho_1, \dots, \rho_n)$, $\beta_n(\rho_1, \dots, \rho_n)$ and $\gamma_n(\rho_1, \dots, \rho_n)$ be non-negative measurable functions of real variables $\rho_1, \rho_2, \rho_3, \dots, \rho_n$ satisfying the following conditions:

$$\lim_{n \rightarrow \infty} \alpha_n(\rho_1, \dots, \rho_n) = 0 \quad (6.2.1)$$

for a sequence $\rho_1, \rho_2, \dots, \rho_n$.

The sum of the series

$$\sum_{n=1}^{\infty} \beta_n(\rho_1, \rho_2, \dots, \rho_n) < \infty \quad (6.2.2)$$

is bounded and converges for any sequence $\rho_1, \rho_2, \dots$

$$\sum_{n=1}^{\infty} \gamma_n(\rho_1, \rho_2, \dots, \rho_n) = \infty \quad (6.2.3)$$

uniformly diverges for any sequence $\rho_1, \rho_2, \dots$, bounded in absolute value, that is, for any sequence $\rho_1, \rho_2, \dots$ such that

$$\sup_{n=1,2,\dots} |\rho_n| < c \quad (6.2.4)$$

c being an arbitrary finite number.

Let θ be a real number, and $T_1, T_2, \dots$ be measurable transformations, satisfying the inequality

$$|T_n(\rho_1, \rho_2, \dots, \rho_n) - \theta| \leq \max [\alpha_n, (1+\beta_n)|\rho_n - \theta| - \gamma_n] \quad (6.2.5)$$

for any real sequence $\rho_1, \rho_2, \dots$

Further let X_1 and $Y_1, Y_2, \dots$ be random variables, and for all $n \geq 1$ let,

$$X_{n+1} = T_n(X_1, X_2, \dots, X_n) + Y_n + g_n(X_1, X_2, \dots, X_n) \quad (6.2.6)$$

where $g_n(\rho_1, \rho_2, \dots, \rho_n)$ are measurable functions such that the sum of the series

$$\sum_{n=1}^N |g_n(\rho_1, \rho_2, \dots, \rho_n)| \quad (6.2.7)$$

uniformly converges and its sum is uniformly bounded for any $\rho_1, \rho_2, \dots$ and that

$$E\{Y_n | X_1, X_2, \dots, X_n\} = 0 \quad (6.2.8)$$

with probability 1. The series

$$\sum_{n=1}^{\infty} E\{Y_n^2\} < \infty \quad (6.2.9)$$

converges and

$$E\{X_1^2\} < \infty \quad (6.2.10)$$

Then, as $n \rightarrow \infty$

$$P\{\lim_{n \rightarrow \infty} X_n = 0\} = 1 \quad (5.2.11)$$

and

$$\lim_{n \rightarrow \infty} E\{(X_n - 0)^2\} = 0 \quad (5.2.12)$$

Extension: The theorem remains valid if α_n , β_n and γ_n in (6.2.5) are replaced by non-negative functions $\alpha_n(\rho_1, \rho_2, \dots, \rho_n)$, $\beta_n(\rho_1, \rho_2, \dots, \rho_n)$ and $\gamma_n(\rho_1, \rho_2, \dots, \rho_n)$, respectively, provided they satisfy the conditions: The functions $\alpha_n(\rho_1, \rho_2, \dots, \rho_n)$ are bounded and

$$\lim_{n \rightarrow \infty} \alpha_n(\rho_1, \rho_2, \dots, \rho_n) = 0 \quad (6.2.13)$$

uniformly for all sequences $\rho_1, \rho_2, \dots, \rho_n$. The function $\beta_n(\rho_1, \rho_2, \dots, \rho_n)$ are measurable and

$$\sum_{n=1}^{\infty} \beta_n(\rho_1, \rho_2, \dots, \rho_n) < \infty \quad (6.2.14)$$

is bounded and convergent for all sequences $\rho_1, \rho_2, \dots, \rho_n$ and the functions $\gamma_n(\rho_1, \rho_2, \dots, \rho_n)$ satisfy,

$$\sum_{n=1}^{\infty} \gamma_n(\rho_1, \rho_2, \dots, \rho_n) = \infty \quad (6.2.15)$$

uniformly for all sequences $\rho_1, \rho_2, \dots, \rho_n$ for which

$$\sup_{n=1,2,\dots} |\rho_n| < L \quad (6.2.16)$$

L being an arbitrary finite number.

In addition to this extension, there are five generalizations to the theorem*. By applying these generalizations and the theorem with extensions, it will be possible to prove the convergence of the algorithm present in this work.

6.3 Convergence using Dvoretzky's Theorem

Consider the algorithm presented here as

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{ (\hat{R}_{r_{n+1}}(\ell))^{1/2} - \hat{x}_n \} \quad n=1,2,\dots \quad (6.3.1)$$

where $\hat{x}_{n+1}$ is the $n+1^{\text{th}}$ estimate of

x the true value being sought

γ_{n+1} is a non-negative real number (gain-sequence)

r_n are the measurements from the distribution of x

and $\hat{R}_{r_{n+1}}(\ell)$ is the sample autocorrelation function.

*Proof of Dvoretzky's Theorem, his extension and generalizations given in Appendix D.

Now subtracting x from both sides of (6.3.1) gives the following equation

$$(\hat{x}_{n+1} - x) = (1 - \gamma_{n+1})(\hat{x}_n - x) + \gamma_{n+1}\psi_{n+1} \quad (6.3.2)$$

$n=1, 2, \dots$

where as before

$$\psi_{n+1} = x \left\{ \left(1 + \frac{\hat{R}_{\xi_{n+1}}(\ell)}{2} \right)^{\frac{1}{2}} - 1 \right\} \quad (6.3.3)$$

consider the transformation

$$(\hat{x}_n - x) = \hat{w}_n \quad (6.3.4)$$

and applying it to (6.3.2) gives the following result;

$$\hat{w}_{n+1} = (1 - \gamma_{n+1})\hat{w}_n + \gamma_{n+1}\psi_{n+1} \quad (6.3.5)$$

Equation (6.3.5) or its equivalent (6.3.2) is a zero seeking algorithm. Now consider the convergence of this transformed algorithm.

Let $\alpha_n, \beta_n, \gamma_n$ and $n=1, 2, 3, \dots$ be non-negative real numbers satisfying condition (6.2.1)

$$\lim_{n \rightarrow \infty} \alpha_n(\rho_1, \rho_2, \dots, \rho_n) = 0$$

condition (6.2.2)

$$\sum_{n=1}^{\infty} \beta_n(\rho_1, \rho_2, \dots, \rho_n) < \infty$$

and condition (6.2.3)

$$\sum_{n=1}^{\infty} \gamma_n(\rho_1, \rho_2, \dots, \rho_n) = \infty$$

Making the selection for γ_n as in (3.4.1), recall this as

$$\Gamma = \{\gamma_i \in \Gamma \mid \gamma_i = \frac{1}{i+\alpha} \quad i=1,2,\dots\},$$

and letting $\alpha=0$, then α_n and β_n can be selected as

$$A = \{\alpha_i \in A \mid \alpha_i = 1/i \quad \forall i=1,2,\dots\} \quad (6.3.6)$$

and

$$B = \{\beta_i \in B \mid \beta_i = 1/i^2 \quad \forall i=1,2,\dots\} \quad (6.3.7)$$

Now let us consider α_n as just defined. Since

$$\lim_{n \rightarrow \infty} \alpha_n = \lim_{n \rightarrow \infty} 1/n$$

and

$$\lim_{n \rightarrow \infty} 1/n \rightarrow 0, \quad (6.3.8)$$

then by (6.3.8) the selection of α_n in (6.3.6) satisfies condition (6.2.1) of the theorem.

Now consider the selection β_n in (6.3.7) and using Cauchy's integral test (see Appendix E for proof), write

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{1}{y^2} dy = \lim_{\tau \rightarrow \infty} [-y^{-1}]_1^{\tau}$$

and hence

$$\lim_{\tau \rightarrow \infty} [-y^{-1}]_1^{\tau} = \lim_{\tau \rightarrow \infty} \left[\frac{1}{1} - \frac{1}{\tau} \right]$$

and thus

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{1}{y^2} dy = 1 \quad (6.3.9)$$

This proves that the integral converges and hence by Cauchy's integral test the summation

$$\sum_{n=1}^{\infty} \beta_n$$

is bounded. By further use of Cauchy's integral test the bounds on the summation can be established by applying

$$\int_1^{k+1} \frac{1}{y^2} dy < \sum_{n=1}^k \frac{1}{n^2} < \int_1^k \frac{1}{y^2} dy + 1 \quad (6.3.10)$$

Taking the limit as k becomes very large and using result (6.3.9) gives the bounds on the summation as

$$1 < \sum_{n=1}^{\infty} \frac{1}{n^2} < 2 \quad (6.3.11)$$

Hence by result (6.3.9) and (6.3.11) the selection of n made in (6.3.7) satisfies condition (6.2.2) of Dvoretzky's Theorem.

Now consider the choice of γ_n as in (3.4.1) and recalled previously. By applying Cauchy's integral test, the following integral can be written:

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{1}{y} dy = \lim_{\tau \rightarrow \infty} \ln y \Big|_1^{\tau}$$

and hence

$$\lim_{\tau \rightarrow \infty} \ln y \Big|_1^{\tau} = \lim_{\tau \rightarrow \infty} \ln \tau$$

and thus

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{1}{y} dy \rightarrow \infty \quad (6.3.12)$$

Since the integral is unbounded, then by Cauchy's integral theorem the sum

$$\sum_{n=1}^{\infty} \gamma_n = \infty \quad (6.3.13)$$

is without bound; hence, by result (6.3.12) and (6.3.13) condition (6.2.3) of the theorem is satisfied.

By the selection of the α -sequence, β -sequence and Γ -sequence as stated in (6.3.6), (6.3.7) and (3.4.1) respectively, all the members of all three sets satisfy condition (6.2.4) of the theorem simply by definition of the sets themselves.

In the preamble to the generalization of Dvoretzky's theorem* it is stated that for a zero estimating scheme like (6.3.5), condition (D-4.1) is stronger than condition (6.2.5). Hence, by using condition (D-4.1) of the preamble instead of condition (6.2.5) in fact strengthens the generality of the theorem. To apply this condition, begin by

identifying $T_n(\rho_1, \rho_2, \dots, \rho_n)$ with $(1 - \gamma_{n+1})^{\wedge} W_n$
 and ρ_n with $\hat{W}_n$

*See Appendix D

then

$$|T_n(\rho_1, \rho_2, \dots, \rho_n)| \leq \max [\alpha_n, (1+\beta_n - \gamma_n) |\rho_n|] \quad (6.3.14)$$

is satisfied since

$$|(1-\gamma_{n+1})\hat{W}_n| = |(1-\frac{1}{n+1})\hat{W}_n|$$

and

$$|(1-\gamma_{n+1})\hat{W}_n| = (1-\frac{1}{n+1})|\hat{W}_n| \quad \text{for } n=1,2,3,\dots \quad (6.3.15)$$

and

$$(1+\beta_n - \gamma_n) |\rho_n| = (1+\frac{1}{n} - \frac{1}{n}) |\hat{W}_n| \quad \text{for } n=1,2,3,\dots \quad (6.3.16)$$

Now

$$n^3 < n^3 + 1 \quad \text{for } n=1,2,3,\dots \quad (6.3.17)$$

But

$$(n^3+1) = (n+1)(n^2-n+1) \quad (6.3.18)$$

Hence

$$n^3 < (n+1)(n^2-n+1) \quad (6.3.19)$$

Now, if both sides of (6.3.19) are multiplied by $\frac{1}{n^2(n+1)} \quad \forall n > 0$ the result is

$$\frac{n}{n+1} < \frac{n^2+1-n}{n^2} \quad (6.3.20)$$

If $1/n+1$ is added and subtracted to the left hand side of (6.3.20) giving

$$\frac{n}{n+1} + \frac{1}{n+1} - \frac{1}{n+1} < \frac{n^2+1-n}{n^2} \quad (6.3.21)$$

and regrouping terms results in the following equation

$$(1-1/n+1) < (1+1/n^2-1/n) \quad (6.3.22)$$

Multiplying both sides of equation (6.3.22) by $|\hat{W}_n|$ gives,

$$(1-\frac{1}{n+1}) |\hat{W}_n| \leq (1+1/n^2-1/n) |\hat{W}_n| \quad (6.3.23)$$

Here the equality is introduced to include the case of $|\hat{W}_n| = 0$. Now since $(1-1/n+1)$ is non-negative for all $n=1,2,3,\dots$, it can then be taken within the modulus sign in equation (6.3.23) giving

$$|(1-1/n+1)\hat{W}_n| \leq (1+1/n^2-1/n) |\hat{W}_n| \quad (6.3.24)$$

This equation (6.3.24) satisfies equation (6.3.14) and hence satisfying the condition (D-4.1) of the theorem. If it should occur that

$$\alpha_n \leq (1+\beta_n-\gamma_n) |\rho_n| \quad (6.3.25)$$

in equation (6.3.14), this condition and (D-4.1) holds by the result in equation (6.3.24). Alternately, if it should occur that

$$\alpha_n > (1+\beta_n-\gamma_n) |\rho_n| \quad (6.3.26)$$

then (6.3.14) and (D-4.1) hold by virtue of (6.3.26) and result (6.3.24). Hence condition (D-4.1) of the theorem is

always satisfied.

Now identifying condition (6.2.6) with equation (6.3.5) yields

$$Y_n = \gamma_{n+1} \psi_{n+1} \quad (6.3.27)$$

where as defined in (3.3.12) and (6.3.3)

$$\psi_{n+1} = x \left\{ \left(1 + \frac{\hat{R}_{\xi_{n+1}}^{(\ell)}}{x^2} \right)^{\frac{1}{2}-1} \right\} \quad (6.3.28)$$

Also, since the expectation operator $E\{.\}$ is linear, it commutes with the summation operator " $\sum$ " which is also linear. Hence the condition (6.2.9)

$$\sum_{n=1}^{\infty} E\{Y_n^2\} < \infty$$

can be rewritten as

$$E\left\{ \sum_{n=1}^{\infty} Y_n^2 \right\} < \infty \quad (6.3.29)$$

combining the contents of equation (6.3.28) with equation (6.3.27) gives

$$Y_n^2 = \gamma_{n+1}^2 x^2 \left\{ \left(1 + \frac{\hat{R}_{\xi_{n+1}}^{(\ell)}}{x^2} \right)^{\frac{1}{2}-1} \right\}^2 \quad (6.3.30)$$

Squaring the contents of the brackets and multiplying through by x^2 gives the equation

$$Y_n^2 = 2\gamma_{n+1}^2 x^2 - 2\gamma_{n+1}^2 x^2 \left(1 + \frac{\hat{R}_{\xi_{n+1}}^{(\ell)}}{x^2} \right)^{\frac{1}{2}} + \gamma_{n+1}^2 \hat{R}_{\xi_{n+1}}^{(\ell)} \quad (6.3.31)$$

Applying equation (6.3.31) to condition (6.3.29) and investigating term by term can be done by first investigating the 1st term of (6.3.31)

$$2x^2 \sum_{n=1}^{\infty} \gamma_{n+1}^2 = 2x^2 \sum_{n=1}^{\infty} \frac{1}{(n+1)^2} \quad (6.3.32)$$

Applying Cauchy's integral test to (6.3.32) gives

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{1}{(1+y)^2} dy = \lim_{\tau \rightarrow \infty} \int_1^{\tau} (y+1)^{-2} dy$$

and hence

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} (y+1)^{-2} dy = \lim_{\tau \rightarrow \infty} -(y+1)^{-1} \Big|_1^{\tau}$$

resulting in

$$\lim_{\tau \rightarrow \infty} -(y+1)^{-1} \Big|_1^{\tau} = \lim_{\tau \rightarrow \infty} \left[-\frac{1}{\tau+1} + \frac{1}{2} \right]$$

and therefore

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{1}{(1+y)^2} dy = \frac{1}{2} \quad (6.3.33)$$

Thus, the series in (6.3.32) has a bounded and convergent sum. By further application of Cauchy's integral test the bounds on the sum can be established. They are given by,

$$\int_1^{k+1} \frac{1}{(y+1)^2} dy < \sum_{n=1}^k \frac{1}{(n+1)^2} < \int_1^k \frac{1}{(y+1)^2} dy + \frac{1}{(y+1)^2} \Big|_{y=1} \quad (6.3.34)$$

Taking the limit as k becomes very large and using result (6.3.33) gives the bounds as

$$\frac{1}{2} < \lim_{k \rightarrow \infty} \sum_{n=1}^k \frac{1}{(n+1)^2} < \frac{1}{2} + \frac{1}{4} \quad (6.3.35)$$

which when multiplied through by $2x^2$ puts the bounds on the first term of (6.3.31) as,

$$x^2 < 2x^2 \sum_{n=1}^{\infty} \frac{1}{(n+1)^2} < \frac{3x^2}{2} \quad (6.3.36)$$

Thus for any finite x the first term of (6.3.31) in (6.3.29) is finite.

Consider the sum of the second term in equation (6.3.31).

It can be written as

$$\lim_{k \rightarrow \infty} \sum_{n=1}^k 2x^2 \gamma_{n+1}^2 \left(1 + \frac{\hat{R}_{\xi_{n+1}}^{(\ell)}}{x^2}\right)^{\frac{1}{2}} = \lim_{k \rightarrow \infty} 2x^2 \sum_{n=1}^k \frac{1}{(n+1)^2} \left(1 + \frac{\hat{R}_{\xi_{n+1}}^{(\ell)}}{x}\right)^{\frac{1}{2}} \quad (6.3.37)$$

where

$$\hat{R}_{\xi_{n+1}}^{(\ell)} = \frac{1}{n+1} \sum_{k=1}^{n+1-\ell} \xi_k \xi_{k+\ell} \quad (6.3.38)$$

Now since in all practical situations the upper value of a measured sample is limited, hence one can write

$$\sup |r_i| \leq L \quad (6.3.39)$$

where L is a large arbitrary, finite number. Hence, the upper value of an element of noise is also limited and thus it can be written that

$$\sup |\xi_i| \leq M \quad (6.3.40)$$

where

$$M = L - x \quad (6.3.41)$$

Hence, combining equation (6.3.37) and equation (6.3.38)

and using condition (6.3.40) gives

$$\lim_{k \rightarrow \infty} 2x^2 \sum_{n=1}^k \frac{1}{(n+1)^2} \left(1 + \frac{\hat{R}_{\xi_{n+1}}(\ell)}{x^2}\right)^{\frac{1}{2}} \leq \lim_{k \rightarrow \infty} 2x^2 \sum_{n=1}^k \frac{1}{(n+1)^2} \left(1 + \frac{(n+1)M^2}{x^2}\right)^{\frac{1}{2}} \quad (6.3.42)$$

Applying Cauchy's integral test to (6.3.42) gives the integral,

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{2x^2}{1(y+1)^2} \left(1 + \frac{(y+1)M^2}{x^2}\right)^{\frac{1}{2}} dy = \lim_{\tau \rightarrow \infty} 2x \int_1^{\tau} \frac{1}{1(y+1)^2} (x^2 + (y+1)M^2)^{\frac{1}{2}} dy \quad (6.3.43)$$

Now consider an integral of the following form and integrate it by parts

$$\int \frac{\sqrt{ax+b}}{x^2} dx = \frac{\sqrt{ax+b}}{-x} + \frac{1}{2} a \cdot \frac{dx}{x\sqrt{ax+b}} \quad (6.3.44)$$

Applying a tabulated integral* to equation (6.3.44) gives the following,

$$\int \frac{\sqrt{ax+b}}{x^2} dx = -\frac{\sqrt{ax+b}}{x} + \frac{a}{2} \left[\frac{1}{\sqrt{b}} \log \frac{\sqrt{ax+b} - \sqrt{b}}{\sqrt{ax+b} + \sqrt{b}} \right] \quad (6.3.45)$$

Applying this result to Cauchy's integral in equation (6.3.43) results in the following equation:

*#72 of Handbook of Mathematical Tables and Formulas by
R. S. Burington

$$\lim_{\tau \rightarrow \infty} 2x \int_1^{\tau} \frac{1}{(y+1)^2} [x^2 + (y+1)M^2]^{\frac{1}{2}} dy = \lim_{\tau \rightarrow \infty} \{M^2 \log \frac{[x^2 + (\tau+1)M^2]^{\frac{1}{2}} - x}{[x^2 + (\tau+1)M^2]^{\frac{1}{2}} + x} - \frac{[x^2 + (y+1)M^2]^{\frac{1}{2}}}{(y+1)}\}_1^{\tau} \quad (6.3.46)$$

Substituting the limits into equation (6.3.46) gives,

$$\lim_{\tau \rightarrow \infty} 2x \int_1^{\tau} \frac{1}{(y+1)^2} [x^2 + (y+1)M^2]^{\frac{1}{2}} dy = \lim_{\tau \rightarrow \infty} \{M^2 \left[\log \frac{[x^2 + (\tau+1)M^2]^{\frac{1}{2}} - x}{[x^2 + (\tau+1)M^2]^{\frac{1}{2}} + x} - \log \frac{[x^2 + 2M^2]^{\frac{1}{2}} - x}{[x^2 + 2M^2]^{\frac{1}{2}} + x} \right] - \left[\frac{[x^2 + (\tau+1)M^2]^{\frac{1}{2}}}{\tau+1} - \frac{[x^2 + 2M^2]^{\frac{1}{2}}}{2} \right]\} \quad (6.3.47)$$

which reduces to

$$\begin{aligned} \lim_{\tau \rightarrow \infty} 2x \int_1^{\tau} \frac{1}{(y+1)^2} [x^2 + (y+1)M^2]^{\frac{1}{2}} dy \\ = \lim_{\tau \rightarrow \infty} \{M^2 \log \left[\frac{\{[x^2 + (\tau+1)M^2]^{\frac{1}{2}} - x\}}{\{[x^2 + (\tau+1)M^2]^{\frac{1}{2}} + x\}} \frac{\{[x^2 + 2M^2]^{\frac{1}{2}} + x\}}{\{[x^2 + 2M^2]^{\frac{1}{2}} - x\}} \right] \\ - \left[\frac{[x^2 + (\tau+1)M^2]^{\frac{1}{2}}}{\tau+1} - \frac{[x^2 + 2M^2]^{\frac{1}{2}}}{2} \right]\} \end{aligned} \quad (6.3.48)$$

In the limit equation (6.3.49) becomes

$$\lim_{\tau \rightarrow \infty} 2x \int_1^{\tau} \frac{1}{(y+1)^2} [x^2 + (y+1)M^2]^{\frac{1}{2}} dy = \frac{[x^2 + 2M^2]^{\frac{1}{2}}}{2} - 1 + M^2 \log \frac{\{[x^2 + 2M^2]^{\frac{1}{2}} + x\}}{\{[x^2 + 2M^2]^{\frac{1}{2}} - x\}} \quad (6.3.49)$$

Hence for a finite M , the second term of (6.3.31) has a finite sum. The bounds on the sum are

$$\begin{aligned} - \left[2x^2 \int_1^{k+1} \frac{1}{(y+1)^2} \left(1 + \frac{(y+1)M^2}{x^2} \right)^{\frac{1}{2}} dy \right] < 2x^2 \sum_{m=1}^k \frac{1}{(n+1)^2} \left(1 + \frac{\hat{R}_{\xi_{n+1}}(\ell)}{x^2} \right)^{\frac{1}{2}} \\ < 2x^2 \int_1^k \frac{1}{(y+1)^2} \left(1 + \frac{(y+1)M^2}{x^2} \right)^{\frac{1}{2}} dy + \frac{x^2}{2} \left(1 + \frac{\hat{R}_{\xi_2}(\ell)}{x^2} \right)^{\frac{1}{2}} \end{aligned} \quad (6.3.50)$$

Taking the limit in equation (6.3.50) as k becomes very large gives after using result (6.3.49) the following equation,

$$-\left[\frac{x^2 + 2M^2}{2} - 1 + M^2 \log \frac{\{[x^2 + 2M^2]^{\frac{1}{2}} + x\}}{\{[x^2 + 2M^2]^{\frac{1}{2}} - x\}} \right] < 2x^2 \sum_{n=1}^{\infty} \frac{1}{(n+1)^2} \left(1 + \frac{\hat{R}_{\xi_{n+1}}^{(\ell)}}{x^2}\right)^{\frac{1}{2}} \\ < \frac{[x^2 + 2M^2]^{\frac{1}{2}}}{2} - 1 + M^2 \log \frac{\{[x^2 + 2M^2]^{\frac{1}{2}} + x\}}{\{[x^2 + 2M^2]^{\frac{1}{2}} - x\}} + \frac{x^2}{2} \left(1 + \frac{M^2}{x^2}\right)^{\frac{1}{2}} \quad (6.3.51)$$

Therefore, by result (6.3.50) and result (6.3.51), it can be seen that the second term of (6.3.31) has a finite and bounded infinite sum for a finite value of M .

Considering the sum of the third term of (6.3.31) gives

$$\sum_{n=1}^{\infty} \gamma_{n+1}^2 \hat{R}_{\xi_{n+1}}^{(\ell)} = \sum_{n=1}^{\infty} \frac{1}{(n+1)^2} \hat{R}_{\xi_{n+1}}^{(\ell)} \quad (6.3.52)$$

Applying Cauchy's integral test to (6.3.52) gives the integral

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{1}{(y+1)^2} \hat{R}_{\xi_{n+1}}^{(\ell)} dy \leq \lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{M^2}{(y+1)^2} dy \quad (6.3.53)$$

where condition (6.3.38) and condition (6.3.41) are used in a substitution. Now the integration yields

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{M^2}{(y+1)^2} dy = M^2 \lim_{\tau \rightarrow \infty} \left\{ -\frac{1}{(y+1)} \right\}_1^{\tau} \quad (6.3.54)$$

and after taking the limit

$$\lim_{\tau \rightarrow \infty} \int_1^{\tau} \frac{1}{(y+1)} \hat{R}_{\xi_{n+1}}^{(\ell)} dy \leq \frac{M^2}{2} \quad (6.3.55)$$

Therefore, for finite values of M , the infinite sum of the third term of (6.3.31) is bounded and convergent. The bounds can be given by further applying Cauchy's integral test resulting in

$$\int_1^{k+1} \frac{1}{(y+1)^2} \hat{R}_{\xi_{n+1}}^{(\ell)} dy < \sum_{n=1}^k \frac{1}{(n+1)^2} \hat{R}_{\xi_{n+1}}^{(\ell)} < \int_1^n \frac{1}{(y+1)^2} \hat{R}_{\xi_{n+1}}^{(\ell)} dy + \frac{M^2}{2} \quad (6.3.56)$$

Taking k as it becomes very large in equation (6.3.56) and using result (6.3.55) yields

$$-\frac{M^2}{2} < \sum_{n=1}^{\infty} \frac{1}{(n+1)^2} \hat{R}_{\xi_{n+1}}^{(\ell)} < M^2 \quad (6.3.57)$$

Hence, by result (6.3.57) the infinite sum of the third term of (6.3.31) is bounded and finite for finite M .

Now considering (6.3.31), equation (6.3.34), equation (6.3.51) and equation (6.3.57) it can be written that

$$C < \sum_{n=1}^{\infty} Y_n^2 < D \quad (6.3.58)$$

where

$$C = x^2 - \frac{[x^2 + 2M^2]^{\frac{1}{2}}}{2} + 1 - M^2 \log \frac{\{[x^2 + 2M^2]^{\frac{1}{2}} + x\}}{\{[x^2 + 2M^2]^{\frac{1}{2}} - x\}} - \frac{M^2}{2} \quad (6.3.59)$$

and

$$D = \frac{3x^2}{2} + \frac{[x^2 + 2M^2]^{\frac{1}{2}}}{2} - 1 + M^2 \log \frac{\{[x^2 + 2M^2]^{\frac{1}{2}} + x\}}{\{[x^2 + 2M^2]^{\frac{1}{2}} - x\}} + \frac{x^2}{2} \left(1 + \frac{M^2}{x^2}\right)^{\frac{1}{2}} + M^2 \quad (6.3.60)$$

Since C and D are constant and finite, then taking

expectations of equation (6.3.58) gives

$$C < E \left\{ \sum_{n=1}^{\infty} Y_n^2 \right\} < D \quad (6.3.61)$$

or considering equation (6.3.29), equation (6.3.61) can be rewritten as

$$C < \sum_{n=1}^{\infty} E \{Y_n^2\} < B \quad (6.3.62)$$

and this satisfies condition (6.2.9) of Dvoretzky's theorem.

Now with Y_n as in (6.3.27)

$$Y_n = \gamma_{n+1} \psi_{n+1}$$

and where

$$\psi_{n+1} = x \left\{ \left(1 + \frac{\hat{R}_{\xi_{n+1}}^{(\ell)}}{x^2} \right)^{\frac{1}{2}} - 1 \right\}$$

as in equation (6.3.28). Hence

$$Y_n = \gamma_{n+1} \left\{ (x^2 + \hat{R}_{\xi_{n+1}}^{(\ell)})^{\frac{1}{2}} - x \right\} \quad (6.3.63)$$

Now for large ℓ and/or n

$$\hat{R}_{\xi_{n+1}}^{(\ell)} \rightarrow 0 \quad (6.3.64)$$

Hence, if expectations of (6.3.63) is taken, the result is

$$E \{Y_n\} = 0 \quad (6.3.65)$$

which satisfied condition (6.2.8) of Dvoretzky's theorem.

Now provided that the condition (6.2.10) on the first

estimate is assured, then all the conditions for Dvoretzky's theorem have been satisfied and therefore, his theorem and results apply. Now it can be said that the algorithm presented in equation (6.3.5) converges with probability 1 and in the mean square sense for all $n=1,2,3,\dots$ such that

$$\lim_{n \rightarrow \infty} E \{ (\hat{W}_n - \theta_w)^2 \} = 0 \quad (6.3.66)$$

and

$$P \{ \lim_{n \rightarrow \infty} \hat{W}_n = \theta_w \} = 1 \quad (6.3.67)$$

where $\theta_w = 0$. Now since algorithm (6.3.5) converges as stated, apply the reverse transformation of (6.3.4), that is,

$$\hat{W}_n = (\hat{x}_n - x) \quad (6.3.68)$$

to algorithm (6.3.5) which gives,

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{ (P_{r_{n+1}}(e))^{1/2} - \hat{x}_n \} \quad n=1,2,3,\dots \quad (6.3.69)$$

This also converges with probability 1 and in the mean square sense such that

$$\lim_{n \rightarrow \infty} E \{ (\hat{x}_n - x)^2 \} = 0 \quad (6.3.70)$$

and

$$P \{ \lim_{n \rightarrow \infty} \hat{x}_n = x \} = 1 \quad (6.3.71)$$

for any finite, zero or non-zero x .

6.4 Summary

In the previous section of this chapter, an alternate method of proof for the algorithm presented was given. Essentially, Dvoretzky's theorem was stated and the algorithm presented here was shown to satisfy all the requisite conditions of the theorem and hence the theorem applies. The theorem, then, is an alternate proof and guarantee that the algorithm is convergent and thus acts as an independent check on the theory put forth in Chapter III.

CHAPTER VII

Conclusions and Future Work

The work of this thesis has, on the whole, been extensively theoretical in terms of proving the convergence of the suggested algorithm. In an attempt to strike a closer balance between abstraction and reality, effort has been expended on indicating the practical benefits of stochastic approximation algorithms as opposed to conventional estimation techniques. Application of these techniques to various facets in electrical engineering has been implied or cited. In particular, the area of learning control has made extensive utilization of the principles involved and applied stochastic approximation techniques to the evaluation of transition and state probabilities for the purpose of cost function prediction. In order to give the reader a proper intuitive feel for the topic and an appreciation for the results of this work on stochastic approximation, the subject has been introduced by giving specific pieces of theory in the historical sequence in which these contributions have been made to the body of knowledge on stochastic approximation. The review of the three major contributions serves to establish a frame-of-reference for this work as well as to set the theme for the work currently being done in the area of stochastic approximation. The more

recent algorithms have also been introduced to establish guides for comparison in terms of convergence and overall performance.

Essentially, then, the concept of sample autocorrelation has been applied to the area of stochastic approximation involved in the area of parameter extraction for noisy environments based on no *a priori* knowledge of process statistics. A pre-requisite before comparison of performances can be made is the establishment of convergence. Convergence was proved for the new algorithm in two ways, each method independent of the other. Chapter III contains a proof of convergence based on the concept that the mean square error of the algorithm vanishes in the limit. Likewise, Chapter VI proves convergence of the same algorithm in the mean square since based on the application of Dvoretzky's theorem. Both methods of proof concur on the convergent nature of the new algorithm.

Having established convergence, evaluation of the relative performance of the three algorithms was made. The vehicle for comparison was the sum of squared error (S.S.E) and the time sum of squared errors (T.S.S.E.) which are the discrete equivalents of the integral of squared errors (I.S.E.) and the time integral of squared errors (T.I.S.E.) in continuous analysis respectively. The simulation results in Chapter V verify the improved performance obtained in terms

of rate of convergence by the new algorithm as compared with the previous two algorithms. This has been achieved only over a limited range of noise conditions, which has been established, but has been done successfully with minimal computation time, using only simple calculations and with minimal computer memory requirements, thereby facilitating applications.

These conclusions suggest a number of areas for future work in the area of stochastic approximation, particularly, to forward the contents of this work. One of the areas is the sensitivity of an optimal Γ -sequence as solved from equations (4.3.10) and (4.3.11). This is no easy task, since a two point nonlinear discrete boundary value problem has to be solved. From such a solution sensitivity with respect to the initial error, E_0 , and/or the first estimate $\hat{x}_0$, has to be determined. Since the problem is not traceable into a closed form, the solutions must be obtained iteratively and then numerical analysis is the only recourse to the sensitivity problem. Along with this study, an analysis on the sensitivity of the rate of convergence to the above Γ -sequence can be evaluated.

The sensitivity study would be of exceptional value in helping develop a strong intuitive feel for the convergence rate of the algorithm. With such a founded feeling, one could attempt a study of a family of closed form solutions

of a sub-optimal Γ -sequence obtained by solving a reduced form of the boundary value problem given in equations (4.3.10) and (4.3.11). This could then be used along with the algorithm to calculate a value of y_n at every instant n as the estimation is proceeding. This would keep the procedure very simple and allow a good sub-optimal convergence to occur.

Another aspect is to make the Γ -sequence adaptive. There are essentially two basic ideas that could be utilized. First, a one-side difference can be estimated and used to regulate the Γ -sequence by some rule. To maintain convergence, restrictions and bounds would have to be developed on the elements of the Γ -sequence. Second, one could make use of any Γ -sequence and change the value of the successive terms only after the difference between the sample information and the latest estimate changes sign. If both ideas were merged, a very effective algorithm could be had with very little extra effort.

In all the work to date, no mention has been made of the requirements or desirability of developing a reasonable and useful bound on the convergence. In some very recent work, Davisson* has developed a technique for developing a probability bound on the convergence of a stochastic approximation algorithm when only a finite number of iterations are taken. It would be interesting to do a study on the

*Davisson, L. D.³⁵ "Probability Bound for Stochastic Approximation" McMaster University, Department of Electrical Engineering, Seminar November 11, 1969.

application of this technique to the algorithms developed in this thesis and those suggested in the above summary.

In conclusion, it can be said that with the suggested introduction of sample autocorrelation to stochastic approximation, an improved algorithm, from the point of view of convergence rate has been developed. It enjoys improved performance over the two previously developed algorithms. It can be applied to any problem that can be formulated as a regression problem with repeated observations. In situations where observation periods are long and when *a priori* statistics are unknown, this and other stochastic approximation algorithms have advantages over conventional estimation procedures. Only short intervals of data need be processed at any one time and then discarded. Only simple computations and very little memory space is required. Thus, stochastic approximation algorithms permit estimation in situations that are prohibitive to conventional estimation techniques.

REFERENCES

- (1) Robbins, H., and S. Monro: "A Stochastic Approximation Method", Ann. Math. Stat., Vol. 22, pg. 400-407, 1951.
- (2) Kiefer, J., and J. Wolfowitz: "Stochastic Estimation of the Maximum of a Regression Function", Ann. Math. Stat., Vol. 23, pg. 462-466, 1952.
- (3) Dvoretzky, A.: "On Stochastic Approximation", Proc. Third Berkeley Symp. Math. Stat. Prob., Vol. 1, pg. 39-55, University of California Press, Los Angeles, 1956.
- (4) Fu, K. S., Nikolic, Z. J., Chien, T. Y., and Wee, W. G.: "On the Stochastic Approximation and Related Learning Techniques", Report TR-EE 66, Purdue University, April, 1966.
- (5) Wolfowitz, J.: "On the Stochastic Approximation Method of Robbins and Monro", Ann. Math. Stat., Vol. 23, pg. 457-461, 1952.
- (6) Blum, J. R.: "Approximation Methods which Converge with Probability One", Ann. Math. Stat., Vol. 25, pg. 382-386, 1954.
- (7) Blum, J. R.: "Multidimensional Stochastic Approximation Methods", Ann. Math. Stat., Vol. 25, pg. 737-744, 1954.
- (8) Chung, K. L.: "On a Stochastic Approximation Method", Ann. Math. Stat., Vol. 25, pg. 463-483, 1954.
- (9) Derman, C.: "An Application of Chung's lemma to the Kiefer-Wolfowitz Stochastic Approximation Procedure", Ann. Math. Stat., Vol. 27, pg. 529, 1956.
- (10) Burkholder, D. L.: "On a Class of Stochastic Approximation Processes", Ann. Math. Stat., Vol. 27, pg. 1044-1059, 1956.
- (11) Sacks, J.: "Asymptotic Distribution of Stochastic Approximation Procedures", Ann. Math. Stat., Vol. 29, pg. 373-405, 1958.

- (12) Dupac, V.: Casopis Pest. Mat., Vol. 82, pg. 47-75, (1957) (in Czechoslovakian). English translation: Massachusetts Institute of Technology, Lincoln Lab., Tech. Report No. 22G-0008, Feb. 21, 1960.
- (13) Sakrison, D. J.: "A Continuous Kiefer-Wolfowitz Procedure for Random Processes", Ann. Math. Stat., Vol. 35, pg. 590-599, 1964.
- (14) Driml, M., and Nedoma, J.: Trans. Second Prague Conference on Information Theory, Statistical Decision Functions and Random Processes, Czech. Academy of Science, Prague, pg. 145, 1960.
- (15) Alberts, A., and Gardner, L.: "Stochastic Approximation and Nonlinear Regression", M.I.T. Press, Cambridge, Massachusetts.
- (16) Sakrison, D. J.: "Efficient Recursive Estimation; Application to Estimating the Parameters of a Covariance Function", International Journal of Engineering Science, Vol. 3, pg. 461-483, 1965.
- (17) Sakrison, D. J.: "Efficient Recursive Estimation of the Parameters of a radar or radio Astronomy Target", I.E.E.E. Trans. Information Theory IT-12, No. 1, 35, 1966.
- (18) Sakrison, D. J.: "Application of Stochastic Approximation Methods to System Optimization", Mass. Inst. of Technology., Research Lab. of Electronics. Tech. Report, No. 391, 1962.
- (19) Sakrison, D. J.: "Iterative Design of Optimum Filters for Non-mean-square-error Performance Criteria", IRE Tran. Inf. Theory IT-9, No. 3, pg. 161-167, 1963.
- (20) Kushner, H.: "Adaptive Techniques for the Optimization of Binary Detection Systems", I.E.E.E. International Convention record, vol. 11, pt. 4, pg. 107-117, 1963.
- (21) Sklansky, J.: "Learning Systems for Automatic Control" I.E.E.E. Trans. Automatic Control, AC-11, Jan., 1966.
- (22) Nikolic, Z. J., and Fu, K. S.: "On Some Reinforcement Techniques and their Relation to the Stochastic Approximation", I.E.E.E. Trans. Auto. Control, pg. 756-758, October, 1966.

- (23) Waltz, M. D., and Fu, K. S.: "A Heuristic Approach to Reinforcement Learning in Control Systems", I.E.E.E. Trans. on Auto. Control, AC-10 No. 4, pg. 390-398, October, 1965.
- (24) Fu, K. S.: "Learning System Theory", From "System Theory", edited by Zadek, L. A., and Polak, E., pg. 440, McGraw Hill, 1969.
- (25) McMurty, G. J., and Fu, K. S.: "A Variable Structure Automation used as a Multi-model Searching Technique", Prox. National Electronic Conference, Vol. 21, 1965.
- (26) McLaren, R. W., and Fu, K. S.: "An Application of Stochastic Automata to the Synthesis of Learning Systems", Tech. Report TR-EE 65-17, School of Electrical Engineering, Purdue University, Lafayette, Ind., September, 1965.
- (27) Rozonoer, L. I.: "L.S. Pontryagin's Maximum Principle in Optimal Control Theory", Automation and Remote Control, Vol. 1, No. 20, October, November, December, 1959.
- (28) Pontryagin, L. S., et al: "The Mathematical Theory of Optimal Processes", International Publishers, New York, N. Y., 1962.
- (29) Fan, L. T., and Wang, C. S.: "The Discrete Maximum Principle", John Wiley & Sons, New York, N. Y., 1964.
- (30) Jordan, B. W., and Polak, E.: "Theory of a Class of Discrete Optimal Control Systems", J. Electronics and Control, Vol. 17, December, 1964.
- (31) Katz, S.: "A Discrete Version of Pontryagin's Maximum Principle", J. Electronics and Control, Vol. 13, pg. 179, 1962.
- (32) Golomb, S. W.: "Sequences with Randomness Properties", Glen L. Martin Co., Baltimore, Md., June, 1955.
- (33) Korn, G. A.: "Random Processes: Simulation and Measurement", McGraw Hill, 1966.
- (34) Hamming, R. W.: "Numerical Methods for Scientists and Engineers", McGraw Hill, N. Y., pg. 34 and 389, 1962.

- (35) Davisson, L. D.: "Convergence Probability Bounds for Stochastic Approximation", correspondence (to be published).

LIST OF SYMBOLS

Roman

a_n	a sequence of non-negative real numbers
A_i	scalar factors of a vector in a hyperspace
B	positive constant
C	arbitrarily large finite number
c_n	an element of a difference sequence; non-negative real
D	positive constant
e_i	unit vector in a hyperspace
$E\{\cdot\}$	expectation operator
E_i	mean difference error $\hat{x}_i - x$
$\underline{f}(\dots\dots\dots)$	a vector state function
$f(\dots)$	a bivariable scalar random function
$g(\dots\dots\dots)$	measurable random multiparameter function
H	Hamiltonian
J	cost function
i, j, k	integer subscripts
K	a positive constant
l	discrete argument of autocorrelation function
L	arbitrarily large finite number
m	integer subscript
$m(\cdot)$	a scalar function of a scalar variable
m_0	a nominal or arbitrary scalar value
M	a positive constant

M	discrete equivalent of the scalar Pontryagin function
n	integer subscript
p, q, k	range limits on subscripts
r_n	samples taken from the distribution of x
R	scalar constant
$\hat{R}_{r_n}(\ell)$	discrete sample autocorrelation function
S_i	control situation
t	time
$T_n(\dots, \dots, \dots)$	a measurable transformation
T	sampling interval
u_i	control variable
v_n^2	expected mean square error
$\hat{W}_n$	a transform space defined by $W_n = (\hat{x}_n - x)$
x	true parameter being sought
$\hat{x}_n$	stochastic approximation generated by the stochastic approximation algorithms
$\hat{X}_n$	estimate of θ in Dvoretzky's theorem
y	dummy integration variable
Y_n	a random variable
Z_i	random entity or element of a stochastic process

Greek

α	scalar parameter
α_k	estimate of θ in Kiefer - Wolfowitz method
α_m	estimate of θ in Robbins - Monro Technique
α_n	non-negative measurable function - Dvoretzky's theorem
β_n	non-negative measurable function - Dvoretzky's theorem
β	positive constant
γ_n	non-negative measurable function - Dvoretzky's theorem
δ, ϵ	arbitrarily small positive scalar members
$\Pi(\cdot)$	a small positive scalar function
λ	Lagrangian multiplier
η	a scalar perturbation
θ	a true scalar value
θ_n	transform of $(\hat{x}_n - x)$
$\theta_{t_{k_f}}$	terminal time cost term
θ_w	zero value sought by transformed algorithm
ϕ	integrand of a cost function
ρ	scalar constant
ρ_n	a scalar variable
ξ_n	the random component of r_n
δ	instantaneous value of a cost function
$\wedge$	symbol for estimated variable
$-$	a vector designation
$\langle \cdot, \cdot \rangle$	denotes an inner product

APPENDIX A

Relationship between Reinforced Learning and Stochastic Approximation

To begin with, an attempt will be made to give a physical and intuitive feel for reinforced heuristic learning. It is felt that such an introduction is necessary before giving a more mathematical formulation of heuristic learning control. Once the formulation has been introduced, the relationship to stochastic approximation will be developed.

Consider what is often referred to in psychology as the T-maze rat experiment. A simple T-maze is constructed, as in Figure A-1 and a rodent such as a mouse or rat is placed at the starting point A.

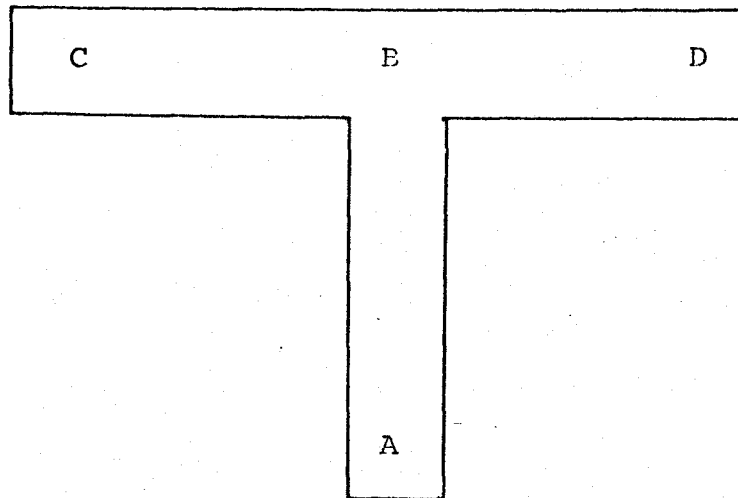


Figure A-1: T-maze Rat Experiment

The maze is enclosed along the complete perimeter as outlined. At the beginning of each experiment food (cheese or grain) is placed in either one or the other arms of the maze at C or D. Say that the food is placed at C and the rat is let go. When the rat reaches point B, it must make a choice as to which way it will go. Call this a situation s_i or a "control situation", that is, a point in the rat's control sequence where a decision has to be made between alternatives in direction. Now assuming the rat is not influenced by olfactory or visual stimuli, chances are even that it will select one or other direction without any preference the first time. If a probability is assigned to indicate this, say P_{ij} is the probability that in situation i a choice will be made to go to j , then there will be two such decisions or transition probabilities and each will be equal to one half. If the rat goes to D, it finds nothing and hence, it was negatively reinforced. In terms of probabilities, it means that P_{bd} was reduced by δ and P_{bc} was increased by δ to keep the sum equal to one. If when the control situation S_b appears again and the rat remembers the previous experience it had, then it is more likely to try the direction C. This is, in fact, what the new transition probabilities say. Now, if in situation S_b the first time the rat had gone to C, then it would have been positively reinforced with the reward of food. In terms of probabilities P_{bc} would have been increased by δ and P_{bd}

reduced by δ to keep the sum unity.

Now if the food is consistently placed at C, the rat will learn to go there when released from A. Similarly, the transition probability P_{bc} will approach 1 and P_{bd} will approach 0 if the experiment were repeated often. Then if one selects the direction or control choice corresponding to the highest transition probability, reinforced learning has been achieved.

Consider a plant described by the differential equation of the form

$$\dot{\underline{x}} = f(\underline{x}, u, \underline{V}, \underline{N}, t) \quad (A-1)$$

where $\underline{x}$ is a state vector defined in state space Ω_x ,

u is the control or control choice,

$\underline{V}$ is the environment vector defined in Ω_x space,

$\underline{N}$ is the output disturbance vector including output measurement noise also defined in Ω_x ,

and t is time.

Now define a measurement vector $\underline{M}$ in space Ω_M as

$$\underline{M}^T = (\underline{x}^T \vdots \underline{V}^T) \quad (A-2)$$

which is obtained from the plant and environment every T seconds, the sampling rate.

The controller learns heuristically to drive the state vector $\underline{x}$ from any set of initial conditions to the neighbourhood of the origin in state space in such a way as to

optimize a pre-defined index of performance IP. The learning is accomplished by building a stimulus-response mapping between element in the space Ω_M and Ω_u . To begin with, the space Ω_M is partitioned into convenient class for practical reasons and the best choice from the set Ω_u for a given class in Ω_M is considered the same for all members of that class. The problem then is to develop the relationship that selects the best control choice for a given control situation.

Now for a control situation S_i , it is not a deterministic problem to select the best choice of control from Ω_u since the control choice at time nT is dependent on the system state at time $(n+1)T$. Let P_{ij} be the probability that the j^{th} element of Ω_u is the best control choice for the situation S_j .

Initially, all $P_{ij} = 1/k$, $j = 1, \dots, k$; $i = 1, \dots, p$, if no prior knowledge is given or assumed. As experiments are carried out, learning proceeds if a transition or subjective probability P_{ij} approaches 1 for a pair u_j, S_i . If the probability P_{ij} exceeds a preset threshold T_p near 0 or 1 then learning is complete and the mapping between Ω_u and Ω_m for a pair u_j, S_i is complete. Otherwise, learning continues and for a given S_i the u_j corresponding to the maximum P_{ij} is used.

Now the probabilities P_{ij} are adjusted ie. reinforced

positively or reinforced negatively based on the evaluation of the IP. The idea being that if a P_{ij} results in a choice of u_j that helps minimize the IP, it is added to, but if it hampers minimization, it is penalized. For this, two learning operators, L_+ and L_- , a positive and negative reinforcing operator respectively, are defined.

For positive reinforcement

$$\begin{aligned} P_{ij} [(n+1)T] &= L_+ \{P_{ij}[nT]\} \\ &= \theta P_{ij}[nT] + (1-\theta) \text{ for } 0 < \theta < 1 \text{ (A-3)} \end{aligned}$$

is used to correspond to a given u_j . At the same time, the negative reinforcement operator is applied to all other choices of u_q , $q = 1, \dots, k$; $q \neq j$.

$$\begin{aligned} P_{ij} [(n+1)T] &= L_- \{P_{ij}[nT]\} \\ &= \theta P_{ij}[nT] \text{ for } 0 < \theta < 1 \text{ (A-4)} \end{aligned}$$

where θ is the learning parameter. A large θ results in slow learning rate because the probabilities P_{ij} change slowly. The converse is also true.

If one begins with some *a priori* knowledge and/or has an intuitive idea of the direction of control based on experience, then the initial distribution of P_{ij} need not necessarily be uniform but could be adjusted to conform with this knowledge. In addition, a subgoal may be introduced

which helps direct the choice of control. It can take a number of forms, one being a cost function made up of estimates of expectations of P_{ij} or combinations thereof.

It still remains to determine the conditions for convergence of (A-3) and (A-4), and their relationship to stochastic approximation. The basic problem still remains, select an optimal u^* from the set of admissible actions,

$$\{u_j; j = 1, 2, \dots, k\} \equiv \Omega_u \quad (A-5)$$

such that an index of performance of the form

$$E\{\zeta | \underline{M}, u_j^*\} = \min_{u_j} [E\{\zeta | \underline{M}, u_j\} : u_j \in \Omega_u] \quad (A-6)$$

is minimized

where $E\{\zeta | \underline{M}, u_j\}$ is the performance index for the action $u_j \in \Omega_u$ applied after the observation $\underline{M} \in \Omega_M$.

$\underline{M}$ is the observed response of the plant

and ζ is an instantaneous performance index evaluation of the action u_j following an observation $\underline{M}$. It is a random variable dependent on the definition of a subgoal

$$\zeta = g(\underline{M}, u_j, \underline{M}^1) \quad (A-7)$$

where $\underline{M}^1 \in \Omega_M$ is the response of the plant due to the action u_j applied after $\underline{M}$ was observed.

The control law is then specified as

$$P\{u^*|\underline{M}\} = 1 \quad (A-8)$$

that is, given all the observations, the probability is 1 that the optimal control will be found. The reinforcement algorithm defined by equations (A-3) and (A-4) is such that following the $(n+1)^{th}$ observation of $\underline{M}$, if the estimate of the IP is such that

$$E_{n_i}\{\tau|\underline{M}, u_j\} = \min_{u_j} [E_{n_j}\{\tau|\underline{M}, u_j\} : u_j \in \Omega_u] \quad (A-9)$$

then the transition probability corresponding to u_i is positively reinforced according to equation (A-3) and all others are negatively reinforced according to equation (A-4). Now equations (A-3), (A-4) and (A-9) can be rewritten as,

$$P_{ij}(n+1) = \theta_{n+1} P_{ij}(n) + (1-\theta_{n+1}) Z_{n+1}(\underline{M}, u_\ell) \quad (A-10)$$

for every control choice $u_\ell \in \Omega_u$ where $0 < \theta_n < 1$ for $n=1, 2, \dots$, and,

$$Z_{n+1}(\underline{M}, u_\ell) = \begin{cases} 1 & \text{if } E_{n_\ell}\{\tau|\underline{M}, u_\ell\} = \min_{u_j} [E_{n_j}\{\tau|\underline{M}, u_j\} : u_j \in \Omega_u] \\ 0 & \text{if } E_{n_\ell}\{\tau|\underline{M}, u_\ell\} \neq \min_{u_j} [E_{n_j}\{\tau|\underline{M}, u_j\} : u_j \in \Omega_u] \end{cases} \quad (A-11)$$

Now after $(n+1)$ performance evaluations of the choices of u following the observations $\underline{M}$, let

$$P [E_{n_i}\{\tau|\underline{M}, u_i\} = \min_{u_j} [E_{n_j}\{\tau|\underline{M}, u_j\} : u_j \in \Omega_u]] = q_i(n+1) \quad (A-12)$$

Then from equations (A-10), (A-11) and (A-12) for u_i ,

$$E \{P_{ij}(n+1) | P_{ij}(n)\} = \theta P_{ij}(n) + (1-\theta)q_i(n+1) \quad (A-13)$$

Theorem* 1

The sufficient conditions for (A-10) and (A-11) to yield

$$P \{ \lim_{n \rightarrow \infty} P_{ij}(n) |_{u=u^*} = 1 \} = 1 \quad (A-14)$$

are:

(1) for every action $u_i \in \Omega_u$ which is not optimal according to (A-6)

$$P \{ \lim_{n \rightarrow \infty} q_i(n) = 0 \} = 1 \quad (A-15)$$

(2) for every $u_l \in \Omega_u$

$$P_{il}(0) > 0 \text{ with } \sum_{i=1}^m P_{il}(0) = 1 \quad (A-16)$$

and (3) for every n and u_i defined by (A-9)

$$P_{ij}(n) = \max_{u_j} \{P_{ij}(n) : u_i \in \Omega_u\} \quad (A-17)$$

Proof: Suppose (A-15) does not hold for at least one sub-optimal control choice $u_i \in \Omega_u$. The expression on the right hand side of (A-13) is positive then with probability 1. This contradicts the assumption that $u_i \neq u^*$ as defined by (A-6). The necessity of (A-15) is thus shown.

Condition (A-15) can be restated thus: given any positive numbers ϵ and δ , an integer K exists such that,

$$P \{ \sup_{n > K} \{q_i(n)\} < \epsilon \} > 1 - \delta \quad (A-18)$$

*From the work of Nikolic and Fu²²

In that case for any $\varepsilon_1 > 0$ and $\delta_1 > 0$ ($0 < \varepsilon_1 < 1$ and $0 < \delta_1 < 1$) an integer K_1 , can be found such that

$$P \{ \inf_{n > K_1} \{ q^*(n) \} > 1 - \varepsilon_1 \} > 1 - \delta_1 \quad (A-19)$$

for the optimal control u^* . Since the sequence

$\{ \max_{u_j} | P_{ij}(n) : u_j \in \Omega_u | \}$ is a monotonically increasing sequence and converges to 1 according to (A-10) and (A-17), for any $\delta_2 > 0$ an integer $K_2 > K_1$ can be found such that

$$P \{ \inf_{n > K_2} \{ P_{ij}^{(n)} |_{u=u^*} \} > 1 - \delta_2 \} > 1 - \varepsilon_1 \quad (A-20)$$

with ε_1 given in (A-19). The condition (A-16) is sufficient to assure $P_{ij}^{(n)} > 0$ for every $n < \infty$ and $u_j \in \Omega_u$.

Now if account is taken of the probability of the system being in a state i at time n , written as $P_n^{(i)}$, then the algorithm can be written as

$$P_{n+1}^{(i)} = \alpha_{n+1} P_n^{(i)} + (1 - \alpha_{n+1}) \lambda_n(i) \quad (A-21)$$

where $\alpha_n = 1 - \theta_n$

and $\lambda_n(i)$

being the ratio of the number of times the system has been in state i out of all n states it has been in that is

$$\lambda_n(i) = n_i / n \quad (A-22)$$

in fact, $\lambda_n(i)$ is directly proportional to $P_{ij}(n)$.

Now let

*From the work of Nikolic and Fu22

$$0 < \alpha_n < 1, \quad \prod_{n=1}^{\infty} \alpha_n = 0 \quad (\text{A-23})$$

and

$$\sum_{n=1}^{\infty} (1-\alpha_n)^2 < \infty \quad (\text{A-24})$$

If equation (A-21) is rewritten after subtracting $P(i)$ from both sides and rearranging gives

$$P_{n+1}(i) - P(i) = \alpha_{n+1} P_n(i) - p(i) + (1-\alpha_{n+1}) \eta_n(i) \quad (\text{A-25})$$

where $P(i) = \lim_{n \rightarrow \infty} P_n(i)$

$$\eta_n(i) = P_n(i) - P(i)$$

$$E \{ \eta_n(i) \} = 0$$

and $E \{ [\eta_n(i)]^2 \} < 1$

Hence, equation (A-25) is in the form of a stochastic approximation algorithm of the Dvoretzky type as outlined in Section (D-8) of Appendix D with

$$T_n[P_n(i)] = \alpha_{n+1}[P_n(i) - p(i)] \quad (\text{A-26})$$

Consequently,

$$P \{ \lim_{n \rightarrow \infty} P_n(i) = p(i) \} = 1 \quad (\text{A-27})$$

by Dvoretzky's theorem.

In addition to the two learning algorithms shown to be of a stochastic approximation nature, other forms of learning can also be interpreted in a similar fashion. For example,

Fu²⁴ shows that the technique of Bayesian inference as used in the learning sense is an algorithm of the Dvoretzky form. Two other works that Fu participated in also deal with similar areas. With M^CMurtry²⁵, learning and stochastic approximation techniques were applied to a multi-model searching technique. With M^CLearen²⁶, they studied an application of stochastic automata to the synthesis of learning systems.

Essentially, then, it can be said that stochastic approximation is a unifying theory for interpreting the various and diverse facets of learning control theory.

APPENDIX B

Proof of Two Stochastic Approximation Algorithms*

For each of Fu's two algorithms, a proof of convergence will be given and a derivation of the inherent optimum properties will also be given.

Consider an algorithm of the form

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{r_{n+1} - \hat{x}_n\} \quad n=1,2,\dots \quad (B-1)$$

where $\hat{x}_n$ is the n^{th} estimate of the true value sought

x is the true value sought

r_n is a sample from a normally distributed random variable

and γ_n is a gain sequence

Here r_n is a linear combination of noise and parameter being sought, that is,

$$r_n = x + \xi_n \quad (B-2)$$

where ξ_n is an element of normal zero mean white noise of variance σ^2 . The form of the gain sequence γ_n is restricted, that is,

$$\gamma_n = \frac{1}{n+K} \quad \text{for } n=1,2,\dots \quad (B-3)$$

where K is an arbitrary non-negative constant. For this definition of γ_n , the following conditions hold true:

$$\gamma_n > 0 \quad (B-4)$$

$$\lim_{n \rightarrow \infty} \gamma_n \rightarrow 0 \quad (B-5)$$

$$\sum_{n=1}^{\infty} \gamma_n = \infty \quad (B-6)$$

$$\sum_{n=1}^{\infty} \gamma_n^2 < \infty \quad (B-7)$$

These all fit nicely into the statement of Dvoretzky's theorem. Now combining equations (B-1) and (B-2) gives

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \{ (x + \xi_{n+1}) - \hat{x}_n \} \quad n=1, 2, \dots \quad (B-8)$$

Now subtracting x from both sides gives,

$$(\hat{x}_{n+1} - x) = (1 - \gamma_{n+1}) \hat{x}_n + \gamma_{n+1} (x + \xi_{n+1}) - x$$

and after recombining and factoring gives,

$$(\hat{x}_{n+1} - x) = (1 - \gamma_{n+1}) (\hat{x}_n - x) + \gamma_{n+1} \xi_{n+1} \quad (B-9)$$

Squaring both sides of equation (B-9) gives

$$\begin{aligned} (\hat{x}_{n+1} - x)^2 &= (1 - \gamma_{n+1})^2 (\hat{x}_n - x)^2 + \gamma_{n+1}^2 \xi_{n+1}^2 \\ &\quad + 2(1 - \gamma_{n+1}) \gamma_{n+1} \xi_{n+1} (\hat{x}_n - x) \end{aligned} \quad (B-10)$$

Take expectations of both sides yields

$$E \{ (\hat{x}_{n+1} - x)^2 \} = (1 - \gamma_{n+1})^2 E \{ (\hat{x}_n - x)^2 \} + \gamma_{n+1}^2 E \{ \xi_{n+1}^2 \} \quad (B-11)$$

since

$$E \{ 2(1-\gamma_{n+1})\gamma_{n+1}\xi_{n+1}(\hat{x}_n - x) \} = 2(1-\gamma_{n+1})\gamma_{n+1}E\{\xi_{n+1}\}E\{\hat{x}_n - x\} \quad (\text{B-12})$$

because

$$\hat{x}_n - x \text{ and } \xi_{n+1}$$

are independent and since

$$E \{ \xi_{n+1} \} = 0 \quad (\text{B-13})$$

then equation (B-11) holds true.

Define

$$E \{ (\hat{x}_{n+1} - x)^2 \} = V_{n+1}^2 \quad (\text{B-14})$$

Then equation (B-11) can be written as

$$V_{n+1}^2 = (1-\gamma_{n+1})^2 V_n^2 + \gamma_{n+1}^2 \sigma^2 \quad (\text{B-15})$$

where for a stationary process,

$$\sigma_{n+1}^2 = E \{ \xi_{n+1}^2 \} = \sigma^2 \text{ for all } n=1,2,\dots \quad (\text{B-16})$$

Now define a transformation

$$T_n(\hat{x}_1, \dots, \hat{x}_n) = (1-\gamma_{n+1})\hat{x}_n + \gamma_{n+1} \quad (\text{B-17})$$

Then equation (B-8) can be written as

$$\hat{x}_{n+1} = T_n(\hat{x}_1, \dots, \hat{x}_n) + \gamma_{n+1}\xi_{n+1} \quad n=1,2,\dots \quad (\text{B-18})$$

Now

$$|T_n(\hat{x}_1, \dots, \hat{x}_n) - x| = |(1 - \gamma_{n+1})\hat{x}_n + \gamma_{n+1}x - x| \quad (B-19)$$

or

$$|T_n(\hat{x}_1, \dots, \hat{x}_n) - x| = (1 - \gamma_{n+1}) |(\hat{x}_n - x)| \quad (B-20)$$

and by defining

$$F_{n+1} \equiv (1 - \gamma_{n+1}) \quad (B-21)$$

$$|T_n(\hat{x}_1, \dots, \hat{x}_n) - x| = F_{n+1} |(\hat{x}_n - x)| \quad (B-22)$$

Noting definition (B-21) and form (B-3) it can be shown that

$$\prod_{n=1}^{\infty} F_n = 0 \quad (B-23)$$

In view of the transformation defined in equation (B-17), then identifying the remaining element of (B-18) with the form (D-2.5) of Dvoretzky's theorem gives

$$Y_n = \gamma_{n+1} \xi_{n+1} \quad (B-24)$$

Hence since

$$E\{Y_n\} = E\{\gamma_{n+1} \xi_{n+1}\} = 0 \quad (B-25)$$

Then condition (D-2.7) is satisfied.

Now consider

$$\sum_{n=1}^{\infty} E\{Y_n^2\} = \sum_{n=1}^{\infty} E\{\gamma_{n+1}^2 \xi_{n+1}^2\} \quad (B-26)$$

which gives

$$\sum_{n=1}^{\infty} E\{Y_n^2\} = \sigma^2 \sum_{n=1}^{\infty} \gamma_n^2 \quad (B-27)$$

Recalling equation (B-7) and substituting into (B-27) gives

$$\sum_{n=1}^{\infty} E\{Y_n^2\} < \infty \quad (B-28)$$

Hence, Dvoretzky's theorem is satisfied and thus

$$\lim_{n \rightarrow \infty} E\{(\hat{x}_{n+1} - x)^2\} = 0 \quad (B-29)$$

and

$$P\{\lim_{n \rightarrow \infty} \hat{x}_{n+1} = x\} = 1 \quad (B-30)$$

In addition to the convergence property just proven using Dvoretzky's theorem, algorithm one has an optimum gamma sequence. To determine this sequence, which makes the mean square error as small as possible at each step, set all derivatives of V_n^2 with respect to γ_n equal to zero and solve for γ_n .

$$\frac{dV_n^2}{d\gamma_n} = -2(1-\gamma_n)V_{n-1}^2 + 2\gamma_n\sigma^2 \quad (B-31)$$

Equating the right hand side of (B-31) to zero gives

$$V_{n-1}^2 = \gamma_n V_{n-1}^2 + \gamma_n \sigma^2 \quad (B-32)$$

and solving for γ_n gives

$$\gamma_n = \frac{V_{n-1}^2}{V_{n-1}^2 + \sigma^2} \quad (\text{B-33})$$

Now iterating a form of equation (B-15)

$$V_n^2 = (1-\gamma_n)^2 V_{n-1}^2 + \gamma_n^2 \sigma^2 \quad (\text{B-34})$$

with equation (B-33) to obtain the optimum V_n^{2*} and γ_n^* .

To begin, let the initial expected mean square error be

$$V_o^2 = E\{(\hat{x}_o - x)^2\} \quad (\text{B-35})$$

Iterating V_n^2 and γ_n alternately, gives

$$V_1^2 = (1-\gamma_1)^2 V_o^2 + \gamma_1^2 \sigma^2 \quad (\text{B-36})$$

and from (B-33)

$$\gamma_1 = \frac{V_o^2}{V_o^2 + \sigma^2} \quad (\text{B-37})$$

or

$$\gamma_1 = \frac{1}{1+K} \quad (\text{B-38})$$

where

$$K = \frac{\sigma^2}{V_o^2} \quad (\text{B-39})$$

substituting from equation (B-38) into (B-36) for γ_1 gives
after slight simplification

$$V_1^2 = \frac{\sigma^2}{1+K} \quad (\text{B-40})$$

Now

$$\gamma_2 = \frac{V_1^2}{V_1^2 + \sigma^2} \quad (\text{B-41})$$

Substituting into (B-41) for V_1^2 from (B-40) and rearranging terms gives

$$\gamma_2 = \frac{1}{2+K} \quad (\text{B-42})$$

Now

$$V_2^2 = (1-\gamma_2)^2 V_1^2 + \gamma_2^2 \sigma^2 \quad (\text{B-43})$$

Substituting from equation (B-40) and (B-42) for V_1^2 and γ_2 respectively, gives, after some manipulation,

$$V_2^2 = \frac{\sigma^2}{2+K} \quad (\text{B-44})$$

Now assume V_n^2 is known, then by equations (B-33) and (B-32)

$$\gamma_{n+1} = \frac{V_n^2}{V_n^2 + \sigma^2} \quad (\text{B-45})$$

and

$$V_{n+1}^2 = (1-\gamma_{n+1})^2 V_n^2 + \gamma_{n+1}^2 \sigma^2 \quad (\text{B-46})$$

Substituting for γ_{n+1} gives

$$V_{n+1}^2 = \frac{\sigma^2 V_n^2}{V_n^2 + \sigma^2} \quad (\text{B-47})$$

using (B-46) and (B-47) are general terms for γ and V^2 sequence can be extended by simple substitutions. Hence

$$\gamma_3 = \frac{1}{3+K}$$

and

$$V_3^2 = \frac{\sigma^2}{3+K}$$

and so on.

Hence, by inspection, it can be seen that the optimum gamma sequence is given as

$$\gamma_n^* = \frac{1}{n+K} \quad (\text{B-48})$$

with the corresponding optimum expected mean square error

$$V_n^{2*} = \frac{\sigma^2}{n+K} \quad (\text{B-49})$$

where K is defined as in equation (B-39).

Now reconsider equation (B-15) but replacing the variance σ^2 with an upper bound B, that is,

$$\sigma^2 \leq B \quad (\text{B-50})$$

Hence

$$V_n^2 \leq F_n^2 V_{n-1}^2 + \gamma_n^2 B \quad (\text{B-51})$$

Now iterating this back to V_0^2 gives

$$\begin{aligned} V_n^2 &\leq F_n^2 F_{n-1}^2 \dots F_1^2 V_0^2 + F_n^2 F_{n-1}^2 \dots F_2^2 \gamma_1^2 B + \dots \\ &\quad + F_n^2 F_{n-1}^2 \dots F_{m+1}^2 \gamma_m^2 B + \dots \\ &\quad + F_n^2 \gamma_{n-1}^2 B + \gamma_n^2 B \quad \text{for } n \geq m \geq 1 \end{aligned} \quad (\text{B-52})$$

If it is assumed that $\gamma_n = 1/n$, then (B-52) can be rewritten as

$$V_n^2 \leq B \left[\frac{1}{n^2} + \left[\frac{2-1}{n} \right]^2 + \dots + \left[\frac{m}{n} \right]^2 \frac{1}{m^2} + \dots + \left[\frac{n-1}{n} \right]^2 \frac{1}{(n-1)^2} + \frac{1}{n^2} \right] \quad (\text{B-53})$$

which reduces to

$$V_n^2 \leq \frac{B}{n} \quad (\text{B-54})$$

Now instead of as in algorithm one given in equation (B-1), the sample mean

$$\frac{1}{n} \sum_{i=1}^n r_i$$

is used as sample information instead of only r_n , then the resulting algorithm is

$$\hat{x}_{n+1} = \hat{x}_n + \gamma_{n+1} \left\{ \frac{1}{n+1} \sum_{i=1}^{n+1} r_i - \hat{x}_n \right\} \quad n=1,2,\dots \quad (\text{B-55})$$

For the moment, consider the random element of noise ξ_n as given in equation (B-2). Let

$$v_n = \sum_{i=1}^n \xi_i \quad (\text{B-56})$$

Then

$$v_n^2 = \left[\frac{1}{n} \sum_{i=1}^n \xi_i \right]^2 \quad (\text{B-57})$$

which can be expanded as

$$v_n^2 = \frac{1}{n} \left[\frac{1}{n} \sum_{i=1}^n \xi_i^2 \right] + \frac{2}{n^2} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \xi_i \xi_j \quad (\text{B-58})$$

or alternately

$$v_n^2 = \frac{1}{n} \left[\frac{1}{n} \sum_{i=1}^n \xi_i^2 \right] + 2 \left[\frac{1}{n} \sum_{i=1}^n \xi_i \right] \left[\frac{1}{n} \sum_{\substack{j=1 \\ j \neq i}}^n \xi_j \right] \quad (\text{B-59})$$

Since

$$\sigma_n^2 = \frac{1}{n} \sum_{i=1}^n \xi_i^2 \quad \text{or} \quad \sigma_n^2 = E\{\lambda_n^2\} \quad (\text{B-60})$$

and for large n ,

$$\sum_{i=1}^n \xi_i \rightarrow 0 \quad \text{or} \quad E\{\xi_i\} = 0 \quad (\text{B-61})$$

Using the first two conditions of (B-60) and (B-61), and substituting into equation (B-59) gives

$$v_n^2 = \frac{1}{n} \sigma_n^2 \quad (\text{B-62})$$

Now taking expectations of this equation gives

$$E\{v_n^2\} = \frac{1}{n} E\{\sigma_n^2\} \quad (\text{B-63})$$

Since the random process is assumed to be stationary

$$E\{\sigma_n^2\} = \sigma^2,$$

and hence

$$E\{v_n^2\} = \frac{1}{n} \sigma^2 \quad (\text{B-64})$$

If into equation (B-55) a substitution is made for r_n from equation (B-2), the result after rearrangement is

$$\hat{x}_{n+1} = (1-\gamma_{n+1}) \hat{x}_n + \gamma_{n+1} x + \gamma_{n+1} v_{n+1} \quad n=1,2,\dots \quad (\text{B-65})$$

Subtracting x from both sides yields

$$(\hat{x}_{n+1} - x) = (1 - \gamma_{n+1})(\hat{x}_n - x) + \gamma_{n+1}v_{n+1}, \quad (B-66)$$

and squaring results in

$$\begin{aligned} (\hat{x}_{n+1} - x)^2 &= (1 - \gamma_{n+1})^2 (\hat{x}_n - x)^2 + 2\gamma_{n+1}v_{n+1}(1 - \gamma_{n+1})(\hat{x}_n - x) \\ &\quad + \gamma_{n+1}^2 v_{n+1}^2 \end{aligned} \quad (B-67)$$

Since the following

$$E\{2\gamma_{n+1}\gamma_{n+1}(1 - \gamma_{n+1})(\hat{x}_n - x)\} = 2\gamma_{n+1}(1 - \gamma_{n+1})E\{\hat{x}_n - x\}E\{v_n\} \quad (B-68)$$

vanish because $E\{v_n\} = 0$ the expectation of equation (B-67) gives

$$E\{(\hat{x}_{n+1} - x)^2\} = (1 - \gamma_{n+1})^2 E\{(\hat{x}_n - x)^2\} + \gamma_{n+1}^2 \frac{\sigma^2}{n} \quad (B-69)$$

or using the definition (B-14), equation (B-69) becomes

$$V_{n+1}^2 = (1 - \gamma_{n+1})^2 V_n^2 + \gamma_{n+1}^2 \frac{\sigma^2}{n} \quad (B-70)$$

Now if the upper bound on σ^2 is used, namely B , then

$$V_{n+1}^2 \leq (1 - \gamma_{n+1})^2 V_n^2 + \gamma_{n+1}^2 \frac{B}{n} \quad (B-71)$$

Iterating the above inequality back to V^2 gives

$$\begin{aligned} V_n^2 &\leq F_n^2 \dots F_1^2 V_0^2 + F_n^2 \dots F_3^2 F_2^2 \gamma_1^2 \frac{B}{1} + \dots + F_n^2 \dots F_{n+2}^2 F_{m+1}^2 \gamma_m^2 \frac{2B}{m} + \dots \\ &\quad + F_n^2 \gamma_{n-1}^2 \frac{2B}{n-1} + \gamma_n^2 \frac{2B}{n} \end{aligned} \quad (B-72)$$

Assuming as before that $\gamma_n = 1/n$ then (B-72) can be written as

$$V_n^2 \leq B \left\{ \frac{1}{n^2} + \left[\frac{3-1}{n} \right]^2 \frac{1}{2^2} + \dots + \left[\frac{m}{n} \right]^2 \frac{1}{m^2} \frac{1}{m} + \dots \right. \\ \left. + \left[\frac{n-1}{n} \right]^2 \frac{1}{(n-1)^2} \frac{1}{n+1} + \frac{1}{n^2} \frac{1}{n} \right\} \quad (B-73)$$

This reduces to

$$V_n^2 \leq \frac{B}{n^2} \left\{ 1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{m} + \dots + \frac{1}{n-1} + \frac{1}{n} \right\} < \frac{B}{n} \quad (B-74)$$

Hence, it can be seen that the expected mean square error of algorithm two decreases faster than that for algorithm one.

Again consider the optimum rate of convergence of algorithm two by equating the derivative of equation (B-70) with respect to γ_n to zero. This gives

$$(1-\gamma_{n+1})V_n^2 = \gamma_{n+1} \frac{\sigma^2}{n+1} \quad (B-75)$$

which when solved for γ_{n+1} , gives

$$\gamma_{n+1} = \frac{V_n^2}{V_n^2 + \frac{\sigma^2}{n}} \quad (B-76)$$

Let the initial expected mean square error be V_0^2 . Substituting into equation (B-76) gives,

$$\gamma_1 = \frac{V_0^2}{V_0^2 + \sigma^2} \quad (B-77)$$

Substituting this into equation (B-70) and solving for V_1^2 gives

$$V_1^2 = \frac{\sigma^2}{1+K} \quad (B-78)$$

where K is defined as in equation (B-39).

Now

$$\gamma_2 = \frac{V_1^2}{V_1^2 + \frac{\sigma^2}{2}} \quad (B-79)$$

and after substituting from (B-78)

$$\gamma_2 = \frac{2}{3+K} \quad (B-80)$$

Substituting for γ_2 and V_1 in equation (B-70) gives

$$V_2^2 = \frac{\sigma^2}{3+K} \quad (B-81)$$

Now assuming V_n^2 is known then

$$\gamma_{n+1} = \frac{V_n^2}{V_n^2 + \frac{\sigma^2}{n+1}} \quad (B-82)$$

Substituting into equation (B-70) for γ_{n+1} gives

$$V_{n+1}^2 = \frac{\sigma^2}{(n+1) + \frac{\sigma^2}{V_n^2}} \quad (B-83)$$

Using these two forms repeatedly gives a sequence for γ and V^2 , for example

$$\gamma_3 = \frac{3}{6+K} \quad (B-84)$$

and

$$V_3^2 = \frac{\sigma^2}{6+K} \quad (B-85)$$

and so on.

By observing the sequence one can now write the optimum forms;

$$\gamma_n^* = \frac{n}{\frac{n(n+1)}{2} + K} \quad (\text{B-86})$$

and

$$v_n^{2*} = \frac{2}{\frac{n(n+1)}{2} + K} \quad (\text{B-87})$$

Hence, it has been shown that both algorithm one and algorithm two converge in the mean square sense in the limit, and that both possess an optimum gain sequence that generates the least expected mean square error at every step.

APPENDIX C

The Discrete Maximum Principle

One of the fundamental problems in science and engineering is the optimization of some function with respect to one or more parameters and usually subject to a number of constraints. The most useful single technique in system theory used to perform the extremization is the calculus of variations. Variational principles have been applied to physical problems, for example, such as wave propagation from the time of Huygens. The Hamiltonian formulation of the variational problem has existed since the early nineteenth century in the work of Hamilton, Jacobi and others. The most significant recent contribution was made by L. S. Pontryagin. His work (27,28) extended the variational method to include problems in which the control or driving function and the state vector are bounded.

The same principle applied to continuous problems has been recently applied to problems involving discrete data systems (29). In reality, the maximum principle is not universally valid for the case of discrete systems (27). Because of restrictions on possible variations of the control signal, the maximum principle must be modified for the general discrete case. Jordan and Polak, in reference (30), discuss the limitations and derive a modified form of the

maximum principle, which is applicable to the general discrete problem. The difficulty is not of as much concern with the advent of digital controllers where now the range of control can easily extend beyond the "saturation" limits of the plant.

The derivation of Pontryagin's maximum principle is done by using the Hamiltonian formulation and is applied to the discrete version of the Bolza problem. Katz³¹ was the first to establish the discrete version of the maximum principle that was valid not only for discrete time but also space.

Given a discrete, nonlinear dynamic system with a state vector, $\underline{x}_k$ and an input vector, $\underline{u}_k$, $\underline{x}_k$ is a "n" vector; $\underline{u}_k$ is an "r" vector. The state of the system at the $(k+1)^{th}$ stage is related to the state at the k^{th} stage by the relationship

$$\underline{x}_{k+1} = \underline{f}(\underline{x}_k, \underline{u}_k, k) \quad (C-1)$$

The process begins at stage k_0 and terminates at stage k_f .

The problem is to find $\underline{x}_k$ and $\underline{u}_k$ such that the cost function

$$J = [\theta(\underline{x}_k, k)] \Big|_{k=k_0}^{k=k_f} + \sum_{k=k_0}^{k_f-1} \phi(\underline{x}_k, \underline{u}_k, k) \quad (C-2)$$

is minimized subject to the constraint of equation (C-1).

Using the Lagrange multiplier, λ_k , an equivalent cost function can be formulated as

$$J' = [\theta(\underline{x}_k, k)] \Big|_{k=k_0}^{k=k_f} + \sum_{k=k_0}^{k_f-1} \{ \phi(\underline{x}_k, \underline{u}_k, k) + \lambda_{k+1}^T \underline{x}_{k+1} - f(\underline{x}_k, \underline{u}_k, k) \} \quad (C-3)$$

Now define the Hamiltonian as

$$H[\underline{x}_k, \underline{u}_k, \lambda_{k+1}, k] = \phi(\underline{x}_k, \underline{u}_k, k) + \lambda_{k+1}^T f(\underline{x}_k, \underline{u}_k, k) \quad (C-4)$$

This gives the cost function

$$J' = [\theta(\underline{x}_k, k)] \Big|_{k=k_0}^{k=k_f} + \sum_{k=k_0}^{k_f-1} [H_k - \lambda_{k+1}^T \underline{x}_{k+1}] \quad (C-5)$$

The cost function J' may be minimized with respect to $\underline{x}_k$ and $\underline{u}_k$ by application of the perturbation methods of the calculus of variations. Let independent perturbations be introduced into the state and input vectors such as

$$\underline{x}_k = \underline{x}_k^* + \epsilon \underline{n}_k \quad (C-6)$$

$$\underline{x}_{k+1} = \underline{x}_{k+1}^* + \epsilon \underline{n}_{k+1} \quad (C-7)$$

$$\underline{u}_k = \underline{u}_k^* + \epsilon \underline{v}_k \quad (C-8)$$

Note that the perturbations at different stages are independent; hence, $\underline{n}_k$, $\underline{n}_{k+1}$, and $\underline{v}_k$ are all mutually independent.

Introducing the perturbations into equation (C-5) gives

$$J' = \theta(\underline{x}_{k_f}^* + \epsilon \underline{n}_{k_f}, k_f) - \theta(\underline{x}_{k_0}^* + \epsilon \underline{n}_{k_0}, k_0) + \sum_{k=k_0}^{k_f-1} [H(\underline{x}_k^* + \epsilon \underline{n}_k, \underline{u}_k^* + \epsilon \underline{v}_k, \lambda_{k+1}, k) + \lambda_{k+1}^T [\underline{x}_{k+1}^* + \epsilon \underline{n}_{k+1}]] \quad (C-9)$$

It is known that for a minimum of J' in (C-9) it is required that

$$\frac{\partial J'}{\partial \epsilon} = 0 \quad (C-10)$$

and

$$\frac{\partial^2 J'}{\partial \epsilon^2} < 0 \quad (C-11)$$

for $\epsilon=0$, independent of the variations. In this development it will be assumed that the second derivative requirement is satisfied for all cost functions and systems of interest.

Now equating to zero the first derivative, equation (C-10), requires that

$$\begin{aligned} \left[\frac{\partial \theta_{k_f}}{\partial x_{k_f}^*} \right]^T \underline{n}_{k_f} - \left[\frac{\partial \theta_{k_0}}{\partial x_{k_0}^*} \right]^T \underline{n}_{k_0} + \sum_{k=k_0}^{k_f-1} \left[\frac{\partial H_k}{\partial x_k^*} \right]^T \underline{n}_k - \sum_{k=k_0}^{k_f-1} \lambda_{k+1}^T \underline{n}_{k+1} \\ + \sum_{k=k_0}^{k_f-1} \left[\frac{\partial H_k}{\partial u_k^*} \right]^T \underline{v}_k = 0 \quad (C-12) \end{aligned}$$

Using the discrete version of integration by parts, the fourth term of equation (C-12) can be written as:

$$- \sum_{k=k_0}^{k_f-1} \lambda_{k+1}^T \underline{n}_{k+1} = - \sum_{k=k_0+1}^{k_f} \lambda_k^T \underline{n}_k$$

or

$$- \sum_{k=k_0}^{k_f-1} \lambda_{k+1}^T \underline{n}_{k+1} = - \sum_{k=k_0}^{k_f-1} \left\{ \left| \lambda_k^T \underline{n}_k \right| - \lambda_{k_f}^T \underline{n}_{k_f} + \lambda_{k_0}^T \underline{n}_{k_0} \right\} \quad (C-13)$$

Using this equation in (C-12), and combining term gives,

$$\begin{aligned} & \left[\begin{bmatrix} \frac{\partial \theta_{k_f}}{\partial \underline{x}_{k_f}} \end{bmatrix}^T - \underline{\lambda}_{k_f}^T \right] \underline{n}_{k_f} - \left[\begin{bmatrix} \frac{\partial \theta_{k_o}}{\partial \underline{x}_{k_o}} \end{bmatrix}^T - \underline{\lambda}_{k_o}^T \right] \underline{n}_{k_o} \\ & + \sum_{k=k_o}^{k_f-1} \left[\begin{bmatrix} \frac{\partial H_k}{\partial \underline{x}_k} \end{bmatrix}^T - \underline{\lambda}_k^T \right] \underline{n}_k + \sum_{k=k_o}^{k_f-1} \left[\begin{bmatrix} \frac{\partial H_k}{\partial \underline{u}_k} \end{bmatrix}^T \right] \underline{v}_k = 0 \end{aligned} \quad (C-14)$$

Since the indicated variables are mutually independent, equation (C-14) requires that

$$\underline{\lambda}_k = \frac{\partial H_k}{\partial \underline{x}_k} \quad (C-15)$$

$$\frac{\partial H_k}{\partial \underline{u}_k} = 0 \quad (C-16)$$

$$\underline{\lambda}_{k_o} = \frac{\partial \theta_{k_o}}{\partial \underline{x}_{k_o}} \quad \text{or} \quad \underline{n}_{k_o}^T \left[\underline{\lambda}_{k_o} - \frac{\partial \theta_{k_o}}{\partial \underline{x}_{k_o}} \right] = 0 \quad (C-17)$$

and

$$\underline{\lambda}_{k_f} = \frac{\partial \theta_{k_f}}{\partial \underline{x}_{k_f}} \quad \text{or} \quad \underline{n}_{k_f}^T \left[\underline{\lambda}_{k_f} - \frac{\partial \theta_{k_f}}{\partial \underline{x}_{k_f}} \right] = 0 \quad (C-18)$$

for the general case in which there are no prescribed constraints on the variables. If the value of any variable is specified, the corresponding variation vanishes, and the corresponding requirement in equations (C-15) through (C-18) does not apply. For example, if $\underline{u}_k$ is a known, deterministic function, then $\underline{v}_k=0$ and the requirement (C-16) does not pertain. Similarly, if any component of $\underline{x}_{k_o}$ or $\underline{x}_{k_f}$ is specified, the corresponding boundary condition on $\underline{\lambda}_k$ does not apply.

Hence, solving the difference equations resulting from (C-15) and (C-16) with the boundary conditions given by (C-17)

and (C-18), assuring that the solution satisfies the state equations (C-1), results in an optimum trajectory that minimizes the cost function J .

APPENDIX D

Dvoretzky's Theorem on Stochastic Approximation

D-1 Introduction

Dvoretzky's theorem is a general theorem that encompasses the works of Robbins - Monro, Kiefer - Wolfowitz and many others. It was the first unifying theory to appear in the area of stochastic approximation and still stands as the corner stone in its field. The basis of the approach to the problem was to consider a convergent deterministic scheme with a random element analagous to noise superimposed on the scheme.

The approach to the theorem was such that Dvoretzky first proved a basic theorem according to strict conditions. He furthered the theory by adding an extension, six generalizations, two resultant corollarys and a special case of his generalized theorem. This approach was taken to make the proofs clearer and more interesting without forsaking universal generality.

A statement of most sections of the theorem will be given, however, since the proof is very lengthy and bears no immediate relationship to the thesis, it will not be given here. For the details of the proof reference (3) is recommended.

D-2 Theorem:

Let α_n , β_n , and γ_n , $n=1,2,\dots$, be non-negative real numbers satisfying

$$\lim_{n \rightarrow \infty} \alpha_n = 0 \quad (D-2.1)$$

$$\sum_{n=1}^{\infty} \beta_n < \infty \quad (D-2.2)$$

and

$$\sum_{n=1}^{\infty} \gamma_n = \infty \quad (D-2.3)$$

Let x be a real number and T_n , $n=1,2,\dots$, be measurable transformations satisfying

$$|T_n(\rho_1, \dots, \rho_n) - x| \leq \max [\alpha_n, (1+\beta_n)|\rho_n - x| - \gamma_n] \quad (D-2.4)$$

for all real $\rho_1, \dots, \rho_n$. Let X_1 and Y_n , $n=1,2,\dots$, be random variables and define

$$X_{n+1} = T_n[X_1, \dots, X_n] + Y_n \quad (D-2.5)$$

for $n \geq 1$. Then the conditions $E\{X_1^2\} < \infty$,

$$\sum_{n=1}^{\infty} E\{Y_n^2\} < \infty \quad (D-2.6)$$

and

$$E\{Y_n | X_1, \dots, X_n\} = 0 \quad (D-2.7)$$

with probability 1 of all n . Then

$$\lim_{n \rightarrow \infty} E\{(X_n - \chi)^2\} = 0 \quad (D-2.8)$$

and

$$P\{\lim_{n \rightarrow \infty} X_n = \chi\} = 1 \quad (D-2.9)$$

This basic part of Dvoretzky's work assumes that the sequences α_n , β_n , and γ_n are independent of the observations

D-3 The Extension:

The theorem remains valid if α_n , β_n , γ_n in (D-2.4) are replaced by non-negative functions $\alpha_n(\rho_1, \dots, \rho_n)$, $\beta_n(\rho_1, \dots, \rho_n)$ and $\gamma_n(\rho_1, \dots, \rho_n)$, respectively, provided they satisfy the conditions:

The functions $\alpha_n(\rho_1, \dots, \rho_n)$ are uniformly bounded and

$$\lim_{n \rightarrow \infty} \alpha_n(\rho_1, \dots, \rho_n) = 0 \quad (D-3.1)$$

uniformly for all sequences $\rho_1, \dots, \rho_n, \dots$; the functions $\beta_n(\rho_1, \dots, \rho_n)$ are measurable and

$$\sum_{n=1}^{\infty} \beta_n(\rho_1, \dots, \rho_n) \quad (D-3.2)$$

is uniformly bounded and uniformly convergent for all sequences $\rho_1, \dots, \rho_n, \dots$; and the functions $\gamma_n(\rho_1, \dots, \rho_n)$ satisfy

$$\sum_{n=1}^{\infty} \gamma_n(\rho_1, \dots, \rho_n) = \infty \quad (D-3.3)$$

uniformly for all sequences $\rho_1, \dots, \rho_n, \dots$, for which

$$\sup_{n=1,2,\dots} |\rho_n| < L \quad (D-3.4)$$

L being an arbitrary finite number.

D-4 Generalizations

Particularization: For $x=0$ condition (D-2.4) is weaker than the following,

$$|T_n(\rho_1, \dots, \rho_n)| \leq \max [\alpha_n, (1+\beta_n-\gamma_n)|\rho_n|] \quad (D-4.1)$$

with α_n , β_n , and γ_n still satisfying (D-2.1), (D-2.2) and (D-2.3) respectively.

Generalization 1: The Extended Theorem remains valid if (D-2.7) is replaced by

$$\sum_{n=1}^{\infty} \sup_{x_1, \dots, x_n} |E\{Y_n | x_1, \dots, x_n\}| < \infty \quad (D-4.2)$$

or even by the condition that

$$\sum_{n=1}^{\infty} E\{Y_n | x_1, \dots, x_n\} \quad (D-4.3)$$

be uniformly bounded and uniformly convergent for all sequences $x_1, \dots, x_n, \dots$.

Generalization 2: Conclusion (D.2.9) of the Extended Theorem remains valid even without the restrictions on X_1 and if (D-2.5) is replaced by

$$X_{n+1} = T_n(X_n) + Y_n^* \quad (D-4.4)$$

with the random variables Y_n^* , $n=1,2,\dots$, satisfying

$$P\{Y_n^* \neq Y_n \text{ for infinitely many } n\} = 0 \quad (D-4.5)$$

thus, in particular, when

$$\sum_{n=1}^{\infty} P\{Y_n^* \neq Y_n\} < \infty \quad (D-4.6)$$

Generalization 3: If (D-2.1) is replaced by

$$\overline{\lim}_{n=\infty} \alpha_n = \alpha \quad (D-4.7)$$

or more generally (D-3.1) by

$$\overline{\lim}_{n=\infty} \alpha_n(\rho_1, \dots, \rho_n) \leq \alpha \quad (D-4.8)$$

uniformly for all sequences $\rho_1, \dots, \rho_n, \dots$ then the Extended Theorem remains valid provided (D-2.8) and (D-2.9) are replaced by

$$\overline{\lim}_{n=\infty} E\{(X_n - \bar{x})^2\} \leq \alpha^2 \quad (D-4.9)$$

and

$$P \overline{\lim}_{n=\infty} |X_n| \leq \alpha = 1 \quad \text{respectively.} \quad (D-4.10)$$

Generalization 4: The Extended Theorem remains valid if the assumptions concerning $\alpha_n(\rho_1, \dots, \rho_n)$ are replaced by the following:

$\alpha_1(X_1)$ is bounded with probability 1,

$$\alpha_n(X_1, \dots, X_n) \geq \alpha_{n+1}(X_1, \dots, X_n, X_{n+1}) \quad (D-4.11)$$

with probability 1 and

$$P\{\lim_{n \rightarrow \infty} \alpha_n(X_1, \dots, X_n) = 0\} = 1 \quad (\text{D-4.12})$$

Special Case: If the transformations T_n of (D-2.5) satisfy

$$|T_n(\rho_1, \dots, \rho_n)^{-\chi}| \leq F_n |\rho_n^{-\chi}| \quad (\text{D-4.13})$$

for F_n , $n=1, 2, \dots$, being a sequence of positive numbers satisfying

$$\prod_{n=1}^{\infty} F_n = 0 \quad (\text{D-4.14})$$

then the basic theorem holds.

APPENDIX E

Cauchy's Integral Test

Fundamental Principle of Monotone Convergence: If a sequence $\{S_n\}$ satisfies $S_n \leq S_{n+1} \leq M$ for each n , where M is some constant, then $\lim_{n \rightarrow \infty} S_n$ exists.

In other words, every bounded increasing sequence has a limit.

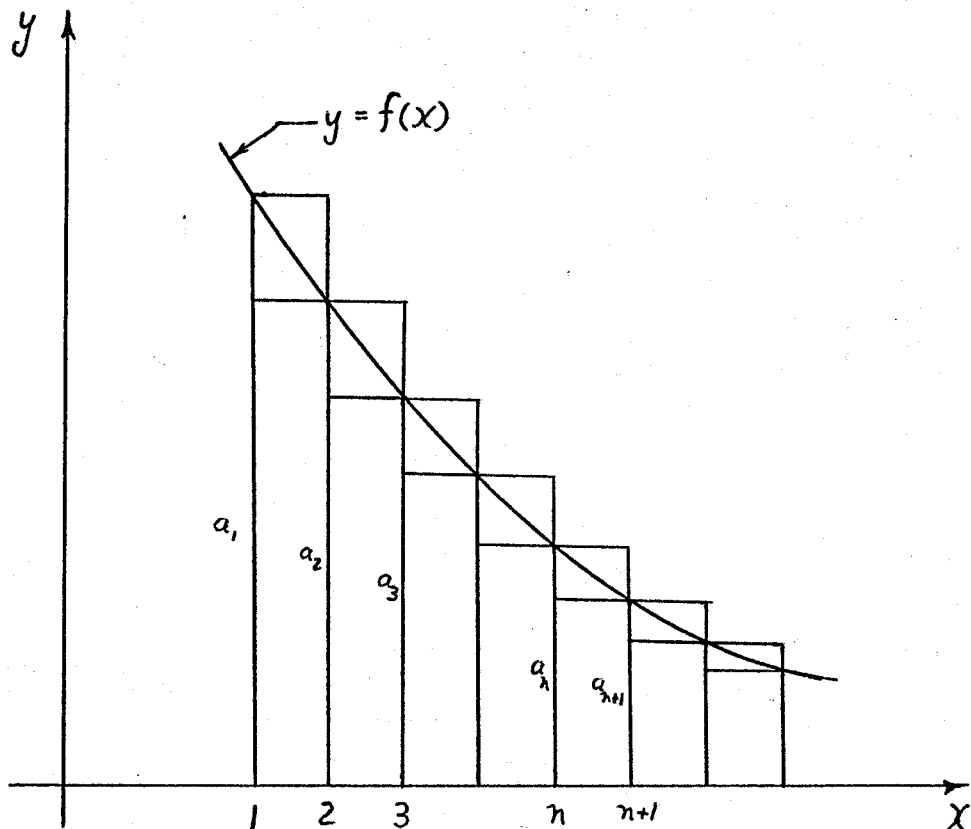


Figure (E-1). Schematic Representation of Locus of $f(x)$

Theorem I: The improper integral

$$\int_1^{\infty} \frac{1}{x^p} dx \quad (E-I.1)$$

converges if, and only if, the constant $p > 1$.

Proof: Consider the integral

$$\int_1^{\infty} \frac{1}{x^p} dx = \lim_{b \rightarrow \infty} \int_1^b x^{-p} dx$$

which gives

$$\lim_{b \rightarrow \infty} \int_1^b x^{-p} dx = \lim_{b \rightarrow \infty} \left[\frac{x^{1-p}}{1-p} \right]_1^b \quad \text{for } p \neq 1$$

resulting in

$$\lim_{b \rightarrow \infty} \left[\frac{x^{1-p}}{1-p} \right]_1^b = \lim_{b \rightarrow \infty} \frac{b^{1-p}-1}{1-p}$$

The question of convergence depends on the behaviour of b^{1-p} as $b \rightarrow \infty$. If the exponent $1-p$ is positive, $b^{1-p} \rightarrow \infty$ and the integral is divergent. But if $1-p$ is negative, then $p-1 > 0$ and hence

$$b^{1-p} = \frac{1}{b^{p-1}} \rightarrow 0 \quad \text{as } p \rightarrow \infty$$

For this case only the integral converges to the value $1/p-1$.

Theorem II: The infinite series $\sum_{k=1}^{\infty} \frac{1}{k^p}$ converges if, and only if, the constant $p > 1$.

This proof follows analogously from Theorem I.

Theorem III: For $x \geq 1$, let $f(x)$ be positive, continuous, and decreasing. Then the series

$$\sum_{n=1}^{\infty} f(n) \quad (\text{E-III.1})$$

and the integral

$$\int_1^{\infty} f(x) dx \quad (\text{E-III.2})$$

both converge or both diverge. In either case, the partial sums are bounded as follows:

$$\int_1^{n+1} f(x) dx < \sum_{k=1}^n f(k) < \int_1^n f(x) dx + f(1) \quad (\text{E-III.3})$$

Proof: Suppose the terms of an infinite series $\sum a_k$ are positive and decreasing; that is, $a_n > a_{n+1} > 0$ for each positive integer n . In this case, there is a continuous decreasing function $f(x)$ such that

$$a_n = f(n) \quad n=1, 2, \dots$$

Each term a_n of the series may be thought of as representing the area of a rectangle of base unity and height $f(n)$ (cf. fig). The sum of the areas of the first n circumscribed rectangles is greater than the area under the curve from 1 to $n+1$, so that

$$a_1 + a_2 + \dots + a_n > \int_1^{n+1} f(x) dx$$

This shows that, if the integral $\int_1^{\infty} f(x) dx$ diverges, the sum

$\sum a_k$ also diverges.

Alternately, the sum of the areas of the inscribed rectangles is less than the area under the curve, so that

$$a_2 + a_3 + \dots + a_n < \int_1^n f(x) dx$$

If the integral converges, since $f(x) > 0$, then

$$\int_1^n f(x) dx < \int_1^\infty f(x) dx \equiv M$$

so that the partial sums are bounded independently of n :

$$S_n = a_1 + a_2 + \dots + a_n$$

$$S_n < M + a_1$$

Since each a_k is positive, these partial sums form an increasing sequence. Hence, the fundamental principle ensures that $\sum a_k$ is convergent.

APPENDIX F

Noise, Spectra and Autocorrelation

The element ξ of the random process considered in the theory can be thought of as noise. The noise to be of value in the simulation was required to be of Gaussian distribution, with stationary statistics and of wide bandwidth, ie, as close to white as possible. Essentially, what was used in the computer simulation was a pseudo-random noise sequence. An outline of the basic process will be given.

First a sequence of pseudo-random pulses³² that 1's and 0's were generated. These were interpreted in groups of 26 (the bit length chosen) as integer numbers. It turns out that the basic property of pseudo-random numbers is that they are evenly distributed between 0 and $2^{26}-1$. In fact, $2^{26}-1$ numbers, excluding all zeros, will be generated before the sequence repeats. The resulting number was then converted to a real number and scaled to the range 0-1.0. Hence, pseudo-random numbers of uniform distribution between 0 and 1.0 had been generated.

To get a normal distribution, use was made of the central limit theorem and the uniformly distributed numbers generated above. An approximation was made by making use of the following formula³⁴,

$$y_n = \frac{\sum_{i=1}^k x_i - \frac{k}{2}}{\sqrt{k/12}} \quad (F-1)$$

where

y_n is a normally distributed number with variance of 1 and mean of zero,

x_i are uniformly distributed $0 < x_i < 1$

and k is the number of x_i used.

This, in fact, is an approximation to numerical convolution.

Adjustment for the required mean and standard deviation is then

$$y_{ns} = \sigma Y_n + \mu \quad (F-2)$$

where y_{ns} is a scaled normally distributed number of

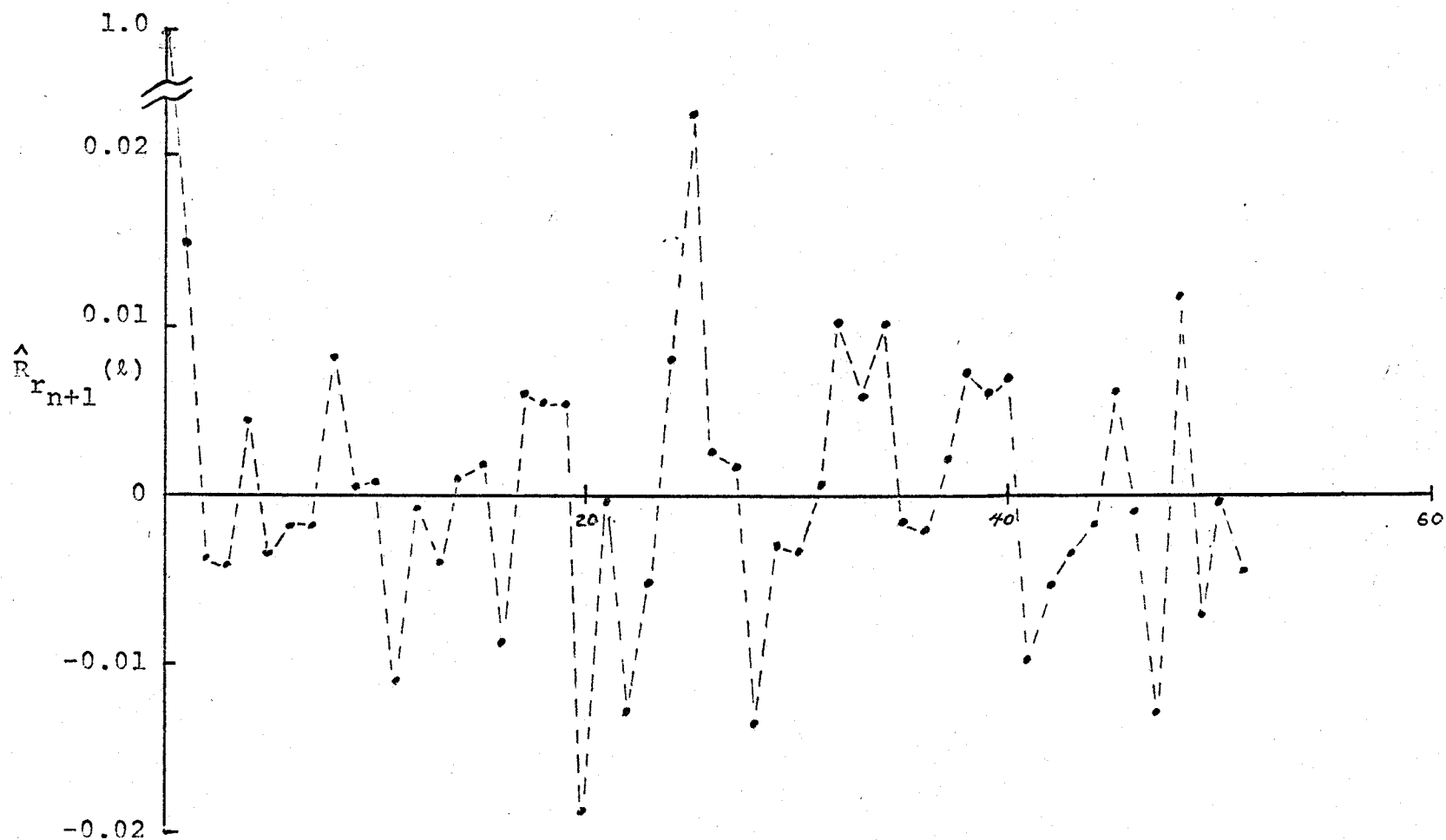
standard deviation σ

and arithmetic mean μ

In the actual simulation k was taken as 12 and this was found to be a good value not requiring much computation time.

The theory presented in this thesis depends much on the fact that the noise is white. To check for this property an autocorrelation was done on a large sample of data and the results shown on Graph (F-1)

Along with this, the power spectrum of the noise produced by the sub-routines was looked at. The fast Fourier transform technique was used and the calculated power spectra was plotted and shown in Graphs (F-2), (F-3), and (F-4).

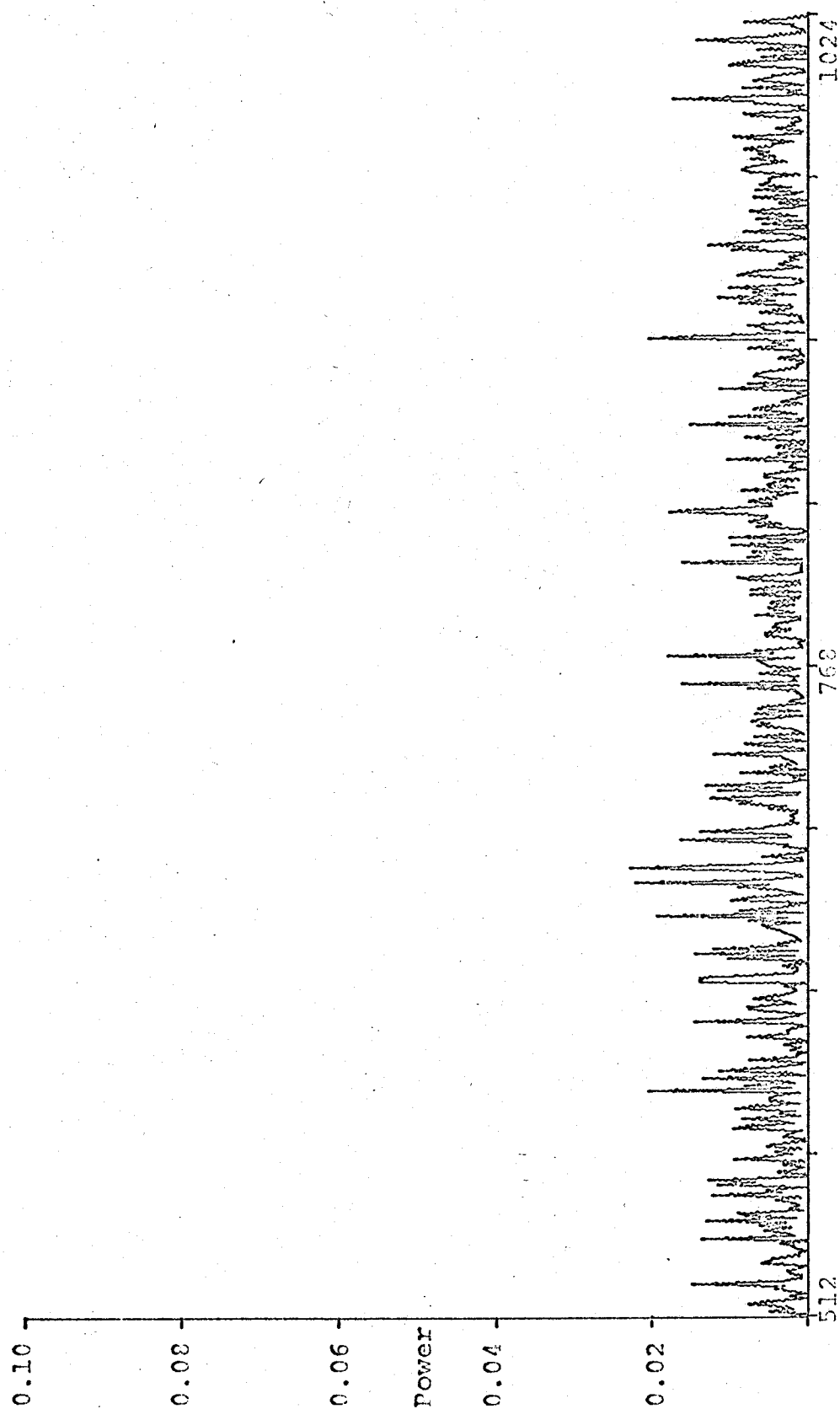


Graph (F-1) Autocorrelation Function of Pseudo-random noise

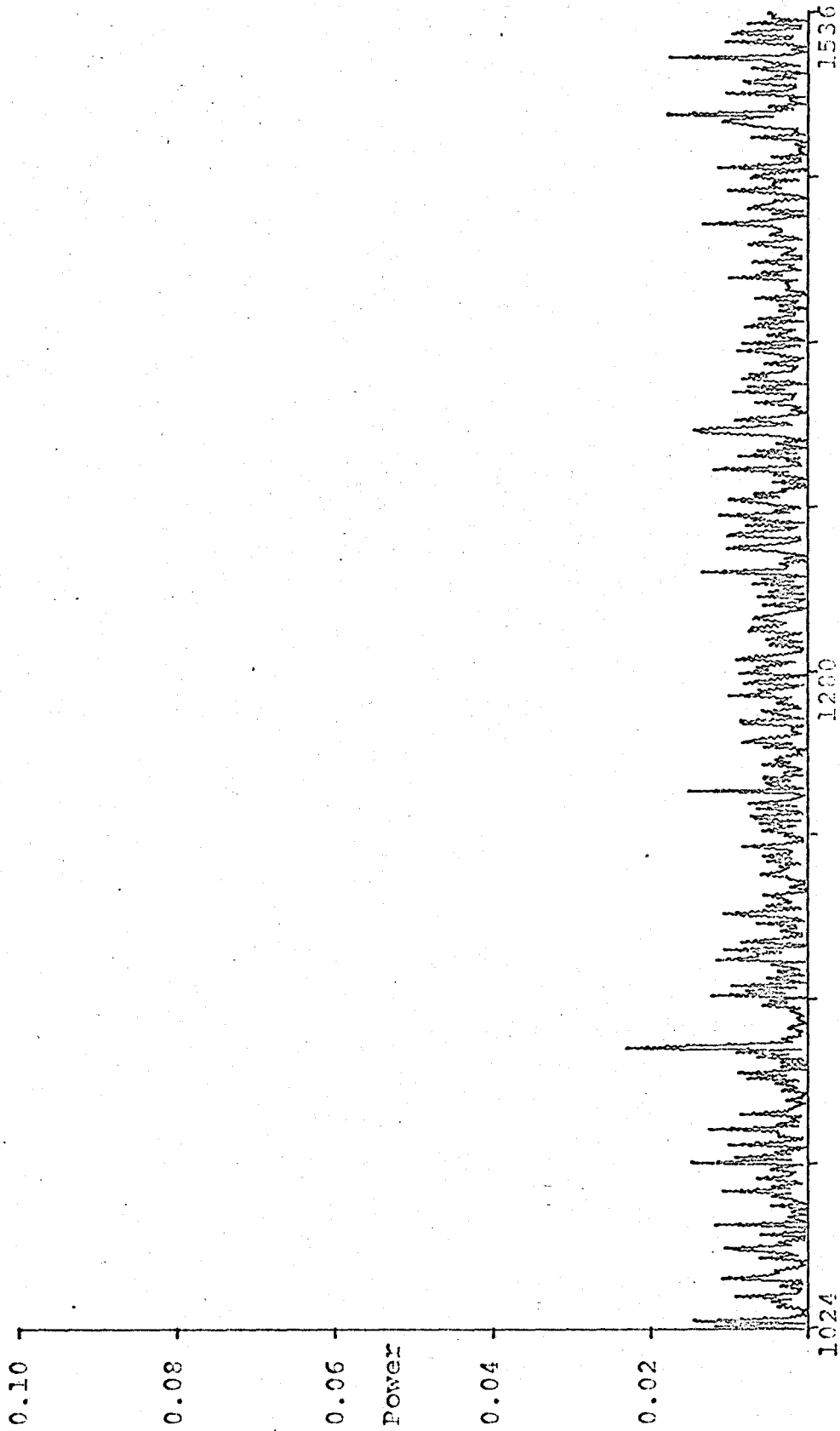
$$\{\hat{R}_{r_{n+1}}(\ell) \text{ vs } \ell; n+1 = 20,000; 0 \leq \ell \leq 51\}$$



Graph (F-2) Fast Fourier Power Spectra of noise



Graph (F-3) Fast Fourier Power Spectra of noise (con't)



Graph (F-4) Fast Fourier Power Spectra of Noise (con't)