

LEVEL

2

AD A105780

(6) AN APPLICATION OF SUBJECTIVE PROBABILITIES TO
THE PROBLEM OF UNCERTAINTY IN COST ANALYSIS

(10) HARLEY R. JORDAN
MICHAEL R. KLEIN

(11) NOV 1975

DTIC
ELECTE
OCT 1 9 1981



OFFICE OF THE CHIEF OF NAVAL OPERATIONS
Resource Analysis Group
Systems Analysis Division (OP-96D)
Room 2C340 - Pentagon
Washington, DC 20350

DTIC FILE COPY

DISTRIBUTION STATEMENT A

Approved for public release;
Distribution Unlimited

406411

X
Dr 10 10 19

ABSTRACT

All cost estimates are characterized by some uncertainty. A device helpful in communicating this uncertainty to the decision maker is a subjective probability distribution of the system cost. A technique--termed the Subjective Probability Estimation Technique (SPET)--is described and a computer program is presented to facilitate its use. This technique permits the analyst to represent his notions about cost uncertainty with the beta or other statistical distributions.

Accession For	
NTIS GRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	<i>for file</i>
By	<i>[Signature]</i>
Distribution/	
Availability Codes	
Dist	Avail and/or Special
<i>A</i>	

ACKNOWLEDGMENTS

The authors are grateful to Mr. Joseph T. Kammerer, Director of the Resource Analysis Group, for the suggestion and encouragement to pursue this research. Along with many valuable comments, Mr. Kammerer provided the algorithm^{1/} for computer program PARAM, which appears in Appendix A of this paper.

Several other individuals offered helpful suggestions during the course of this research, particularly Mr. Carl Wilbourn of the Resource Analysis Group and Mr. Kenneth Linder of the Office of Program Analysis and Evaluation, in the Department of Health, Education, and Welfare.

Those familiar with the literature on uncertainty in cost analysis will recognize a basic similarity between the technique described in this paper and the one expounded by Dienemann.^{2/} This paper is a refinement and expansion of certain elements of his earlier research.

-
- 1/ See J. T. Kammerer, ASW Force Level Study - Equipment Readiness: Models, Computer Simulation and Results (Washington, DC: Office of the Chief of Naval Operations, 1968).
 - 2/ Paul F. Dienemann, Estimating Uncertainty Using Monte Carlo Techniques (Santa Monica, CA: The RAND Corp., RM-4854-PR, 1966).

CONTENTS

INTRODUCTION	1
SUBJECTIVE PROBABILITY ESTIMATION TECHNIQUE.	5
System Decomposition.	5
Selecting the Subjective Probability Distributions.	6
Combining the Subjective Probability Distributions.	12
The Independence Assumption: A Problem	12
SUMMARY	15
APPENDIX A - Estimation of Parameters a and b: Computer Program PARAM	A-1
APPENDIX B - Graphical Aids for Selecting the Uncertainty Coefficient	B-1
APPENDIX C - Computer Program SPET	C-1
REFERENCES	C-19

INTRODUCTION

The cost analyst is faced with many uncertainties as he attempts to estimate the costs associated with a new, undeveloped system. He may wonder:

- (1) Will the physical characteristics of the system remain unchanged as the development process proceeds?
- (2) Will there be any unforeseen problems in the development process?
- (3) Will the economic state of the firms or industry responsible for system development and production continue to change as forecasted?
- (4) Is the quality of the historical cost data sufficiently high to inspire confidence in the estimates made with it?
- (5) Have the cost-estimating relationships been properly specified?

To the extent the analyst is unable to obtain complete answers to these and similar questions, his cost estimates will be enshrouded with uncertainty. Since it is impossible to obtain definitive answers to all these questions, his cost estimates will always be characterized by some uncertainty.

The analyst can treat this uncertainty in one of several ways. He can choose to ignore it and simply report to the decision maker the estimate which represents the "most likely" or "best" estimate of cost, as in the case of this hypothetical guided missile system:

FIGURE 1

"BEST" UNIT COST ESTIMATE OF A HYPOTHETICAL
GUIDED MISSILE SYSTEM



This approach, however, belies the existence of a range of possible costs; when the system is finally acquired, any

one of the innumerable costs within this range may have been realized. Such an oversimplification may mislead the decision maker by causing him to place excessive confidence in the best cost estimate. An illustration of this potential pitfall is provided in the following example:

Suppose two alternative systems that do the same task are being compared. Suppose, too, that on balance, the differences in effectiveness, performance, growth potential, maintainability and similar considerations between the two systems are small, so that the choice is primarily one of cost. Suppose one system is estimated to cost \$1.25 million and the other \$1.50 million. Without an indication of the possible high and low values, the \$1.25 million alternative would be the logical choice. But the choice becomes more difficult if, as shown in the accompanying tabulation, the \$1.25 million cost is qualified with an estimate of a possible high value of \$2.00 million, whereas for the other alternative the limits are estimated to be tighter, and the possible high value is only \$1.60 million. The case for

<u>System</u>	<u>Cost Estimate</u> <u>(Millions of Dollars)</u>		
	<u>Lowest</u>	<u>Most Likely</u>	<u>Highest</u>
First Alternative	1.00	1.25	2.00
Second Alternative	1.40	1.50	1.60

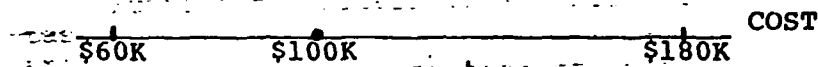
the alternative with the "most likely" cost of \$1.25 million is now more dubious, because its uncertainty spread is greater, extending on the high end to greater costs than the other alternative.^{3/}

The analyst can avoid the problem associated with a single best cost estimate by supplying the decision maker with estimates of the lowest and highest possible costs in addition to the best estimates:

^{3/} W. Sutherland, Fundamentals of Cost Uncertainty Analysis (McLean, VA: Research Analysis Corp., RAC-CR-4 1971), pp 3-4.

FIGURE 2

"BEST" AND RANGE ESTIMATE OF THE UNIT COST
OF A HYPOTHETICAL GUIDED MISSILE SYSTEM



This approach has merit in that it reflects the range of the cost uncertainty. However, it gives little information about the nature of the uncertainty, e.g., whether all the values in the range are almost equally likely to occur, or whether the values closer to the best estimate are much more likely to occur than those near the extremities. Such knowledge could be valuable to the decision maker.

One device helpful in communicating knowledge of both the range and nature of the uncertainty is the probability distribution of the total system cost. To illustrate, consider a probability distribution of the cost of the hypothetical guided missile system:

FIGURE 3

PROBABILITY DISTRIBUTION OF THE UNIT COST
OF A HYPOTHETICAL GUIDED MISSILE SYSTEM



Note that the range of, say, a 95 percent confidence interval [\$80K, \$140K] is significantly smaller than the full range [\$60K, \$180K]. The knowledge that the analyst is 95 percent confident that the cost will occur in the much smaller interval [\$80K, \$140K] will permit the decision maker to act with a more precise idea of the probable cost of the system than he could otherwise (providing, of course, that he trusts the analyst's judgement).

A popular source for the probability distribution of the cost of a weapon system is the prediction interval obtained from cost-estimating relationships (CERs) developed by regression analysis. Unfortunately, this method of probability analysis has a serious limitation: there is no provision for the analyst to incorporate into the anal-

ysis his notions about the stochastic behavior of the system cost. For example, most cost analysts have a good idea of the lower bound on the cost, but are less certain about the upper bound. This suggests a probability distribution that is positively skewed:

FIGURE 4

POSITIVELY SKEWED PROBABILITY DISTRIBUTION



The user of the prediction interval obtained from classical regression analysis, however, has to accept a symmetric probability distribution -- the normal probability distribution -- often against his better judgement:

FIGURE 5

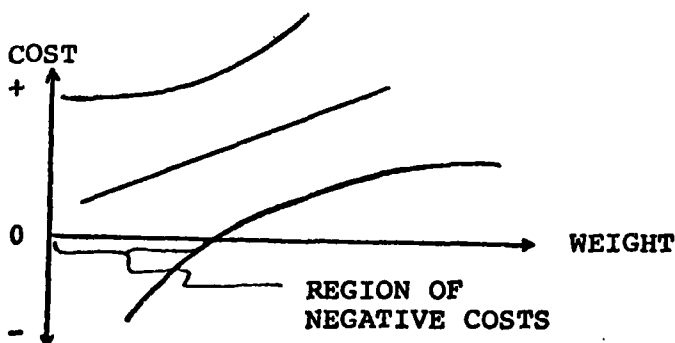
SYMMETRIC PROBABILITY DISTRIBUTION



In addition to this limitation, the paucity (and variability) of data used in regression analysis of weapon system costs often results in prediction intervals with lower bounds which include an extensive region of negative costs:

FIGURE 6

PREDICTION INTERVAL ABOUT A
HYPOTHETICAL REGRESSION OF COST ON WEIGHT



The fact that this model predicts the impossible -- negative costs -- again reveals the above-stated limitation of this approach to probability analysis. What the analyst needs is a technique that will permit him to retain the "most-likely" estimate of system cost and incorporate his a priori beliefs into the prediction interval.

This paper describes such a technique. The authors have entitled it the Subjective Probability Estimation Technique (SPET). This technique is based on the same principles used by Program Evaluation Review Technique (PERT) analysts years ago to treat time-estimating uncertainty.^{4/} As its name suggests, SPET accounts for the fact that the probability distribution of the cost of a new system is by necessity subjective since repeated observations on the cost of the system, from which an objective probability distribution could be inferred, cannot be made (there is only one observation on the cost of a new system -- the final acquisition cost of the system -- and when this observation is made the need for an estimate terminates). The analyst implements SPET in three steps by:

- (1) decomposing the system under examination into several subsystems whose costs are additive;
- (2) selecting the subjective probability distribution that best represents his knowledge and judgement about the cost of a subsystem;
- (3) combining the subjective probability distribution of each subsystem cost into a subjective probability distribution of total system cost.

The remainder of this paper discusses the theoretical and practical aspects of these three steps.

SYSTEM DECOMPOSITION

When a cost analyst estimates the cost of a complex system he generally breaks the system down into several

^{4/} F. S. Hillier and G. J. Lieberman, Introduction to Operations Research (San Francisco, CA: Holden-Day, Inc. 1967), pp. 227-229-232. See also K. R. MacCrimmon and C. A. Ryavec, An Analytical Study of the PERT Assumptions (Santa Monica, CA: The RAND Corp. RM-3408-PR 1962).

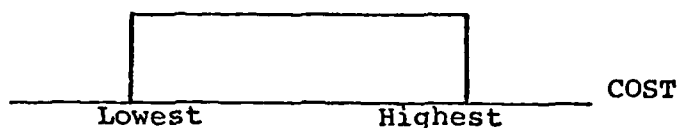
subsystems and estimates the cost of each subsystem. The breakdown of the system is usually determined by the analyst's knowledge of the system and the form of his data base. Using the technique described in this paper, the analyst will also develop a subjective probability distribution describing his uncertainty as to the cost of these subsystems. Of course, if some subsystem cost is known precisely, no uncertainty is involved and this cost is treated as a constant.

SELECTING THE SUBJECTIVE PROBABILITY DISTRIBUTION

A probability distribution can be selected to represent any imaginable combination of factual knowledge and subjective notions an analyst might have about the cost of a subsystem. For example, suppose the analyst has a good idea of what the lowest and highest possible costs for a subsystem could be, but he feels that all costs within that range are equally likely. His subjective probability distribution for the cost of this subsystem can be represented quite adequately with the uniform probability density function:

FIGURE 7

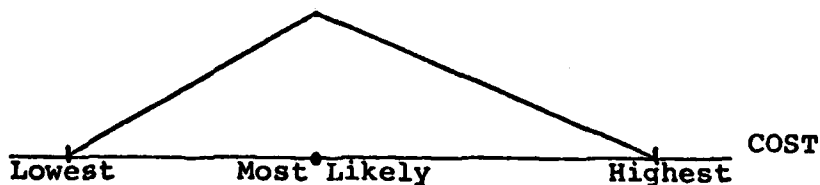
UNIFORM PROBABILITY DENSITY FUNCTION



Suppose that instead of feeling that all costs within the range are equally likely the analyst feels that a particular cost within the range is more likely to be realized than any other. He could represent the cost with a triangular distribution:

FIGURE 8

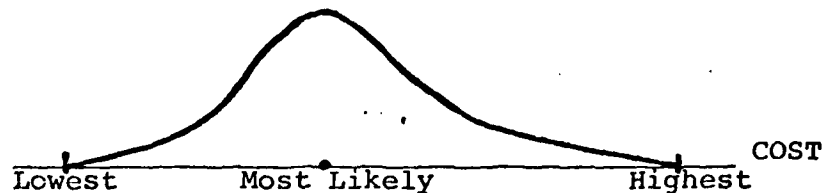
TRIANGULAR PROBABILITY DISTRIBUTION



or the beta distribution:

FIGURE 9

BETA PROBABILITY DISTRIBUTION



or any other distribution with a single maximum value.

The beta distribution is one of the most popular among cost analysts for representing subjective probability distributions. The popularity of the beta is due to its several appealing characteristics. One of these characteristics is that its range may be restricted to positive values; costs similarly are positive. Another is that the beta has a finite, rather than infinite range; it is reasonable to suppose that the cost is bounded by finite upper and lower bounds. Finally, the beta distribution can be expressed in an infinite variety of skewed and symmetric forms which provide the analyst considerable choice when specifying the particular shape of the distribution.^{5/}

Because of the popularity of the beta distribution, the discussion in the remainder of this section will center on it. In the next section, however, a computer program is described that permits the analyst to use any imaginable subjective probability distribution to represent subsystem cost.

^{5/} Another commonly used distribution is the Weibull. Most of the appealing properties of the beta are also found in the Weibull distribution. For examples of the use of the Weibull in treating uncertainty in cost analysis see D. F. Schaefer, et. al., A Monte Carlo Simulation Approach to Cost-Uncertainty Analysis, (McLean, VA: Research Analysis Corp.; RAC-TP-349, 1969) and W. H. Sutherland, A Method for Combining Asymmetric Three-Value Predictions of Time or Cost (McLean, VA: Research Analysis Corp.; RAC-P-65, 1972).

The usual expression for the beta probability density function (pdf) is:^{6/}

$$f(x) = Cx^a(1 - x)^b; 0 < x < 1; a, b > 0; \quad [1]$$

where $C = \Gamma(a + b + 2)/[\Gamma(a + 1)\Gamma(b + 1)]$

= the inverse of the complete beta function

$$\text{and } \Gamma(t) = \int_0^\infty z^{t-1}e^{-z}dz, t > 0.$$

This version of the beta pdf will be called the normalized beta pdf, since the range of x is the unit interval.

A simple linear transformation, $x^* = L + (H - L)x$, where L and H are the lowest and highest values of x^* , respectively, extends the range of x in equation [1] to any finite interval, yielding a generalized beta pdf:

$$g(x^*) = [C/(H - L)^{a+b+1}](x^* - L)^a(H - x^*)^b \quad [2]$$

where C is as defined in equation [1] and $L < x^* < H$, $a, b > 0$.

The four parameters of the generalized beta pdf are a , b , L , and H .^{7/} Therefore, a unique pdf is defined for every four-tuple (a, b, L, H) . The values the analyst assigns to these parameters can be obtained through certain estimation procedures.

The analyst may estimate L and H directly from his and other expert knowledge of the subsystem's technology, contractor (builder), industry, etc. After analyzing

^{6/} H. J. Larson, *Introduction to Probability Theory and Statistical Inference* (New York: John Wiley and Sons, Inc., 1969), p. 305; B. W. Lindgren, *Statistical Theory* (New York: The MacMillan Co., 1968), p. 373.

^{7/} Note that L and H serve only to specify the origin and range of x^* , whereas a and b determine the shape of the pdf of x^* .

these data, he chooses L and H so that the cost of the subsystem could never be less than L or greater than H.

Estimates of the parameters a and b, however, cannot be obtained in a direct manner. One way to estimate them is to obtain two functionally independent equations in a and b and then solve them for these parameters. In Appendix A, two such equations are proposed and a computer program facilitating their solution is documented. Potential users of this method for estimating a and b are cautioned; some combinations of the analyst-supplied inputs result in distributions that cannot be represented by a beta random variable. This problem can be avoided by using another pair of equations. Consider the mode $M(x^*)$ and the variance $V(x^*)$ of the generalized beta distribution:^{8/}

$$M(x^*) = (aH + bL)/(a + b), \quad L \leq M(x^*) \leq H \quad [3]$$

$$V(x^*) = [(a + 1)/(b + 1)(H - L)^2]/[a + b + 2]^2 (a + b + 3)], \quad 0 \leq V(x^*) \leq (H - L)^2/12 \quad [4]$$

Using cost-estimating relationships or other methods, the analyst can estimate the mode or most likely value of the subsystem cost.^{9/} By evaluating much of the information he has about the subsystem cost he can estimate its

^{8/} See Footnote 2 in Appendix A for the derivation of these formulae.

The range of the mode is obtained from its definition. The range of the variance is due, in part, to the fact that the lower bound on the variance of any distribution is zero, and, in part, to the fact that the beta distribution converges to the uniform as parameters a and b approach zero. The upper bound on the variance of the generalized beta distribution is, therefore, the variance of the (generalized) uniform distribution, namely

$$V(x^*) = (H - L)^2/12 \quad [5]$$

^{9/} Under most circumstances the most-likely value is the analyst's point estimate of the subsystem cost.

variance.^{10/} Having obtained estimates for these two statistics, he can solve equations [3] and [4] simultaneously to determine unique values (estimates) for a and b.

The principal difficulty with the technique proposed is in the estimation of the variance. It is difficult for the analyst to interpret his beliefs concerning the uncertainty surrounding a point estimate in terms of the beta variance. As an aid in this process, the analyst may consider a related variable, which is termed an uncertainty coefficient in this paper. The uncertainty coefficient represented with the letter "U" is a normed linear measure of the analyst's uncertainty. If there is no uncertainty in the estimate, $U = 0$; if there is total uncertainty on the whole range (i.e., all values are equally likely), $U = 1$. In most instances, the analyst can assign a reasonable value to U. The variance of x^* can then be determined from the relationship:

$$V(x^*) = [(H - L)U]^2/12, \quad 0 < U < 1. \quad [6]$$

It is difficult for the analyst to visualize the distribution he has chosen from his estimates of the parameters. However, since the shape of the beta distribution is determined by two of the parameters (a and b) which are in turn determined by the values of the mode and the uncertainty coefficient, it is possible to get a reasonable

^{10/} It has been proposed that the analyst assume that the range of the cost variable is equal to six standard deviations, yielding $V(x^*) = [(H - L)/6]^2$. The basis for this assumption is that "most" of the probability associated with distributions such as the normal distribution is contained in the interval ± 3 standard deviations from the mean. The authors of this paper do not find this to be a very strong motivation. The total range of the uniform pdf is contained in an interval of ± 1.75 standard deviations and this distribution is a limiting form of the beta distribution. Further, the significance of the behavior of infinite, symmetric distributions such as the normal pdf in deriving properties of the (generally) finite, asymmetric beta distribution is questionable. It seems more logical to allow the analyst to input his knowledge of the uncertainty via the uncertainty coefficient. However, if the user prefers this device he should input the value $U = 0.577$ when using program SPET.

idea of the distributional form from a set of normalized graphs of beta pdfs with different modes and uncertainty coefficients. Such a set appears in Appendix B.

Appendix B contains ten groups of graphs of normalized beta pdfs. Each group contains graphs of three pdfs with the same mode but distinct uncertainty coefficients. The modal values vary from set to set, beginning with .05 and increasing by .05 until .50 when reading the abscissa from left to right, or beginning with .95 and decreasing by .05 until .50 when reading from right to left. To use Appendix B, the analyst computes the estimated normalized mode (N) from his estimates of L, H, and M with the relationship

$$N = (M - L) / (H - L) \quad [7]$$

and then selects the group of pdfs whose normalized mode is closest to this computed N. From the three graphs in this subset the analyst can see how the pdf varies with the uncertainty coefficient and get a reasonable idea of the shape of the distribution he has chosen. Alternatively, he may look at the set before choosing the uncertainty coefficient and use the information he gains to help him select the value for U.

Example

Assume an analyst is studying the cost of some subsystem "S". By analogy with other systems, or by some other technique, he determines that the cost of S will be something greater than \$7,000 but less than \$12,000. Further, utilizing a CER, or by some other technique, he estimates its most likely value at \$10,500. He calculates the normalized mode "N" from the equation:

$$N = \frac{M - L}{H - L} = \frac{10,500 - 7,000}{12,500 - 7,000} = .64$$

The set of graphs corresponding to this system is found between pages B-19 and B-21. These graphs are read from right to left.

After the analyst has developed the distributions of each subsystem he faces the problem of determining the distribution of their sum (the total cost of the system).

COMBINING SUBJECTIVE PROBABILITY DISTRIBUTIONS

Of the several techniques that have been employed by cost analysts to combine statistical distributions representing their subjective probability distributions, two of the more popular are derivation of moments^{11/} and Monte Carlo simulation.^{12/} Each of these techniques has its advantages: the former can be done with tables and a desk calculator, whereas the latter requires access to an electronic computer. The latter, however, is much faster and easier to use. For this reason Monte Carlo simulation is the technique employed in this research. Appendix C documents Program SPET, a computer program for adding independent statistical distributions by means of Monte Carlo simulation (SPET also performs other calculations discussed in the next section).

Program SPET has been used successfully on an interactive time-sharing computer system. Basically the user enters the parameters of the statistical distributions selected by the analyst and the program generates frequency distributions and summary statistics of the total system cost. Details and an example of the inputs, outputs, and operation of the program can be found in Appendix C.

THE INDEPENDENCE ASSUMPTION: A PROBLEM

When the analyst decomposes the system he is costing into several subsystems, it is very unlikely that the costs of the various subsystems are always statistically independent of one another. For example, the cost of the propulsion system of a guided missile is probably correlated with the cost of its payload. The correlation would be positive if an increase in payload cost was due to an increase in payload size which, in turn, would require a more powerful and hence more costly propulsion system. On the other hand, the correlation would be negative if an increase in payload cost was due to a reduction in payload

^{11/} See S. Sobel, A Computerized Technique to Express Uncertainty in Advanced System Cost Estimates (Bedford, MA: The Mitre Corp., TM-3728, 1963), and W. H. Sutherland, A Method for Combining Asymmetric Three-Value Predictions of Time or Cost (McLean, VA: Research Analysis Corp., RAC-P-65, 1972).

^{12/} P. F. Dienemann, op. cit., and D. F. Schoefer, op. cit.

size brought about by miniaturization which, in turn, would require a less powerful and hence less costly propulsion system. It is impossible to determine a priori whether the correlation between the costs of any two subsystems is positive or negative, but the experience of the authors suggests that in the majority of cases it will be positive.

If an analyst assumes that the statistical distributions (random variables) representing subsystem costs are independent when in fact they are positively correlated, he will underestimate the variance of their sum. To see this, consider the expression for the sum (S) of the variance of n random variables (X_i):

$$V(S) = \sum_{i=1}^n V(X_i) + 2 \sum_{i=1}^{n-1} \sum_{j=i+1}^n \text{Cov}(X_i, X_j) \quad [8]$$

The assumption of independence implies that $\text{Cov}(X_i, X_j) = 0$ for all $i \neq j$, which in turn implies that the expression

$$2 \sum_{i=1}^{n-1} \sum_{j=i+1}^n \text{Cov}(X_i, X_j) = 0.$$

If the X_i are dependent, this term could be positive, negative, or zero. If it is positive, its deletion from [8] by assuming independence results in an understatement of $V(S)$. An understatement of $V(S)$ could be a serious problem because it results in a confidence interval about the mean of the total system cost distribution that is smaller than it should be. This might cause the decision maker to posit unwarranted confidence in the estimate.

Assuming the random variables representing subsystem costs are positively correlated, the magnitude of the underestimate of $V(S)$ is directly related to the number of variables in the sum. To illustrate this fact, consider the following example:

There are two positively correlated random variables, X_1 and X_2 . Then

$$V(S) = V(X_1) + V(X_2) + 2\text{Cov}(X_1, X_2) \quad [9]$$

Now assume $X_1 = Z_1 + Z_2$. Then

$$\begin{aligned}
 V(S) &= V(Z_1) + V(Z_2) + V(X_2) + 2 \text{Cov}(Z_1, Z_2) + \\
 &\quad 2 \text{Cov}(Z_1, X_2) + 2 \text{Cov}(Z_2, X_2) \\
 &= V(Z_1) + V(Z_2) + V(X_2) + 2 \text{Cov}(X_1, X_2) + \\
 &\quad 2 \text{Cov}(Z_1, Z_2)
 \end{aligned}
 \tag{10}$$

Assuming independence in the case of two variables results in the deletion of $2 \text{Cov}(X_1, X_2)$ from $V(S)$. However, in the case of three variables the assumption of independence results in the deletion of not only $2 \text{Cov}(X_1, X_2)$ but $2 \text{Cov}(Z_1, X_2)$ as well. Clearly, if the covariances are positive, $V(S)$ is understated more in the case of three variables than in the case of two.

Therefore in the case where the random variables representing subsystem costs are positively correlated not only is the variance of the total system cost underestimated, but the magnitude of the underestimate is directly related to the number of variables making up the sum.

The obvious solution to this problem is to incorporate into computer Program SPET provisions for the consideration of probable correlation among the variables and then proceed to estimate the nature of the correlation. Although the former idea presents no problem, the latter appears to be a most difficult task. It is not clear at this point how to systematically estimate the correlation among the variables representing subsystem costs. Hopefully a credible technique for doing such will become apparent to someone.

Program SPET has been designed to provide some insight into the significance of the independence assumption in two ways. First, by using the same random number in sampling from all the subsystem distributions, a distribution of total cost which the printout titles "Dependent Beta" is derived. The technique imposes a functional relationship between all the variables. Note that this functional relationship is not an arbitrary relationship but is imposed by the forms of the subsystem pdfs developed by the analyst. Specifically, if $F_i(X_i)$ is the cumulative distribution function of the i th subsystem then:

$$\sum_{i=1}^n X_i = X_1 + \sum_{i=2}^n F_i^{-1}[F_1(X_1)]
 \tag{11}$$

where F_i^{-1} is the functional inverse of F_i . The functions F_i and F_i^{-1} are much too complex to derive the specific form of the relationship imposed but the technique does impose a very real positive correlation between the variables.

As a second means of examining uncertainty without imposing the independence assumption, SPET prints the statistics for a uniform distribution of total cost on the interval between the sum of the minimum costs of the subsystems and the sum of the maximum cost of the subsystems. This is meant to serve as a "worst case." SPET also prints the total cost distribution assuming each of the subsystem's costs is uniformly and independently distributed.

SUMMARY

Several conceptual points have been discussed, among them:

- (1) the nature of uncertainty in cost analysis;
- (2) the value of treating uncertainty explicitly in cost analysis;
- (3) a problem inherent in using the classical linear regression model as a basis for statements on cost uncertainty;
- (4) the subjective nature of cost uncertainty;
- (5) the properties of the beta distribution and how they can be used to facilitate cost uncertainty analysis; and
- (6) the dependence of subsystem costs and its impact on statements about uncertainty.

The basic practical contribution of this paper is a computer program for generating statements on cost uncertainty that permits the analyst to input any imaginable probability distribution to represent a subsystem cost.

APPENDIX A

ESTIMATION OF PARAMETERS a AND b:

COMPUTER PROGRAM PARAM

Introduction

One way an analyst can estimate parameters a and b of the generalized beta probability density function (pdf)^{1/} is to simultaneously solve two functionally independent equations in a and b . Consider the mode -- the "most likely" value -- of the generalized beta pdf:

$$M(x^*) = (aH + bL)/(a + b)^{2/} \quad [A1]$$

Using a cost-estimating relationship, or other methods, the analyst can estimate the mode (M) of the subsystem cost he wants to represent with a beta pdf. Substitution of M into equation [A1], along with L and H , yields one equation in a and b :

$$M = (aH + bL)/(a + b) \quad [A7]$$

Another equation in a and b can be obtained from an estimate of the probability (P) that the cost of the sub-

1/ See equation [2] on page 8.

2/ This expression can be obtained by solving $d[g(x^*)]/dx^* = 0$ for x^* , where $g(x^*)$ is given by equation [2]. However, it is easier to obtain expressions for the mode as well as the mean, $E(x^*)$, and the variance, $V(x^*)$, of the generalized beta pdf by means of simple algebraic operations on the expressions for these statistics derived from the more familiar normalized beta pdf. The mode, mean, and variance of the normalized beta pdf are:

$$M(x) = a/(a + b) \quad [A2]$$

$$E(x) = (a + 1)/(a + b + 2) \quad [A3]$$

$$V(x) = [(a + 1)(b + 1)]/[(a + b + 2)^2(a + b + 3)] \quad [A4]$$

Note that $M(x) = a/(a + b) = M[(x^* - L)/(H - L)] = [M(x^*) - L]/(H - L)$; solving for $M(x^*)$ yields equation [A1]. Proceeding similarly for $E(x^*)$ and $V(x^*)$ yields

$$E(x^*) = (aH + bL + H + L)/(a + b + 2) \quad [A5]$$

$$V(x^*) = [(a + 1)(b + 1)(H - L)^2]/[(a + b + 2)^2(a + b + 3)] \quad [A6]$$

system will lie within a subinterval of its range [L, H]. For convenience, this subinterval is taken as the interval from L to the midpoint between L and M, i.e., $[(L + M)/2]$. Then an equation in a and b results from the relationship

$$P = [C/(H - L)^{a+b+1}] \int_L^{\frac{L+M}{2}} (x^* - L)^a (H - x^*)^b dx^* \quad [A8]$$

where C, x, a, and b are defined in equation [2].

Equations [A7] and [A8] comprise two functionally independent equations in a and b, and when solved simultaneously determine unique values (estimates) for a and b.

Program Description

Program PARAM is written in FORTRAN IV for use in conjunction with a PDP-10 computer in an interactive time-sharing mode. It is designed to solve equations [A7] and [A8] for parameters a and b, given values for L, H, M, and P. This is accomplished in the following sequence:

1 - The inputs L, H, and M are normalized. Denoting the normalized counterparts of L, H, and M by l, h, and m, respectively,

$$\begin{aligned} l &= 0 \\ h &= 1 \\ m &= (M - L)/(H - L) \end{aligned}$$

2 - A standard root-finding technique^{3/} is employed in Subroutine Beta to find the values for a and b that satisfy

$$P = C \int_0^{m/2} x^a (1 - x)^b dx \quad [A9]$$

where P is supplied by the user, m is determined from a user supplied datum (M), and C, x, a, and b are as defined in equation [2]. Note that [A9] is the normalized version of [A8].

^{3/} The root-finding technique is known as the "false position" method and is found in most texts on numerical analysis.

Subroutine Beta calls Function Subprogram Gamma to compute values for $\Gamma(n)$. This is accomplished by use of the relation

$$\Gamma(n+1) = n\Gamma(n) \quad [A10]$$

and an interpolation procedure.

By repeated application of [A10], $\Gamma(n)$ for any $n > 0$ can be expressed as the product

$$(n-1)(n-2)\dots\Gamma(n^*) \quad [A11]$$

where $1 \leq n^* \leq 2$. Since $\Gamma(n)$ is well-behaved in the range $1 \leq n^* \leq 2$, the values of $\Gamma(n^*)$ can be approximated using a "table look-up" interpolation procedure. Function Subprogram Gamma uses such an interpolation device in conjunction with relation [A11] to evaluate $\Gamma(n)$.

3 - The four-tuple (a,b,L,H) is printed as output.

User Instructions

The following example demonstrates the use of Program PARAM. Note that the user supplied portions of the example are underlined:

PROGRAM PARAM

ESTIMATES OF L,H,M,P
3000,14000,10500,.14

ALPHA IS 1.07
BETA IS 1.50
LOW IS 8000.00
HIGH IS 14000.00

CP UNITS 2

EXIT

Program Listing

```

00010      A=BB
00020      220  PP=C.
00030      NMID=0.
00040      RETURN
00050      50  WRITE(II,00)
00060      60  FORMAT(1H-,28HBETA PARAMETERS OUT OF RANGE)
00070      RETURN
00080      END
00090      C
00100      FUNCTION GAMMA(GP)
00110      REAL G(22)
00120      DATA G/1.,.9735,.95135,.93304,.91817,
00130      + .9064,.89747,.89115,.88726,.88565,
00140      + .88623,.88687,.89352,.90012,.90864,
00150      + .91936,.93138,.94561,.96177,.97983,1.,1./
00160      II=16
00170      JJ=16
00180      ERROR=0.
00190      SUM=1.
00200      IF(GP-57.4) 10,10,20
00210      10  IF(GP-1.00E-20) 20,30,30
00220      20  WRITE(II,40)
00230      40  FORMAT(1H-,28HGAMMA PARAMETER OUT OF RANGE)
00240      ERROR=ERROR+1.
00250      RETURN
00260      30  IF(GP-2.) 50,50,60
00270      60  SUM=SUM*(GP-1.)
00280      GP=GP-1.
00290      GO TO 30
00300      50  IF(GP-1.) 70,80,80
00310      70  SUM=SUM/GP
00320      GP=GP+1.
00330      80  I=(GP-1.)/.05+1.
00340      XI=1-I
00350      GPL=1.+XI*.05
00360      GAMFN=G(1)+(G(I+1)-G(I))*(GP-GPL)/.05
00370      GAMMA=GAMFN*SUM
00380      RETURN
00390      END

```

00410		IF (GP-57.4) 40,40,50
00420	40	C1=GAMMA(GP)
00430		GP=A+1.
00440		C2=GAMMA(GP)
00450		C3=C1/C2
00460		GP=B+1.
00470		C4=GAMMA(GP)
00480		C=C3/C4
00490	70	XMM=XMM+2.
00500		DO 80 I=2,3
00510		FT(I)=F(I)**A*(1.-T(I))**3
00520		T(I+1)=T(I)+SP
00530	80	CONTINUE
00540		T(2)=T(4)
00550		XSUM=FT(1)+4.*FT(2)+FT(3)
00560		FT(1)=FT(3)
00570		XINTEG=XINTEG+XSUM
00580		IF (XMM-INTEG) 70,90,90
00590	90	AA=P-C*SP/3.*XINTEG
00600		IF (ABS(AA)-.0001) 100,100,110
00610	110	IF (AA*AAS) 120,130,130
00620	130	IF (AA) 140,140,150
00630	140	AA1=AA
00640		GP1=B
00650		GO TO 160
00660	150	AA2=AA
00670		GP2=B
00680	160	GO TO (170,180),JJ1
00690	120	JJ1=2
00700		GO TO 130
00710	170	B=W
00720		W=W*2.
00730		GO TO 190
00740	180	B=(AA2*GP1-AA1*GP2)/(AA2-AA1)
00750	190	AAS=AA
00760		GO TO 30
00770	100	IF (PP-NMIDD) 200,200,210
00780	200	GO TO 220
00790	210	3B=B
00800		B=A

```

00010      REAL L,M
00020      II=16
00030      JJ=16
00040      WRITE(II,10)
00050      10  FORMAT(1H-,20,HESTIMATES OF L,H,M,P/)
00060      READ(JJ,20) L,H,M,P
00070      20  FORMAT(4F)
00080      CALL BETA(L,H,M,P,A,B)
00090      WRITE(II,30) A,B,L,H
00100      30  FORMAT(1H-,9HALPHA IS ,F8.2/10H BETA IS ,F8.2/
00110      +    10H LOW IS ,F8.2/10H HIGH IS ,F8.2)
00120      WRITE(II,40)
00130      40  FORMAT(1H-)
00140      END
00150      C
00160      SUBROUTINE BETA(L,H,M,P,A,B)
00170      REAL L,M,INTEG,NMODE,NMID,NMIDD,FT(3),F(4)
00180      II=16
00190      JJ=16
00200      INTEG=40.
00210      SEARCH=20.
00220      NMODE=(M-L)/(H-L)
00230      NMID=NMODE/2.
00240      IF(P-NMID) 10,10,20
00250      20  PP=P
00260      NMIDD=NMID
00270      P=1.-P
00280      NMID=1.-NMID
00290      NMODE=1.-NMODE
00300      10  SP=NMID/INTEG
00310      AAS=-1.
00320      JJI=1
00330      X=SEARCH*(1.-NMODE)
00340      B=0.
00350      30  A=B*NMODE/(1.-NMODE)
00360      XMM=0.
00370      XINTEG=0.
00380      F(2)=SP
00390      FT(1)=0.
00400      GP=A+B+2.

```

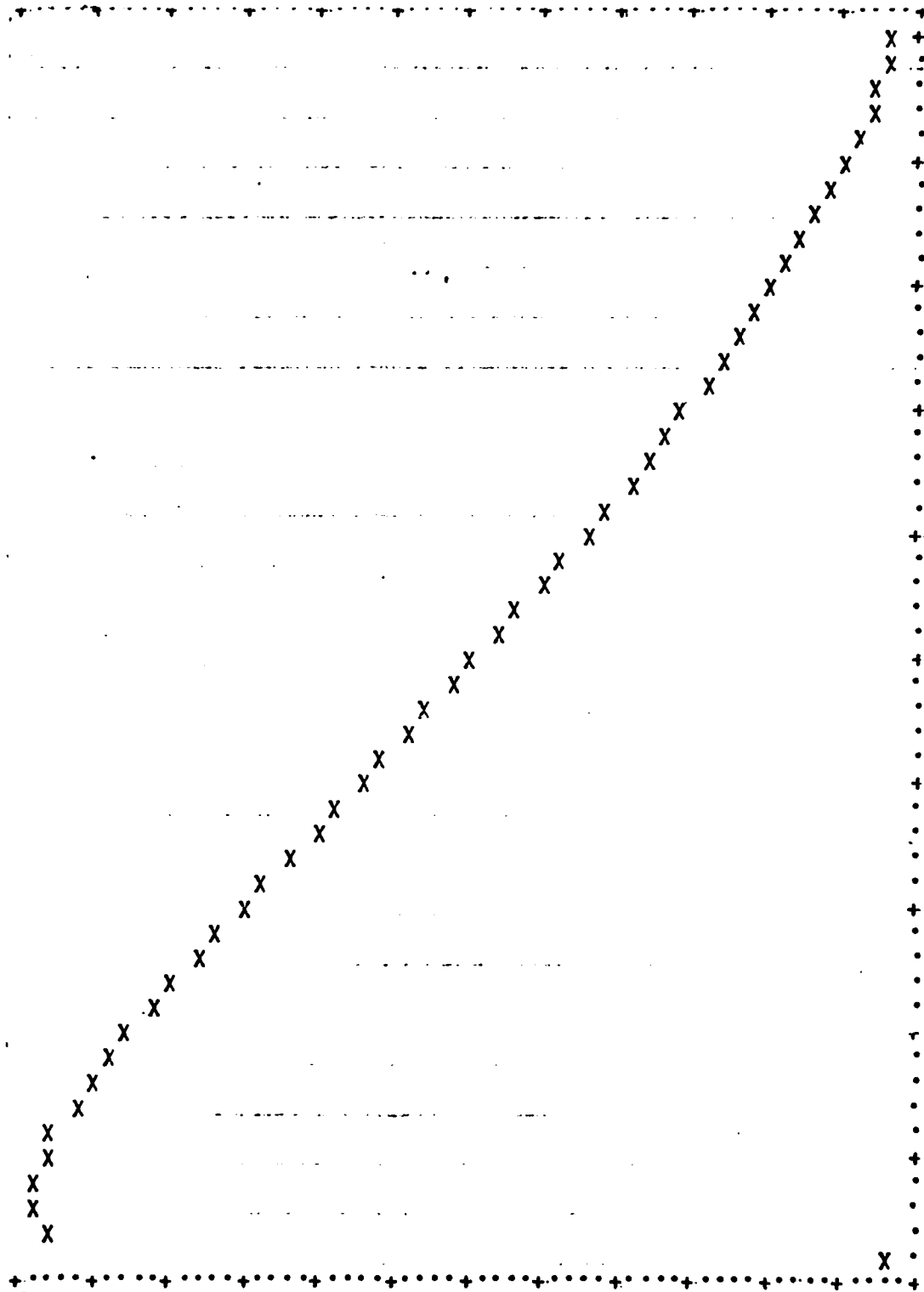
APPENDIX B

GRAPHICAL AIDS FOR SELECTING
THE UNCERTAINTY COEFFICIENT

LEFT-TO-RIGHT
Normalized Mode = .05

HIGH VARIANCE
Uncertainty Coefficient = .75

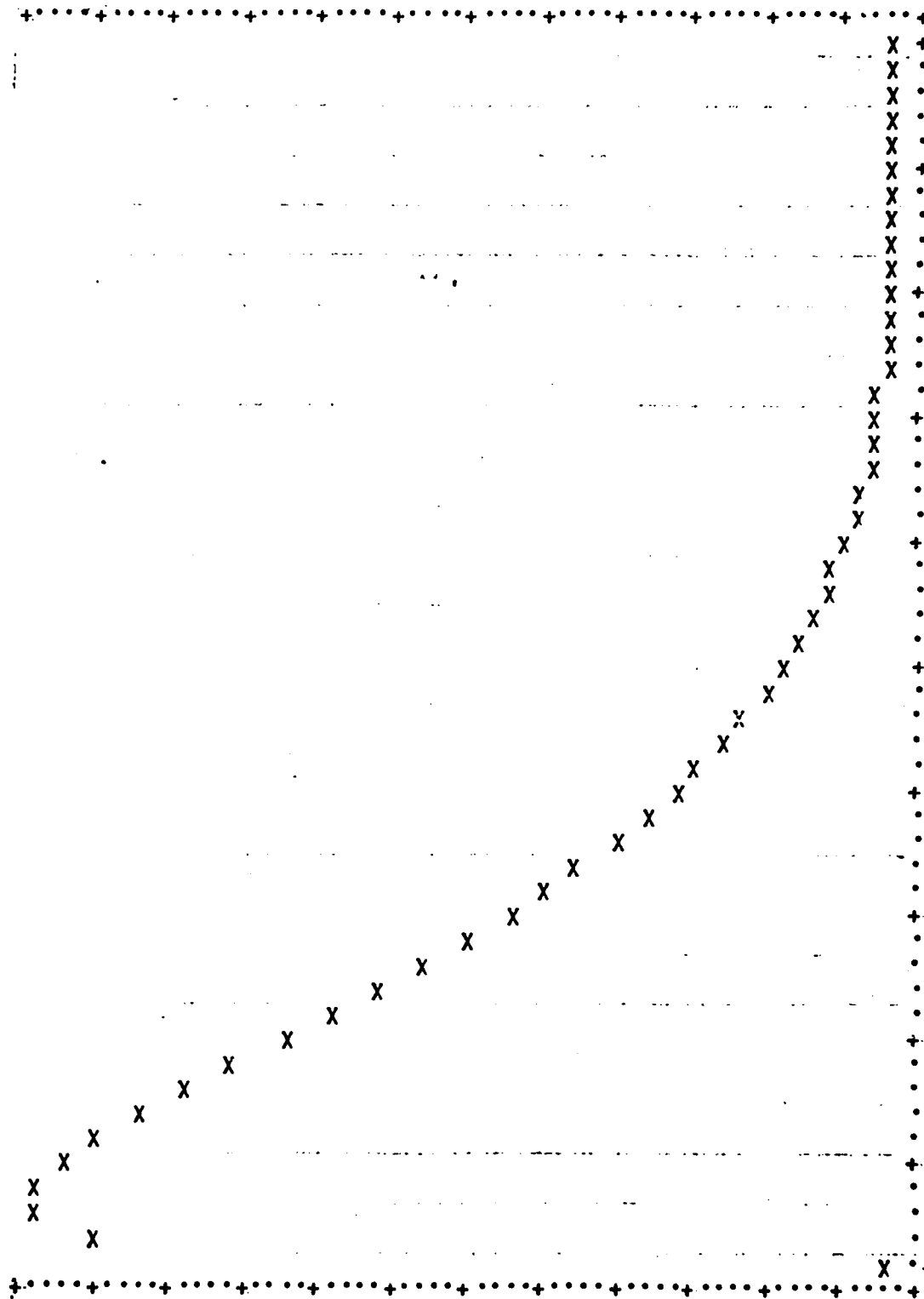
RIGHT-TO-LEFT
Normalized Mode = .95



LEFT-TO-RIGHT
Normalized Mode = .05

MEDIUM VARIANCE
Uncertainty Coefficient = .50

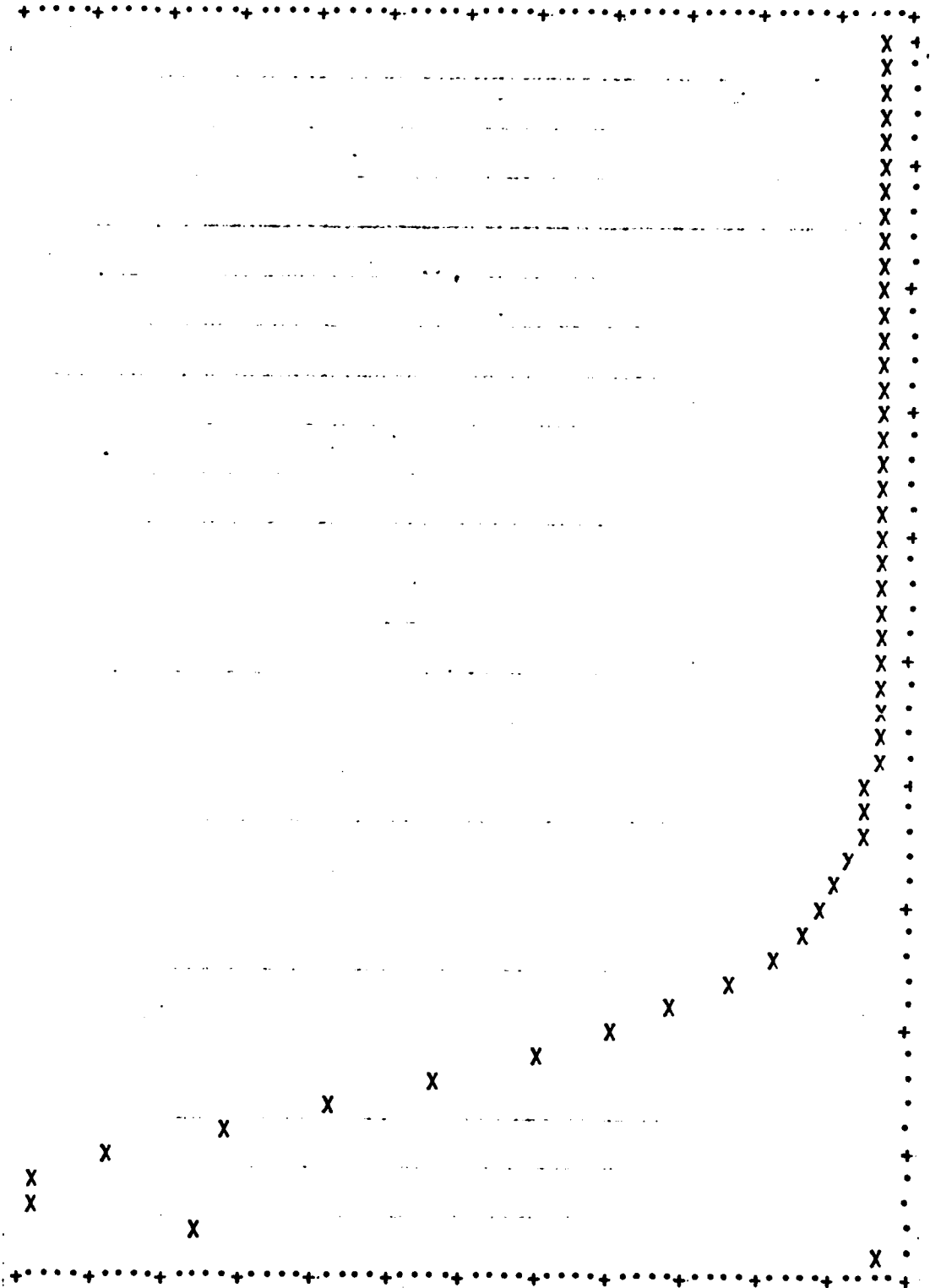
RIGHT-TO-LEFT
Normalized Mode = .95



LEFT-TO-RIGHT
Normalized Mode = .05

LOW VARIANCE
Uncertainty Coefficient = .25

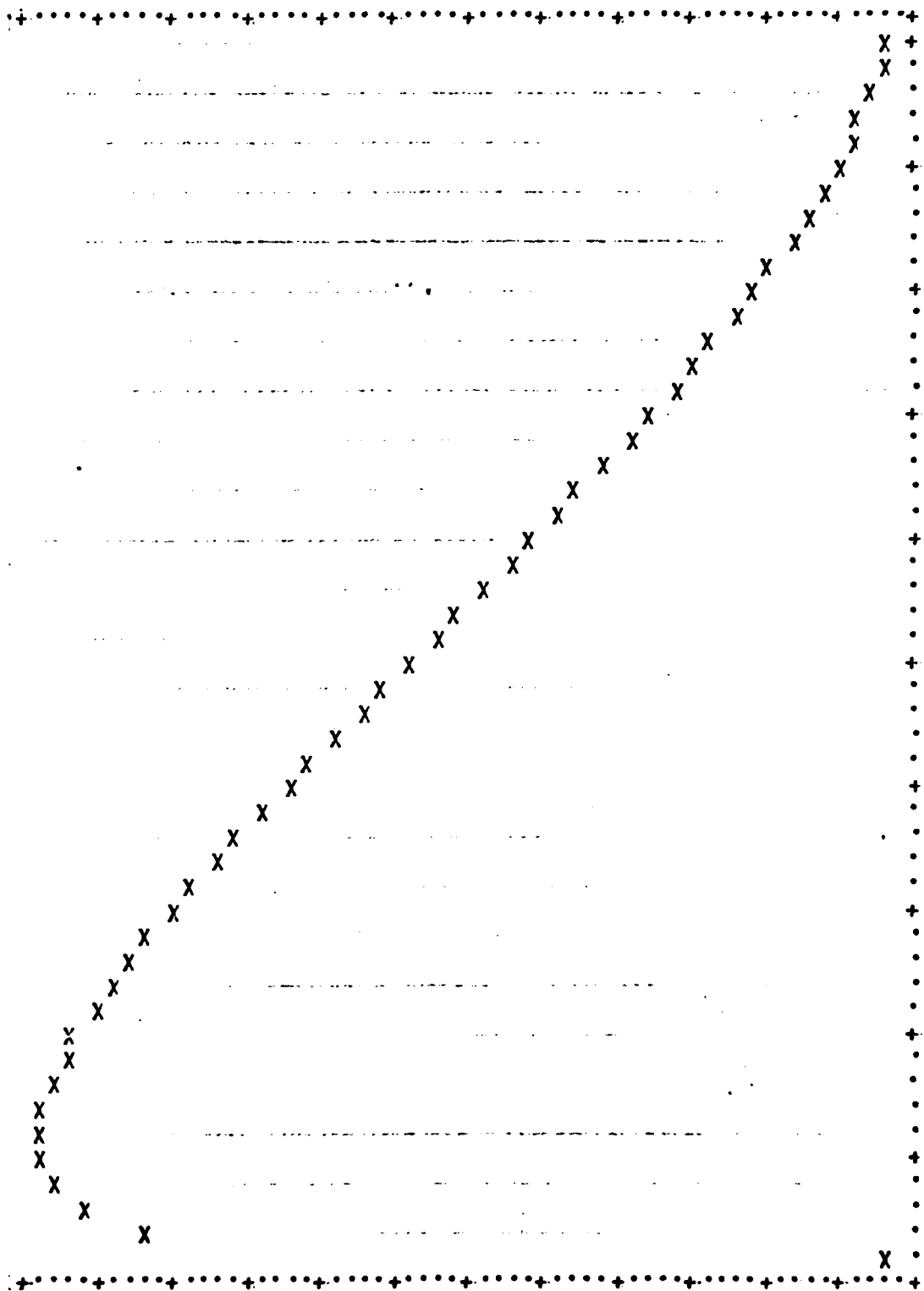
RIGHT-TO-LEFT
Normalized Mode = .95



LEFT-TO-RIGHT
Normalized Mode = .10

HIGH VARIANCE
Uncertainty Coefficient = .75

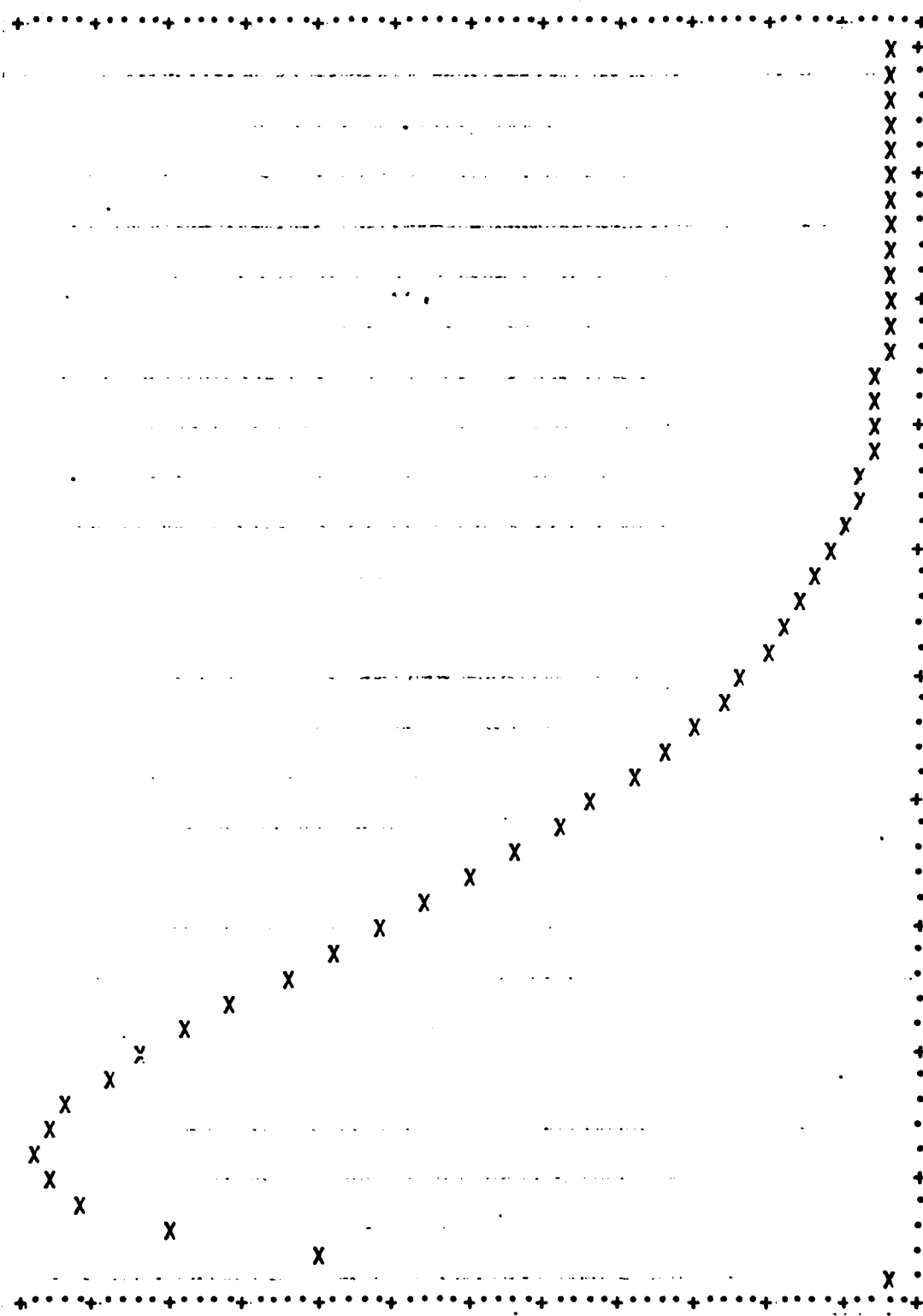
RIGHT-TO-LEFT
Normalized Mode = .90



RIGHT-TO-LEFT
Normalized Mode = .90

MEDIUM VARIANCE
Uncertainty Coefficient = .50

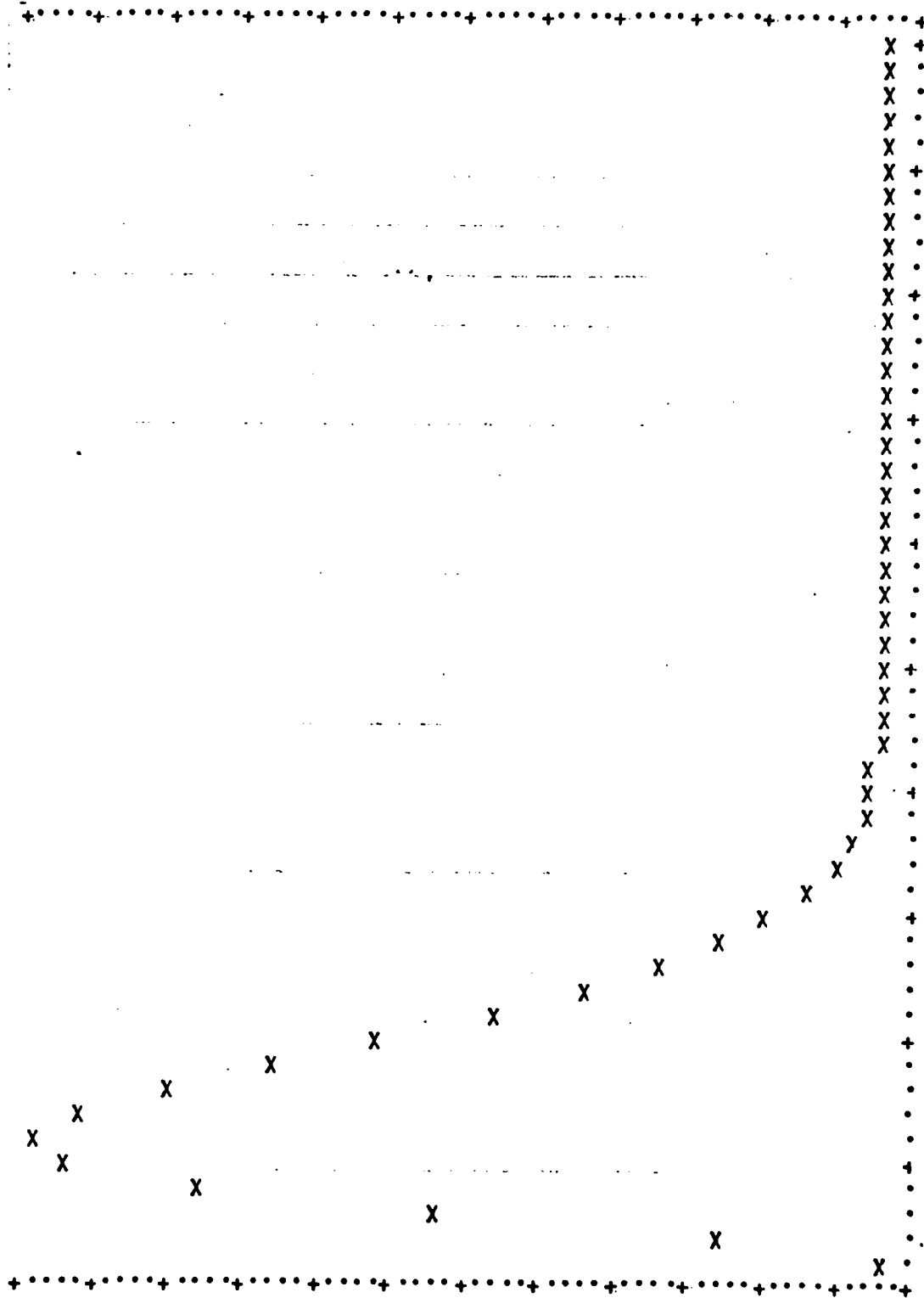
LEFT-TO-RIGHT
Normalized Mode - .10



LEFT-TO-RIGHT
Normalized Mode = .10

LOW VARIANCE
Uncertainty Coefficient = .25

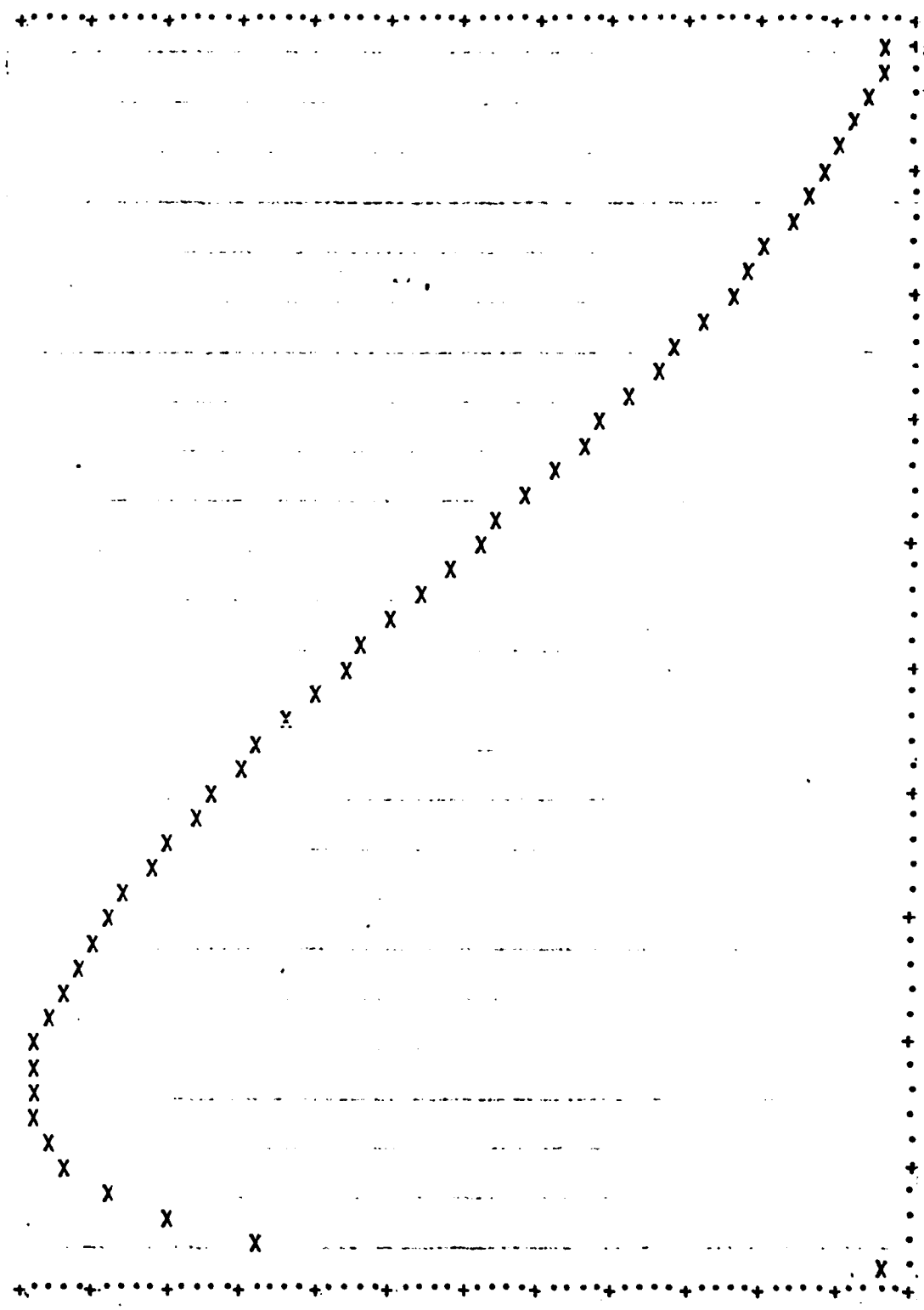
RIGHT-TO-LEFT
Normalized Mode = .90



LEFT-TO-RIGHT
Normalized Mode = .15

HIGH VARIANCE
Uncertainty Coefficient = .75

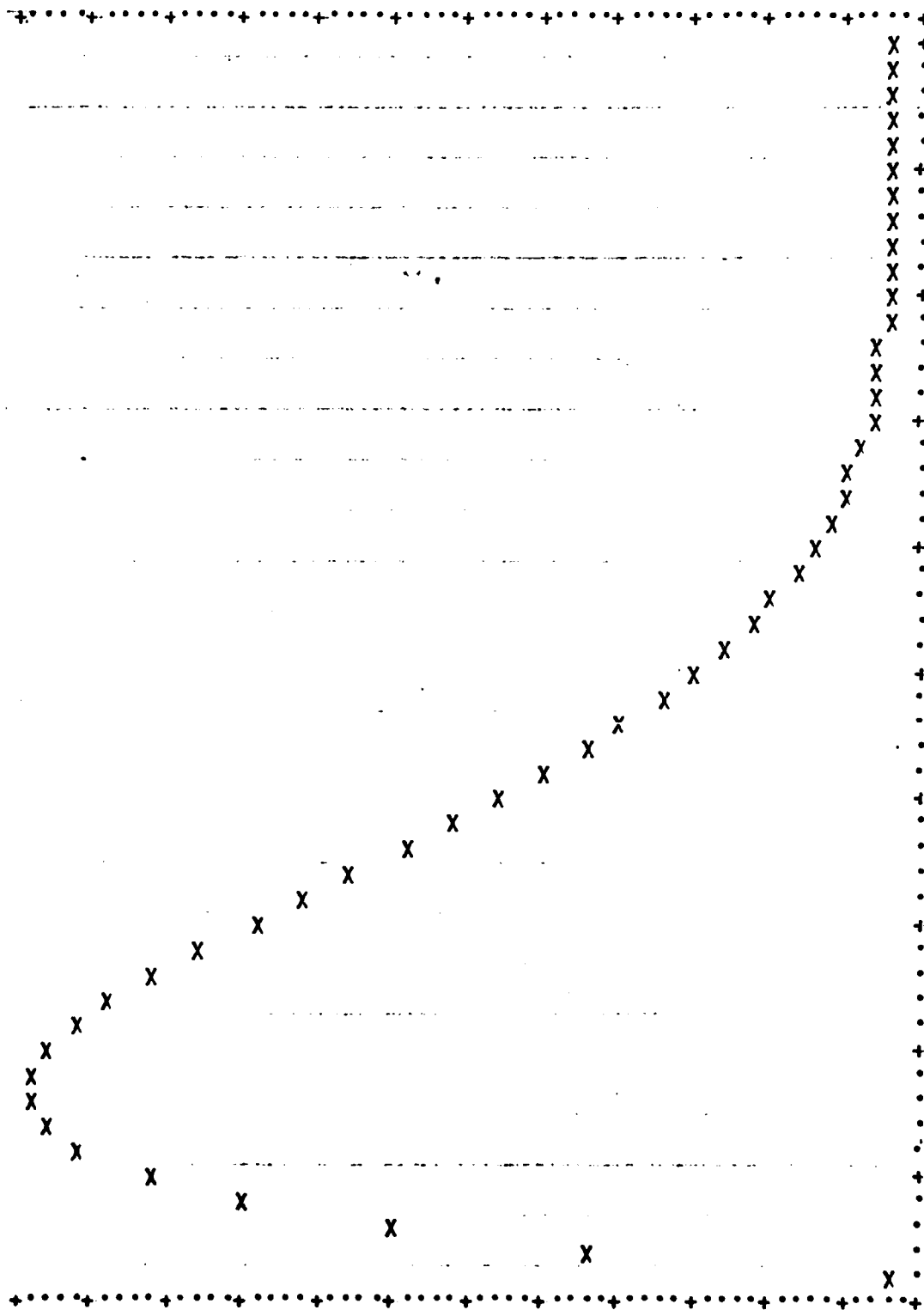
RIGHT-TO-LEFT
Normalized Mode = .85



LEFT-TO-RIGHT
Normalized Mode = .15

MEDIUM VARIANCE
Uncertainty Coefficient = .50

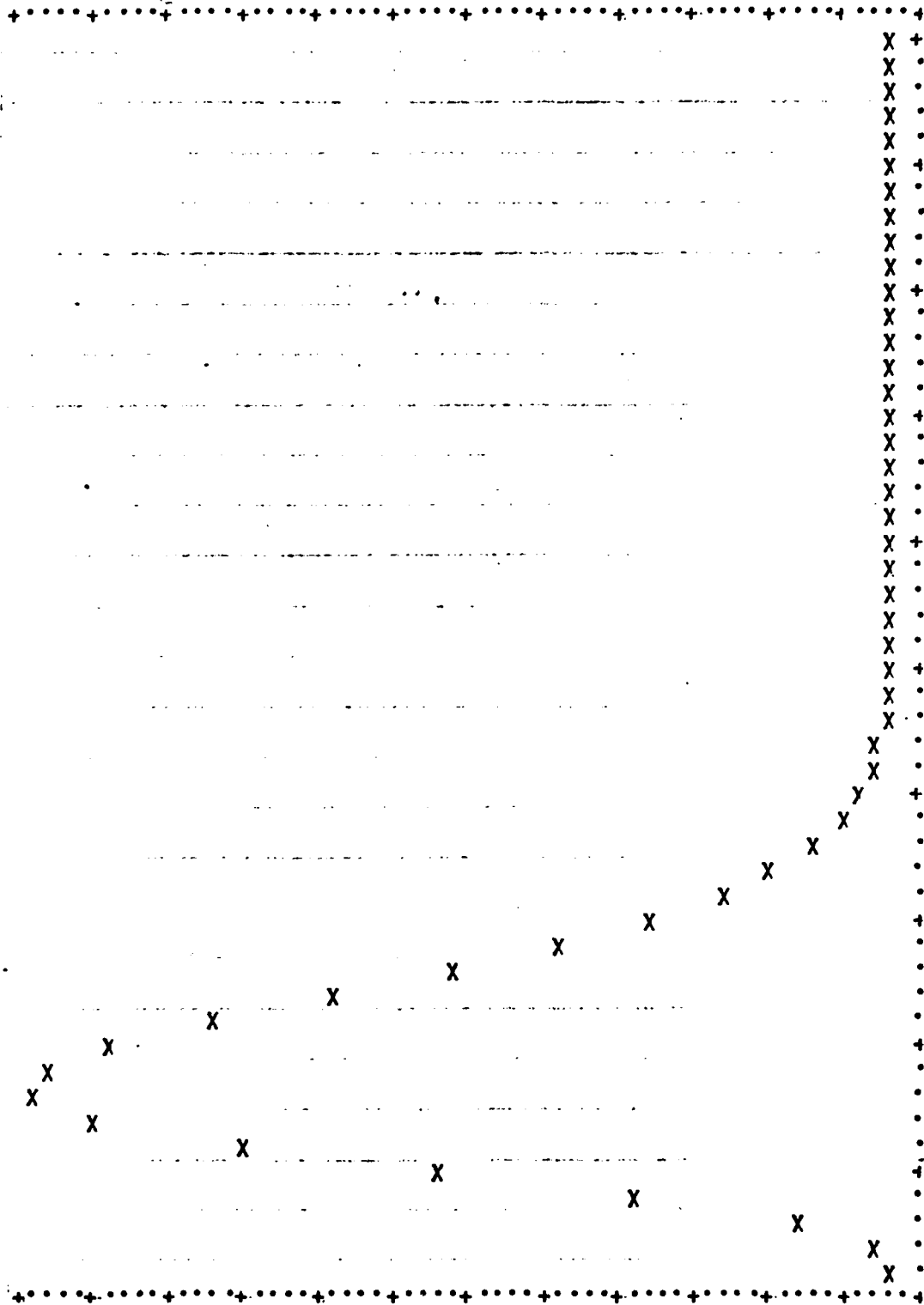
RIGHT-TO-LEFT
Normalized Mode = .85



RIGHT-TO-LEFT
Normalized Mode = .85

LOW VARIANCE
Uncertainty Coefficient = .25

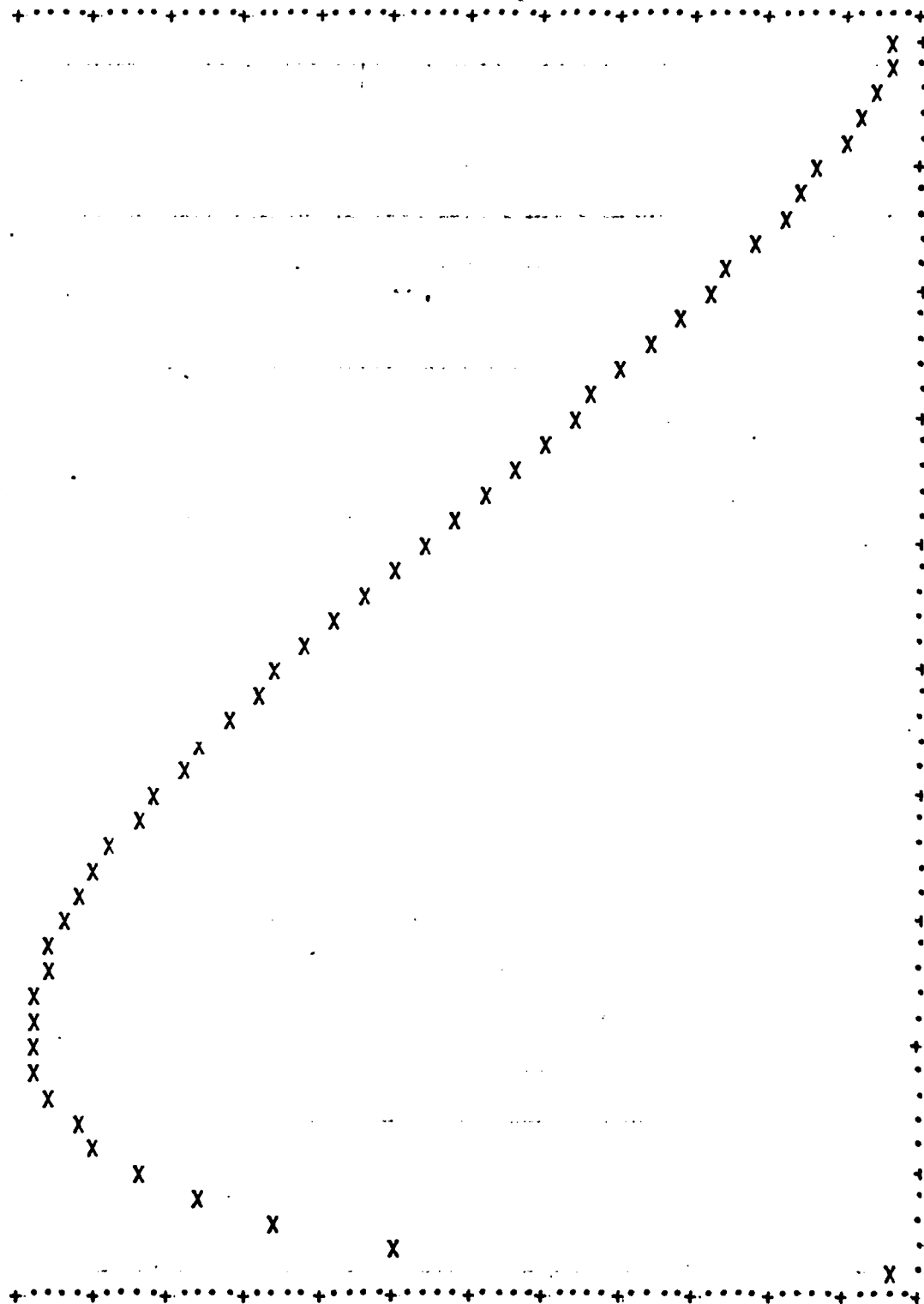
LEFT-TO-RIGHT
Normalized Mode = .15



RIGHT-TO-LEFT
Normalized Mode = .80

HIGH VARIANCE
Uncertainty Coefficient = .75

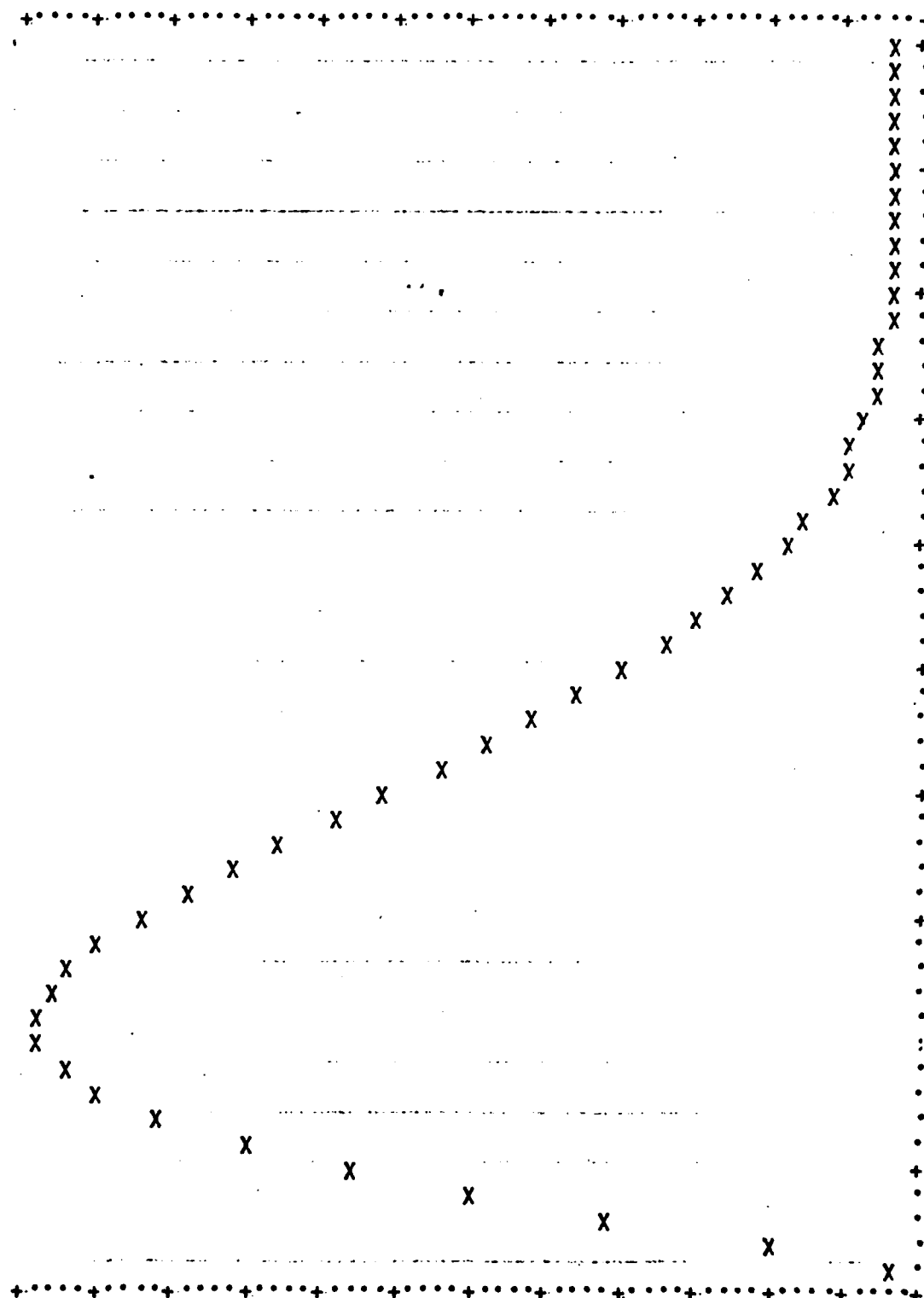
LEFT-TO-RIGHT
Normalized Mode = .20



RIGHT-TO-LEFT
Normalized Mode = .80

MEDIUM VARIANCE
Uncertainty Coefficient = .50

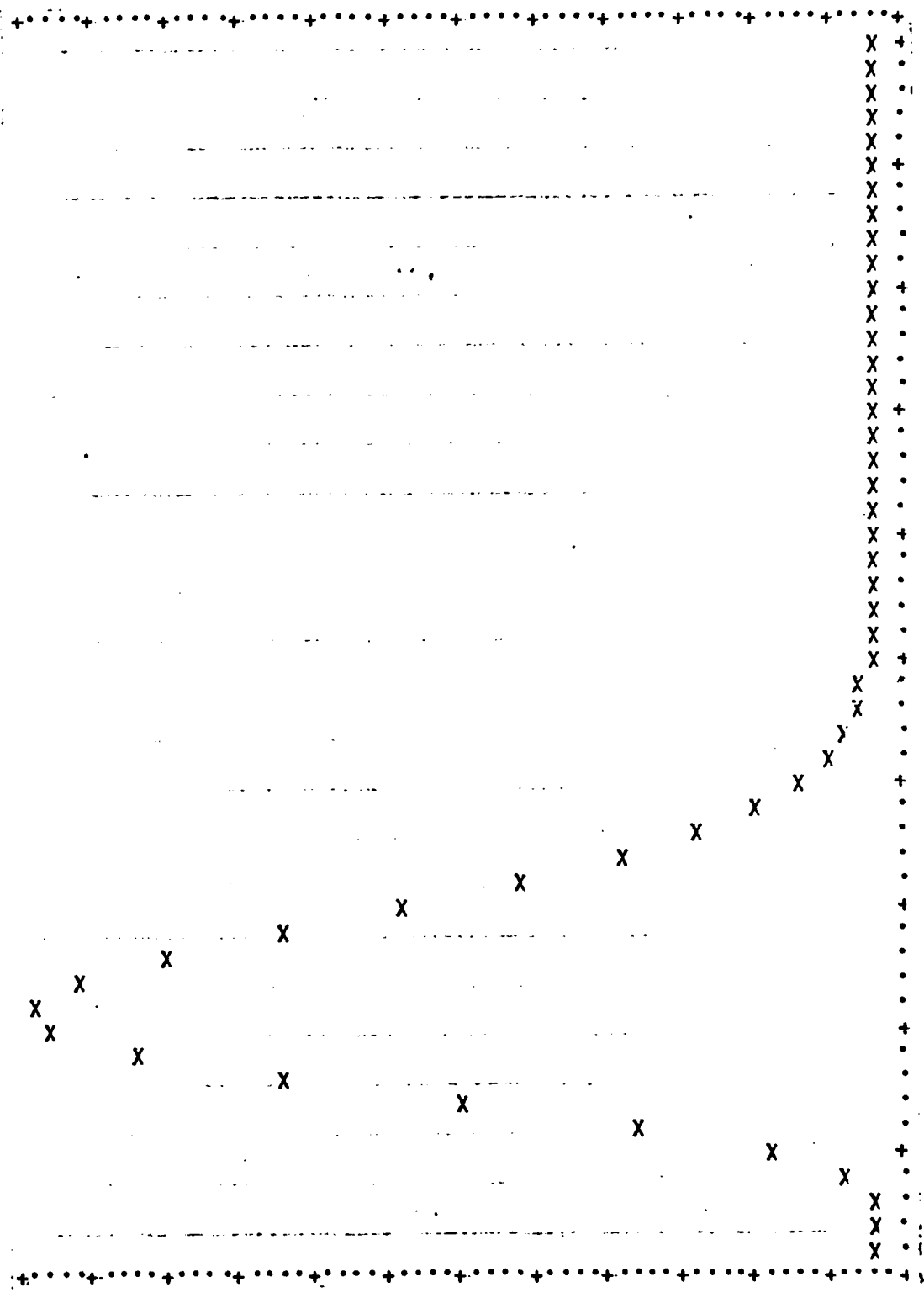
LEFT-TO-RIGHT
Normalized Mode = .20



RIGHT-TO-LEFT
Normalized Mode = .80

LOW VARIANCE
Uncertainty Coefficient = .25

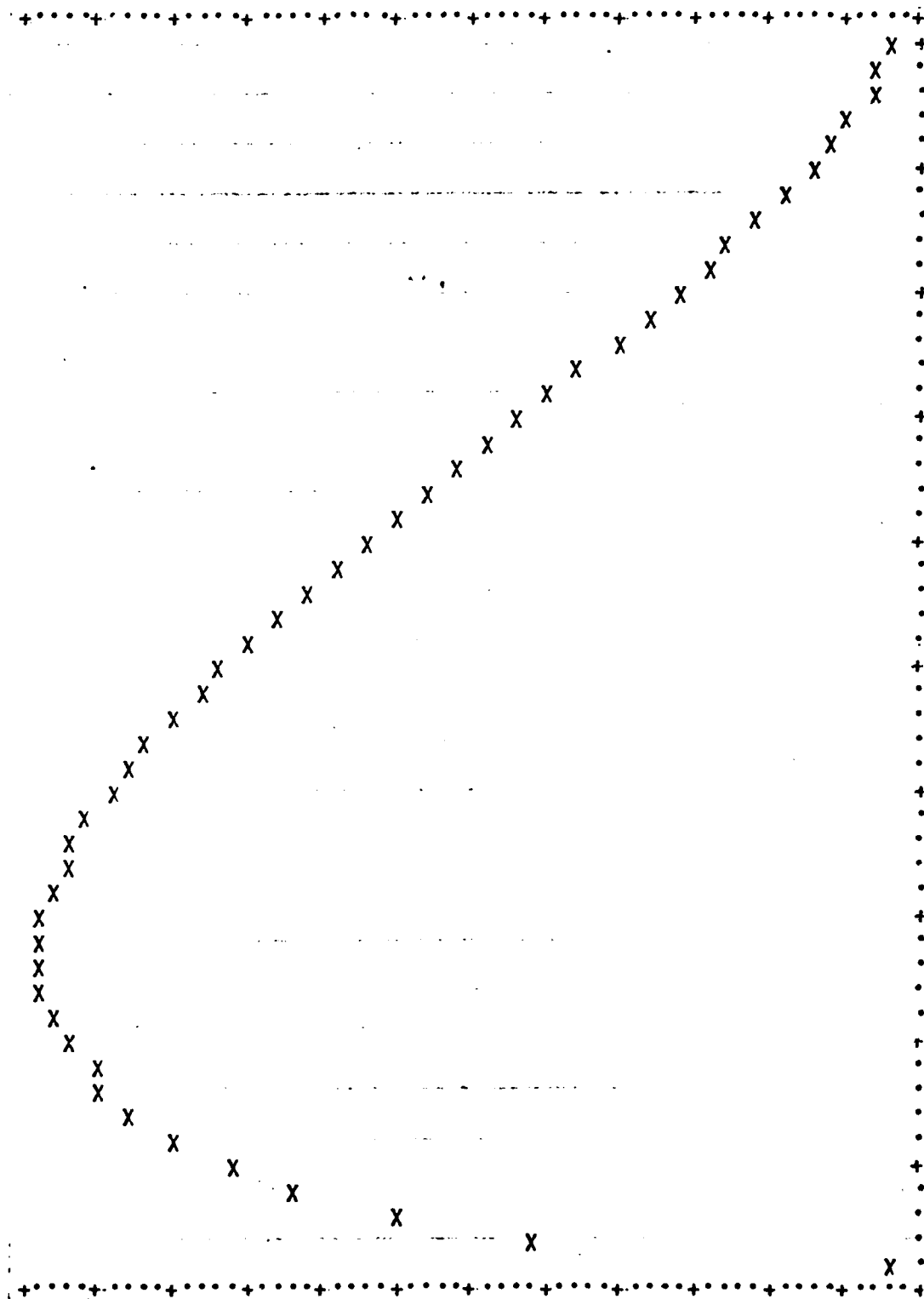
LEFT-TO-RIGHT
Normalized Mode = .20



RIGHT-TO-LEFT
Normalized Mode = .75

HIGH VARIANCE
Uncertainty Coefficient = .75

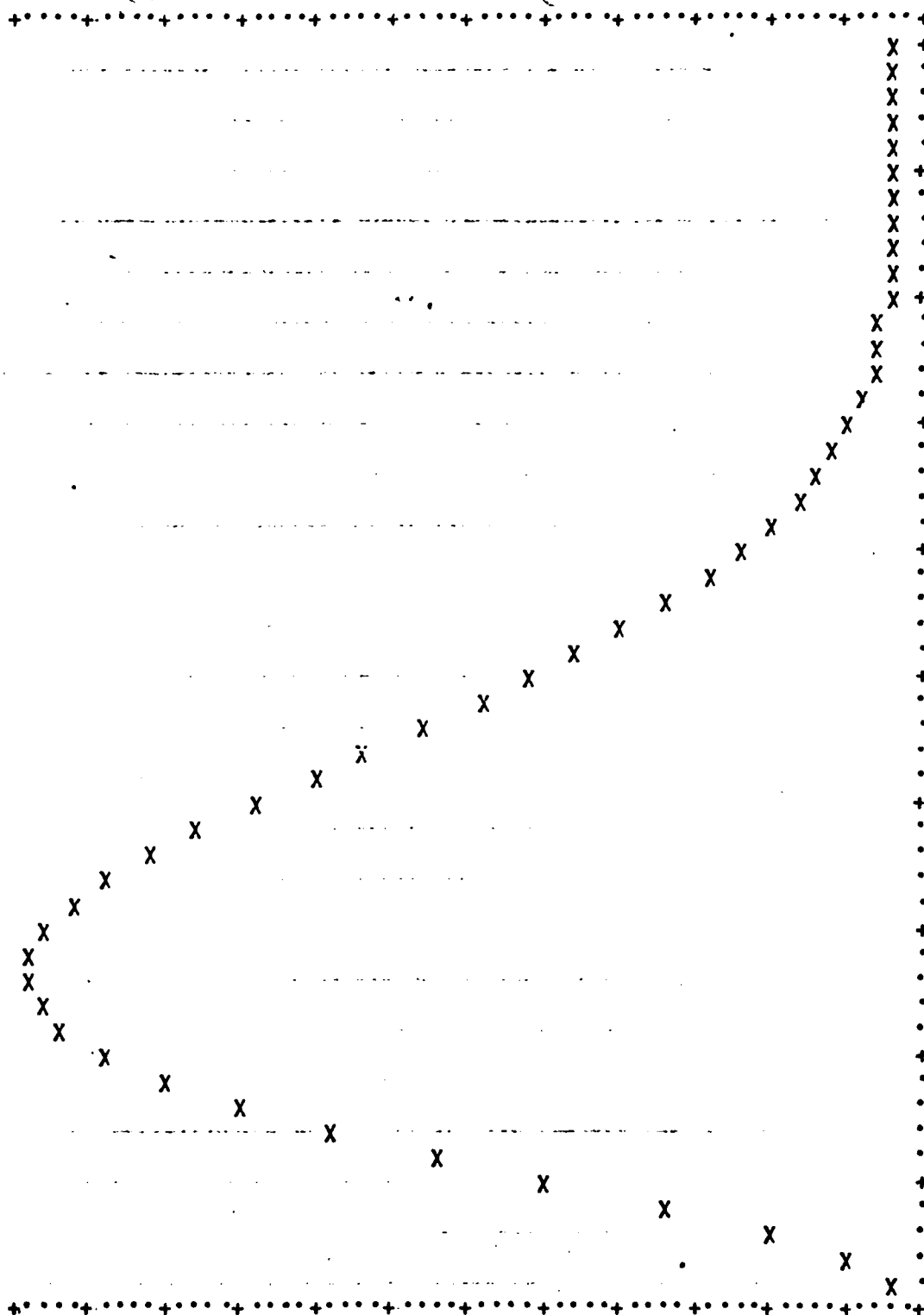
LEFT-TO-RIGHT
Normalized Mode = .25



RIGHT-TO-LEFT
Normalized Mode = .75

MEDIUM VARIANCE
Uncertainty Coefficient = .50

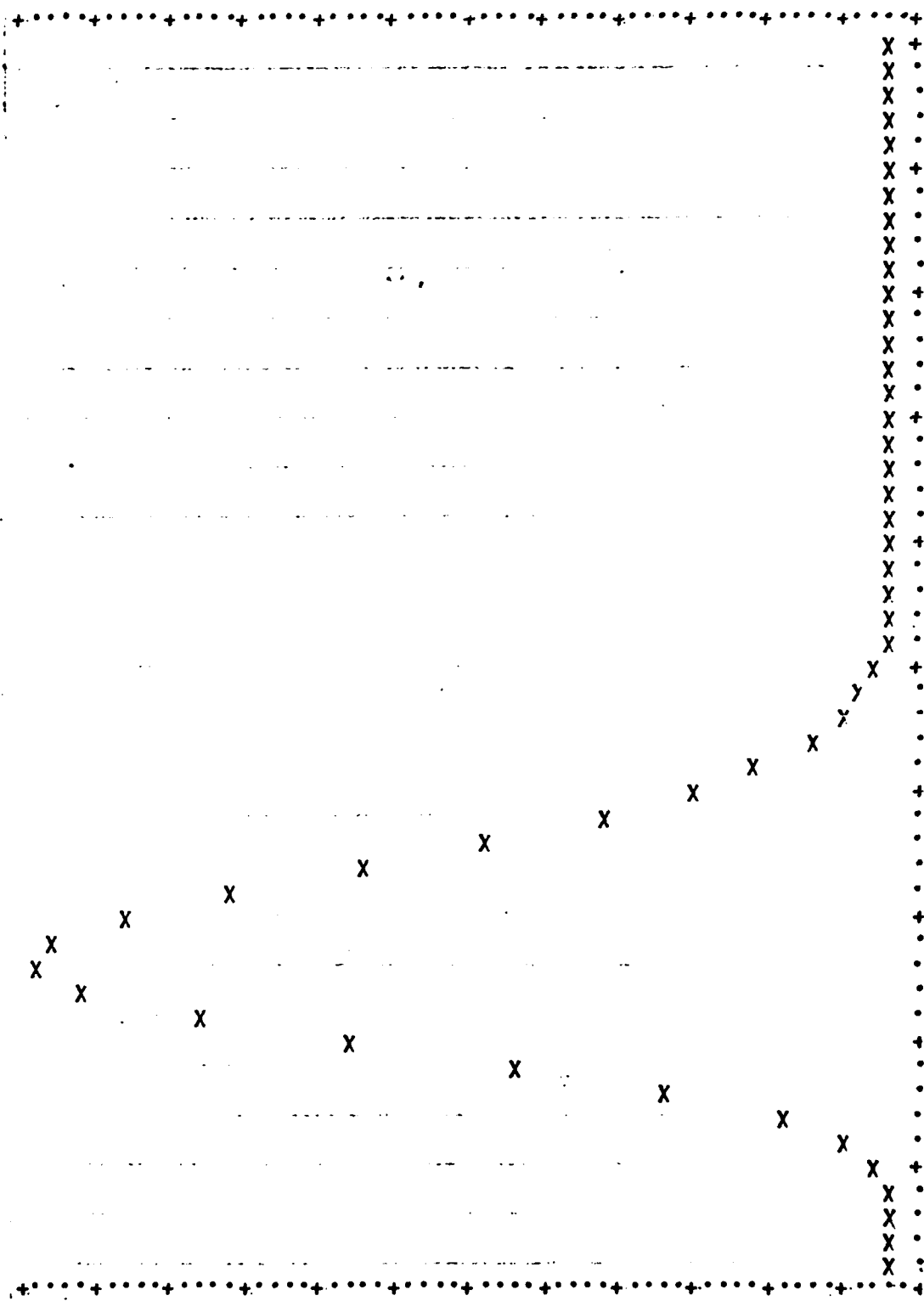
LEFT-TO-RIGHT
Normalized Mode = .25



RIGHT-TO-LEFT
Normalized Mode = .75

LOW VARIANCE
Uncertainty Coefficient = .225

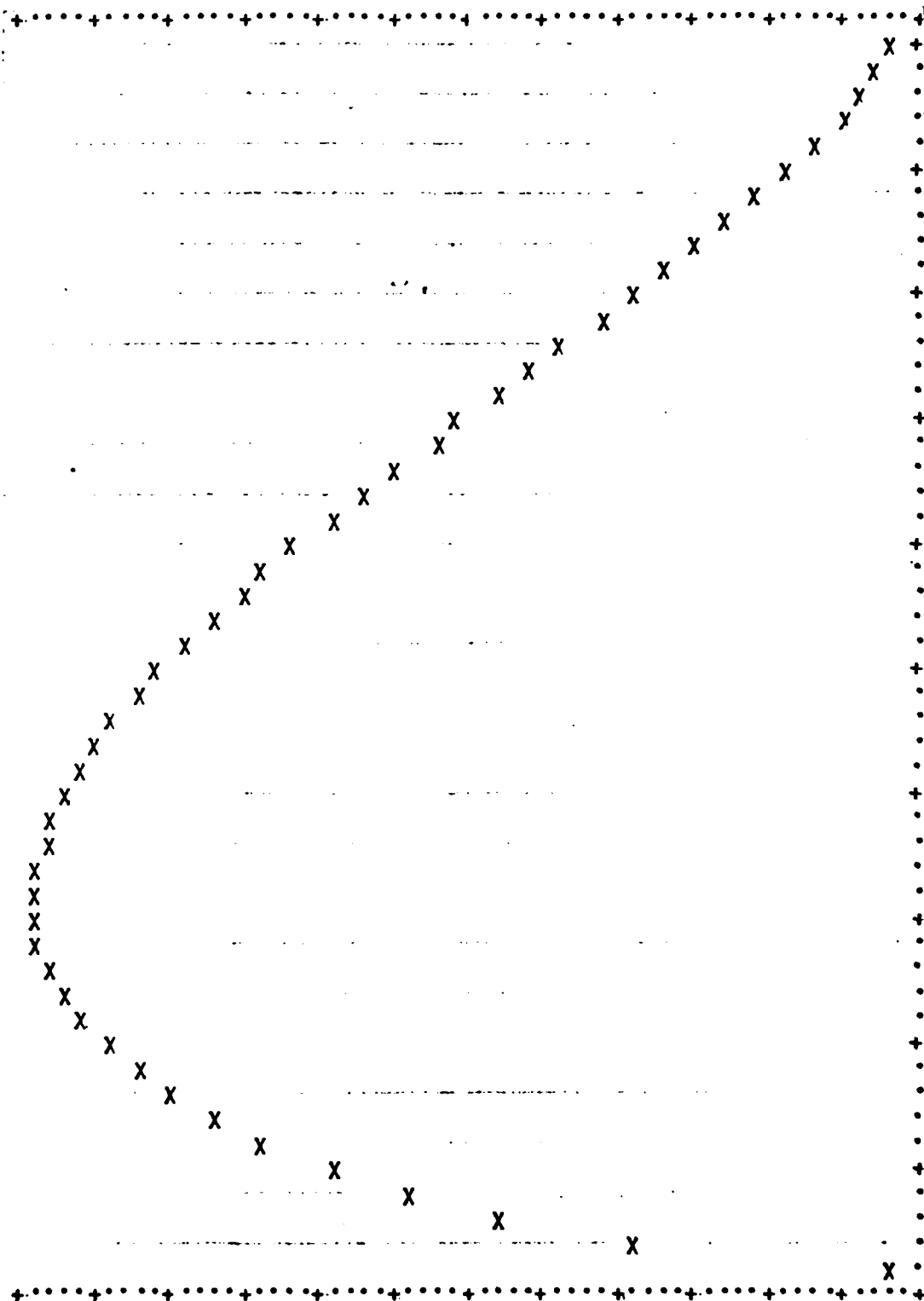
LEFT-TO-RIGHT
Normalized Mode = .25



RIGHT-TO-LEFT
Normalized Mode = .70

HIGH VARIANCE
Uncertainty Coefficient = .75

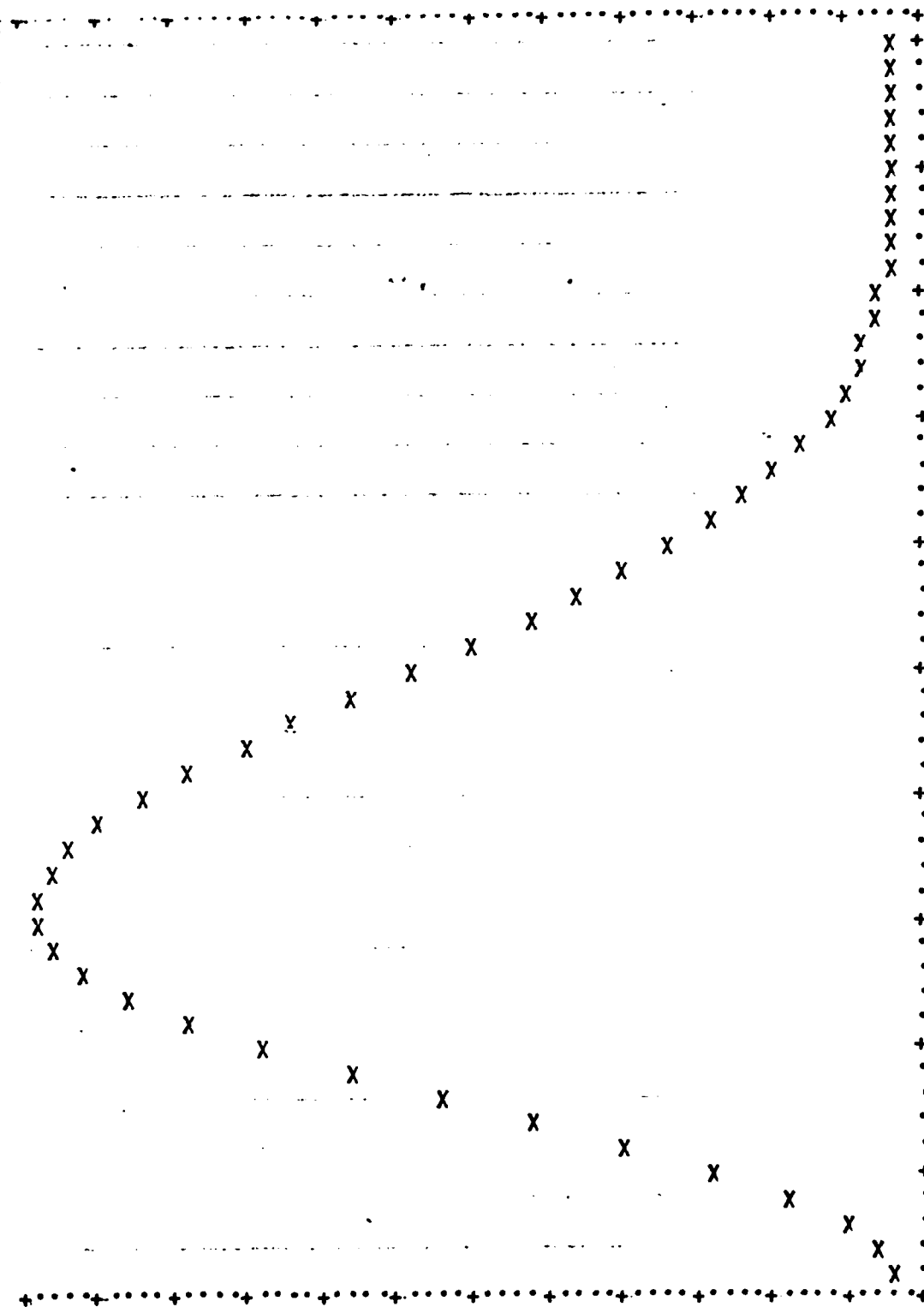
LEFT-TO-RIGHT
Normalized Mode = .30



RIGHT-TO-LEFT
Normalized Mode = .70

MEDIUM VARIANCE
Uncertainty Coefficient = .50

LEFT-TO-RIGHT
Normalized Mode = .30



LEFT-TO-RIGHT

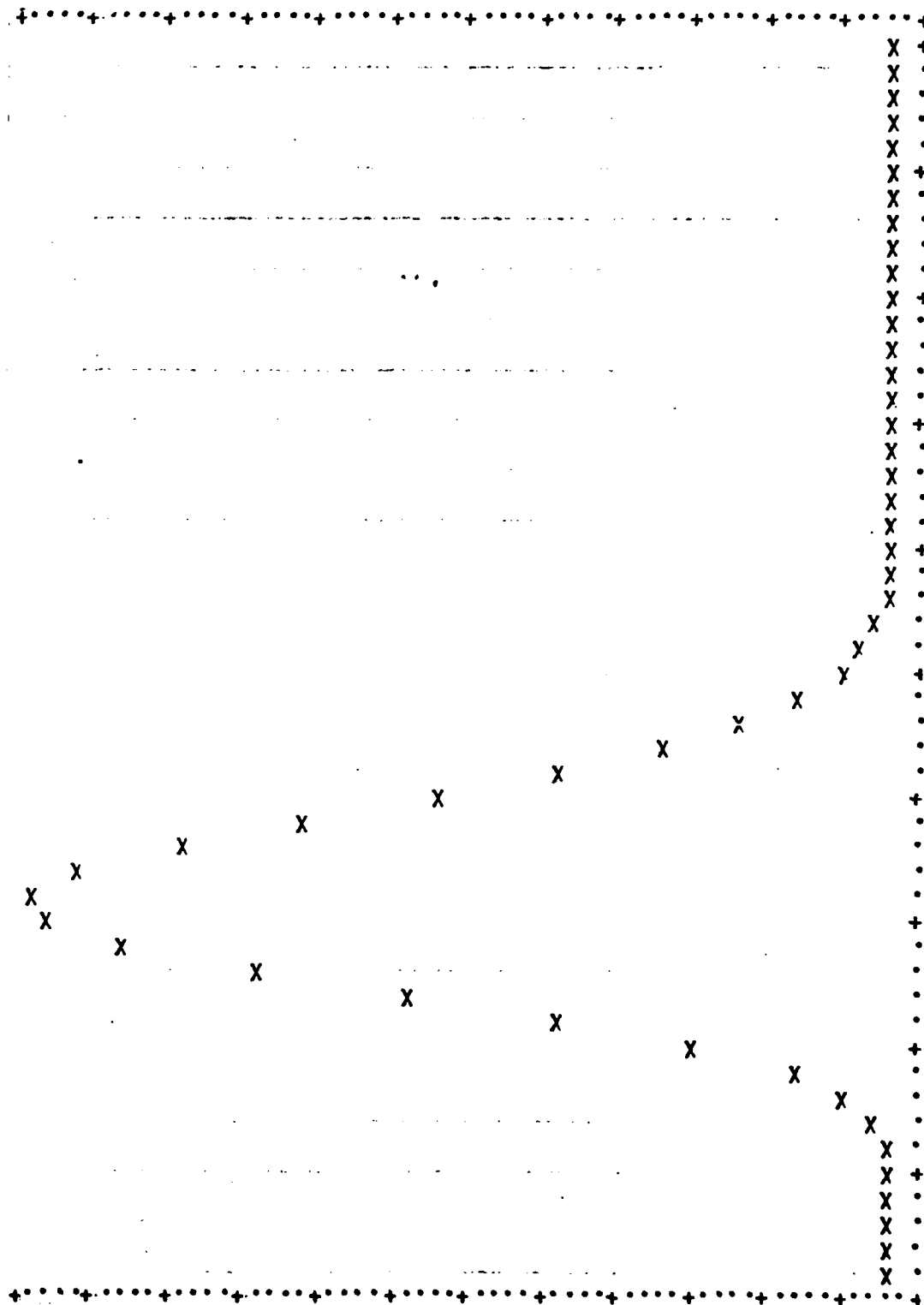
Normalized Mode = .30

LOW VARIANCE

Uncertainty Coefficient = .25

RIGHT-TO-LEFT

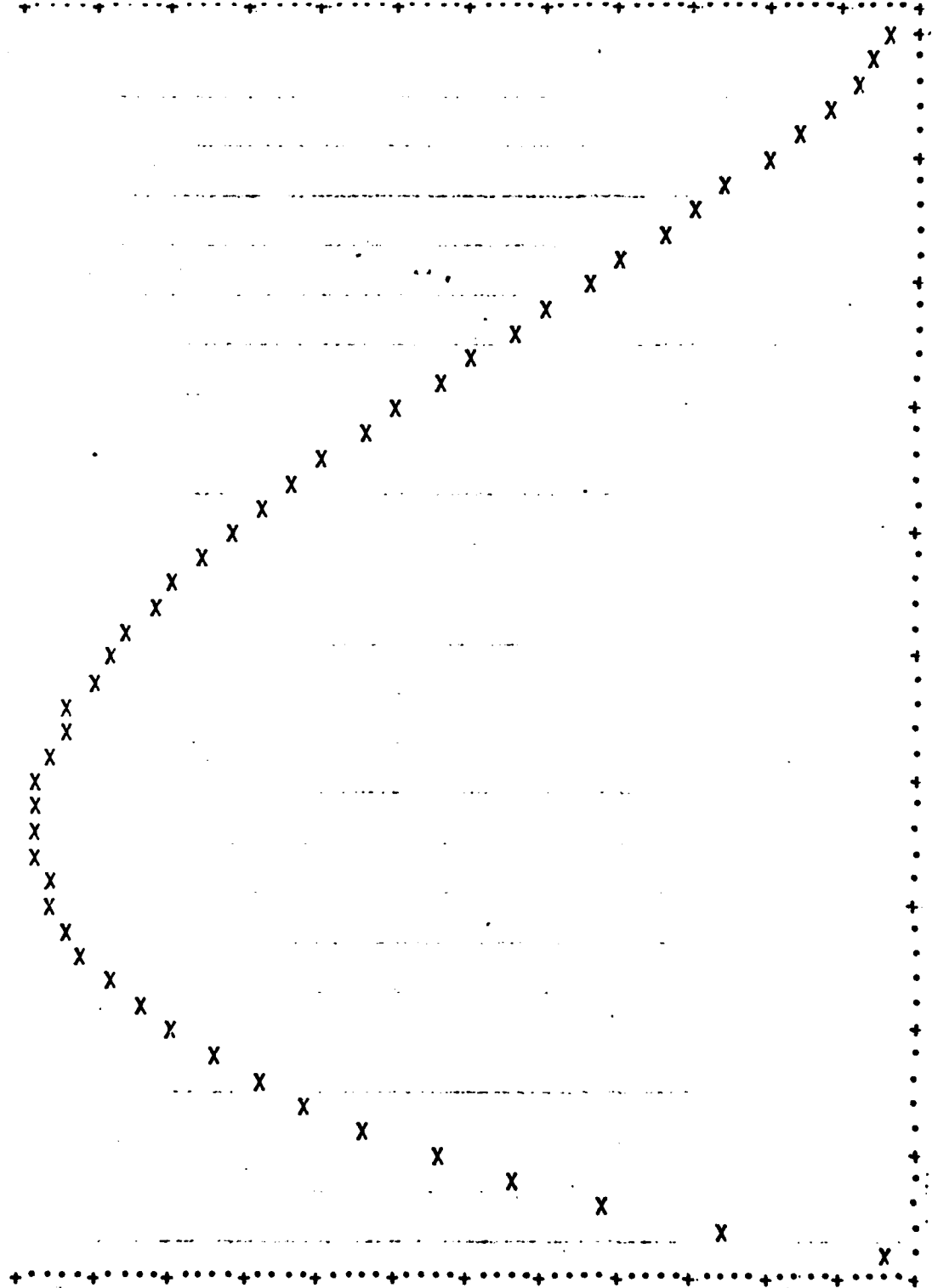
Normalized Mode = .70



RIGHT-TO-LEFT
Normalized Mode = .65

HIGH VARIANCE
Uncertainty Coefficient = .75

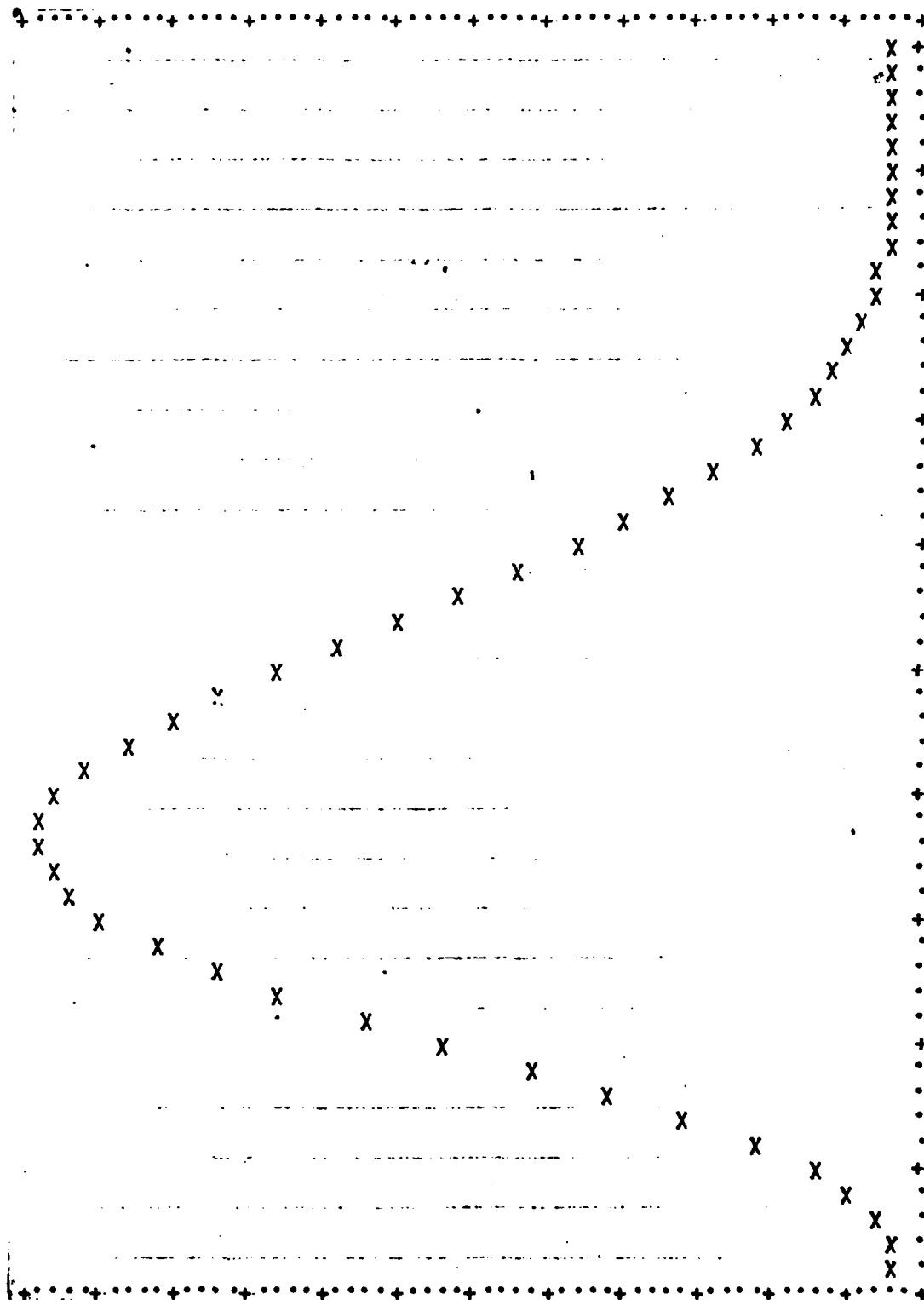
LEFT-TO-RIGHT
Normalized Mode = .35



RIGHT-TO-LEFT
Normalized Mode = .65

MEDIUM VARIANCE
Uncertainty Coefficient = .50

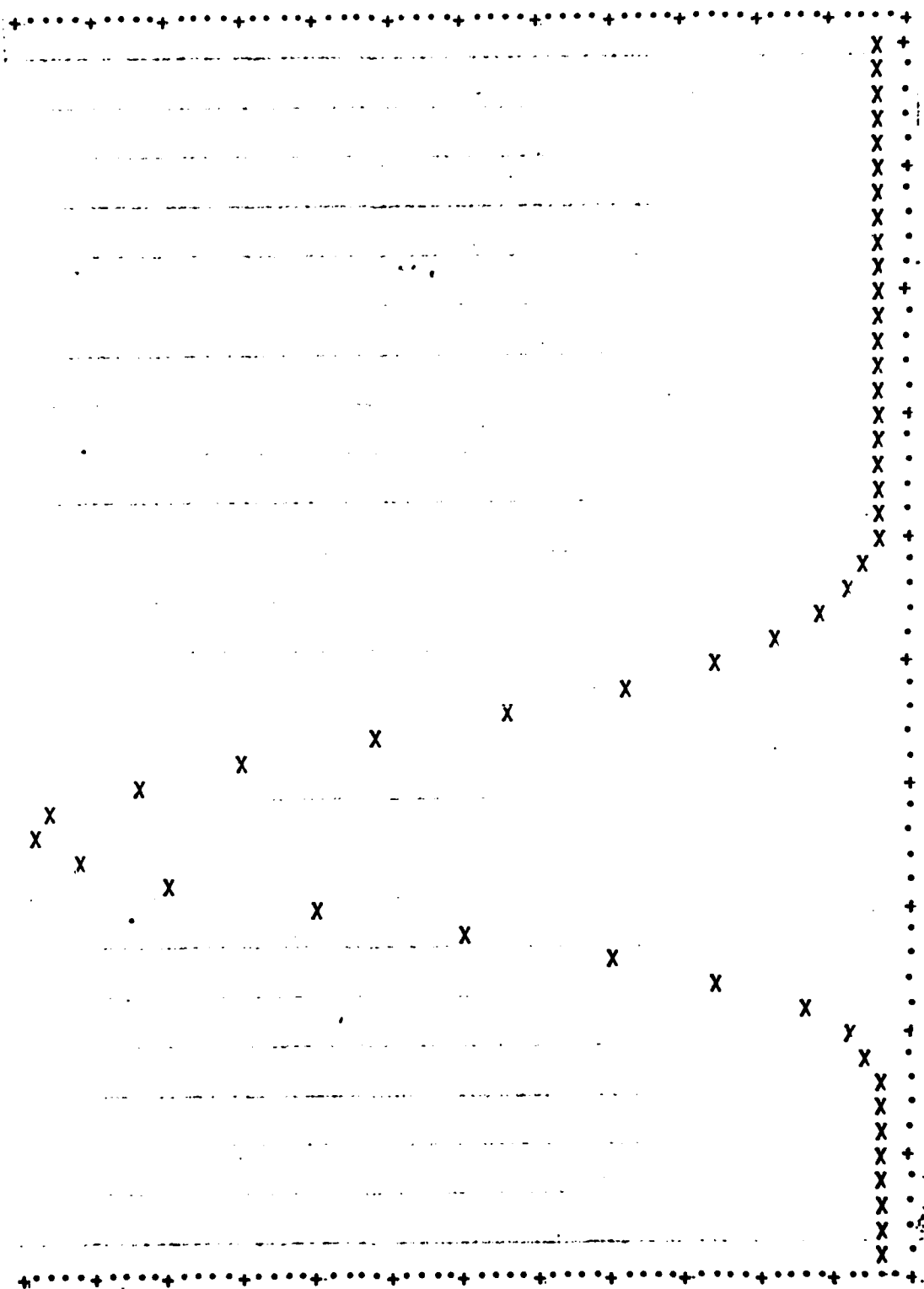
LEFT-TO-RIGHT
Normalized Mode = .35



RIGHT-TO-LEFT
Normalized Mode = .65

LOW VARIANCE
Uncertainty Coefficient = .25

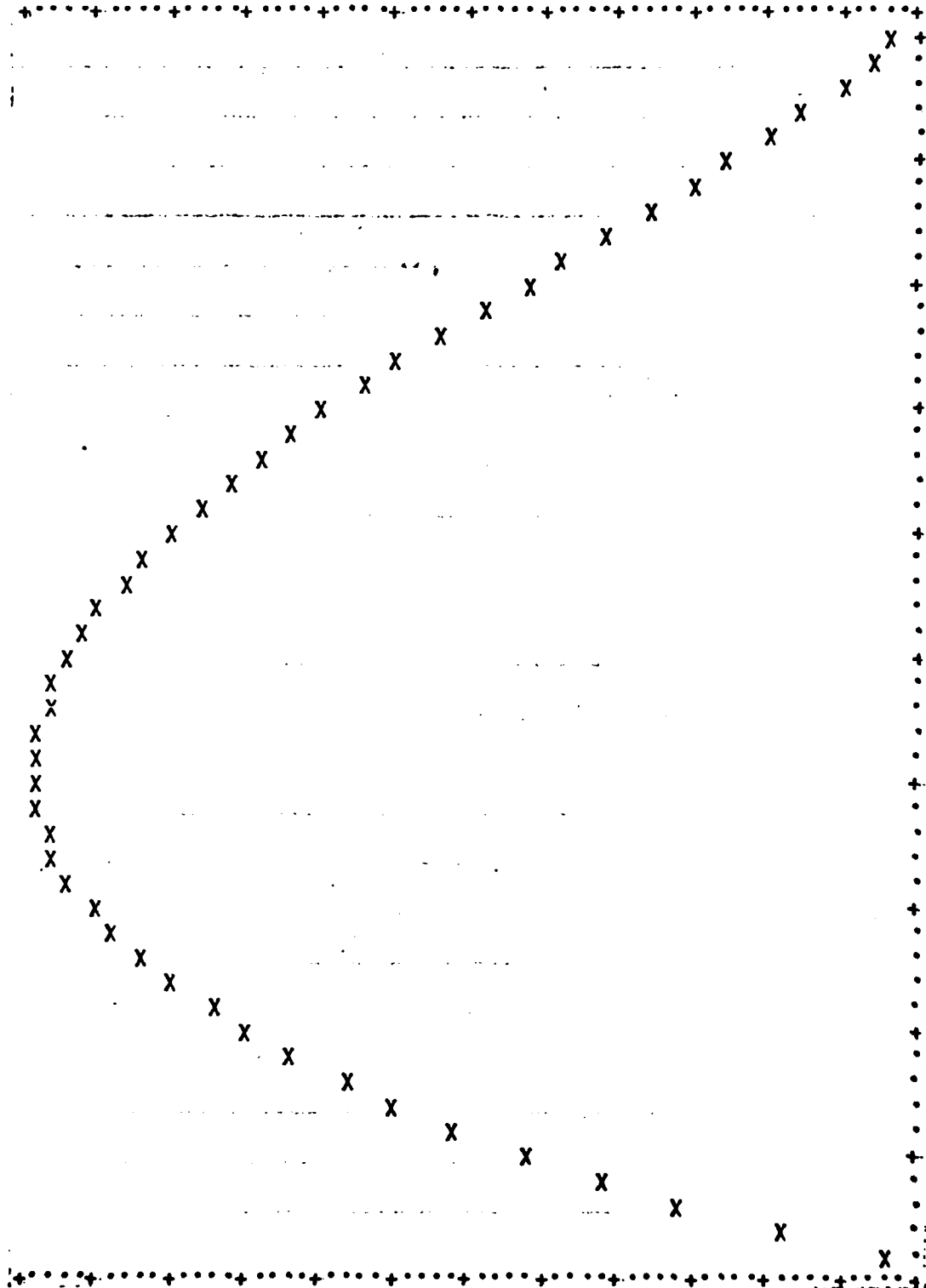
LEFT-TO-RIGHT
Normalized Mode = .35



LEFT-TO-RIGHT
Normalized Mode = .40

HIGH VARIANCE
Uncertainty Coefficient = .75

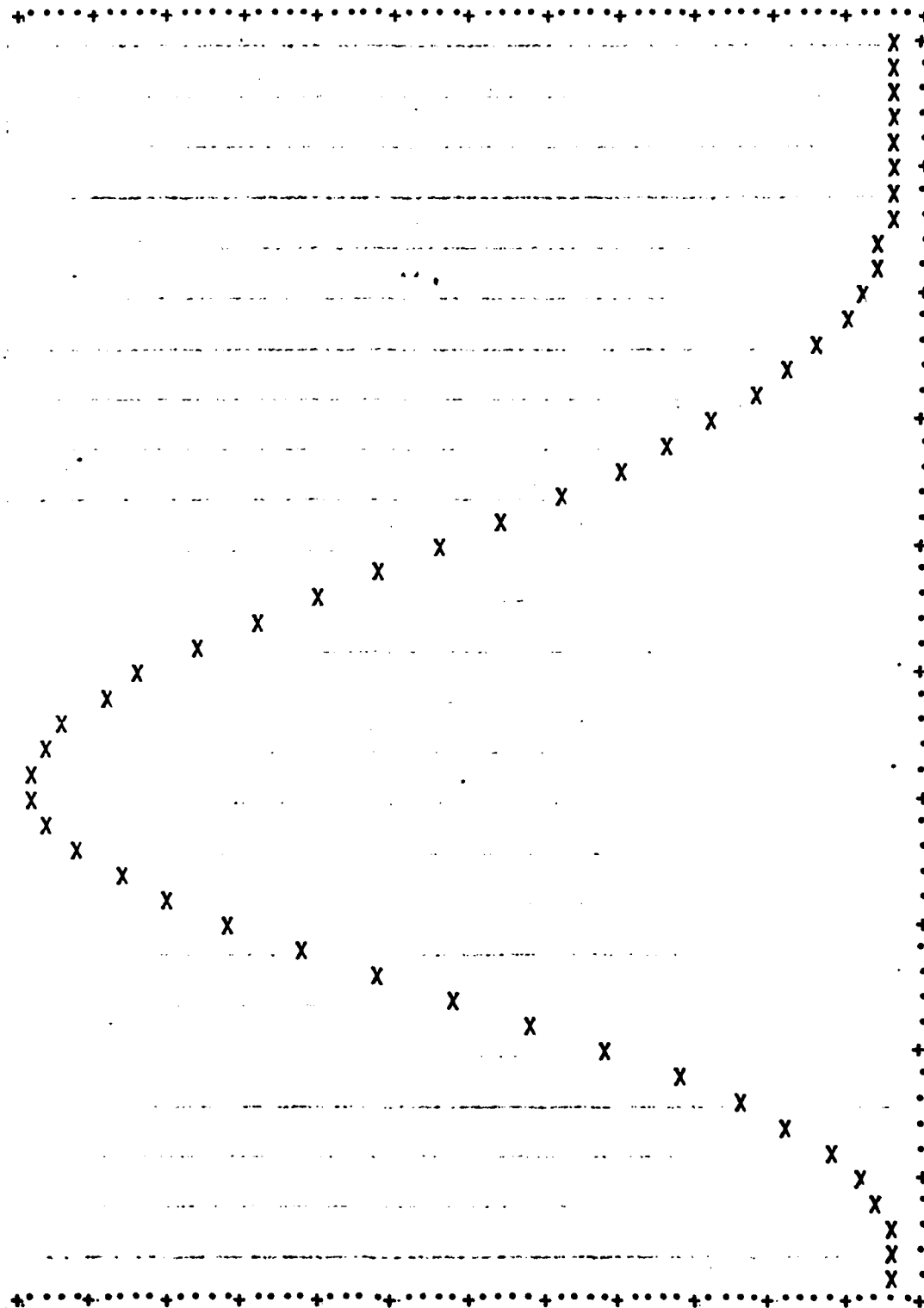
RIGHT-TO-LEFT
Normalized Mode = .60



RIGHT-TO-LEFT
Normalized Mode = .60

MEDIUM VARIANCE
Uncertainty Coefficient = .50

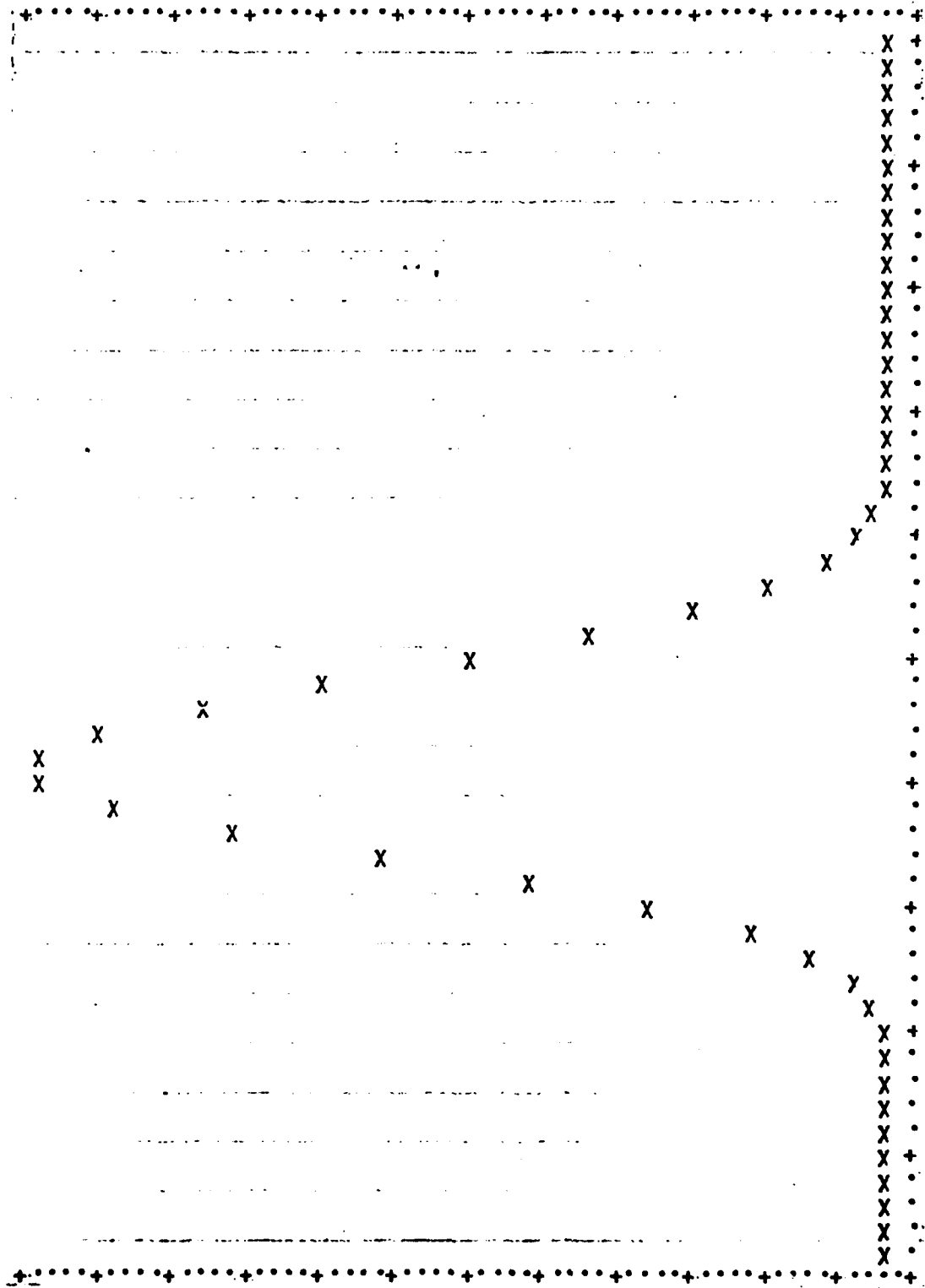
LEFT-TO-RIGHT
Normalized Mode = .40



LEFT-TO-RIGHT
Normalized Mode = .40

LOW VARIANCE
Uncertainty Coefficient = .25

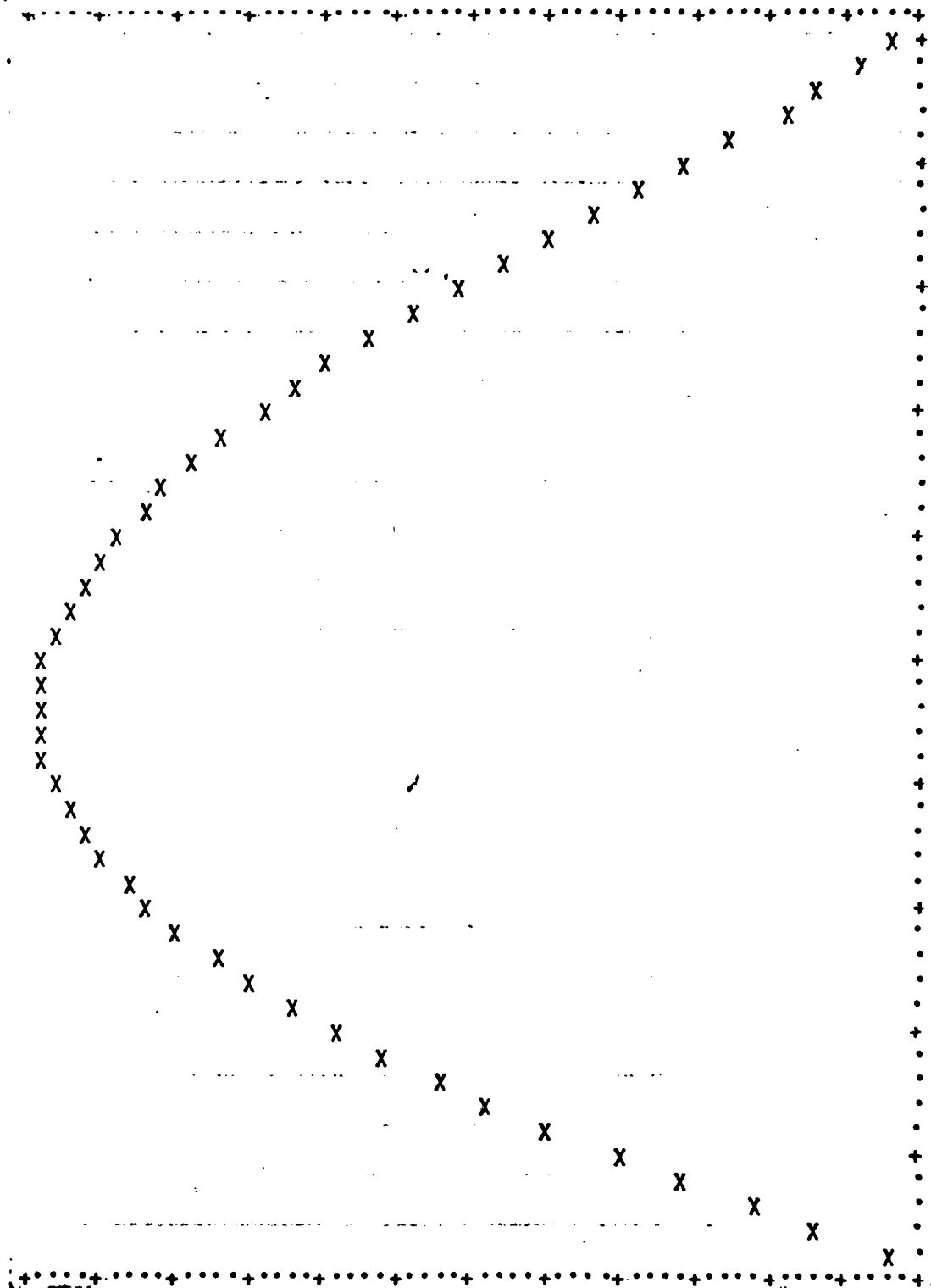
RIGHT-TO-LEFT
Normalized Mode = .60



LEFT-TO-RIGHT
Normalized Mode = .45

HIGH VARIANCE
Uncertainty Coefficient = .75

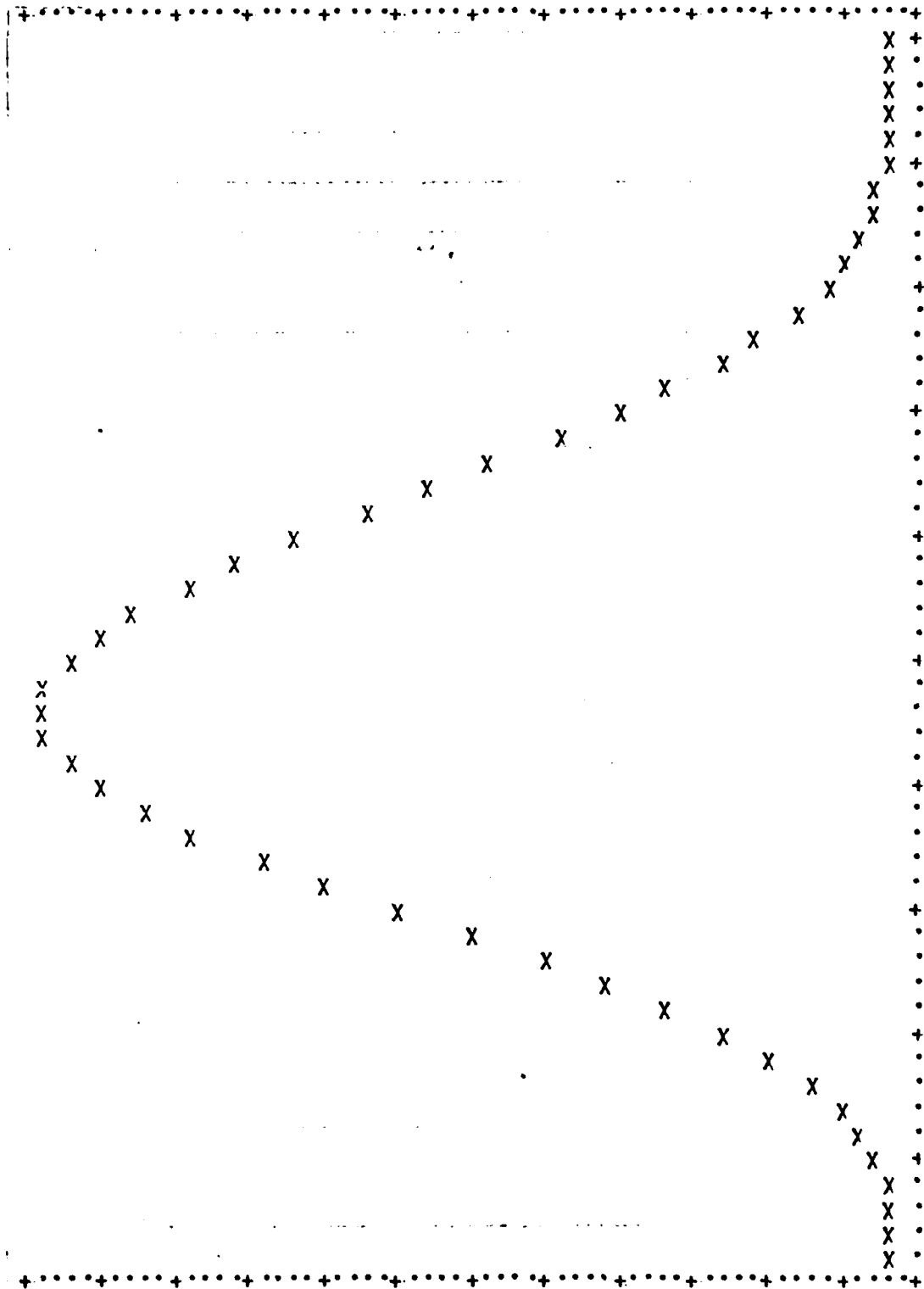
RIGHT-TO-LEFT
Normalized Mode = .55



LEFT-TO-RIGHT
Normalized Mode = .45

MEDIUM VARIANCE
Uncertainty Coefficient = .50

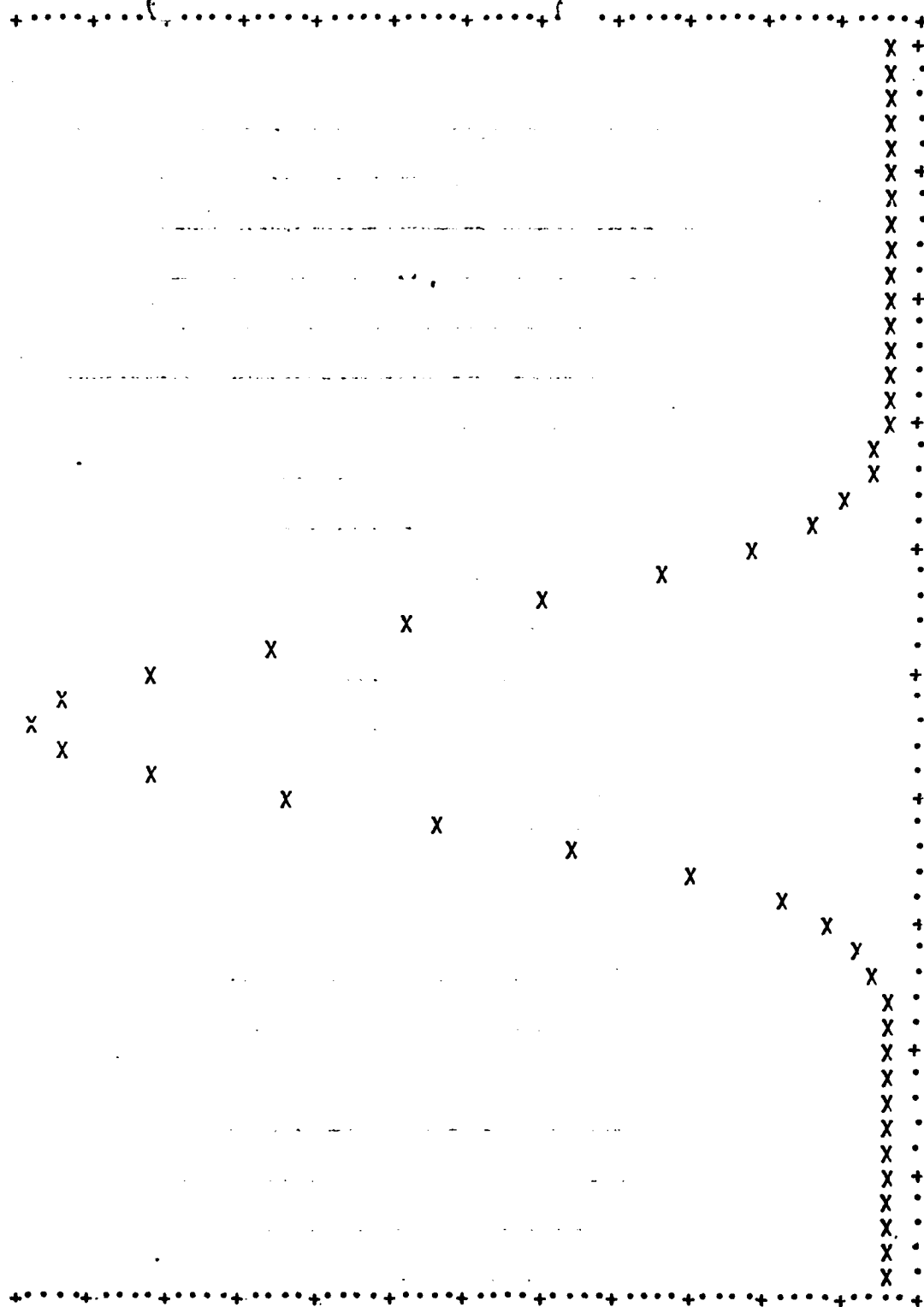
RIGHT-TO-LEFT
Normalized Mode = .55



LEFT-TO-RIGHT
Normalized Mode = .45

LOW VARIANCE
Uncertainty Coefficient = .25

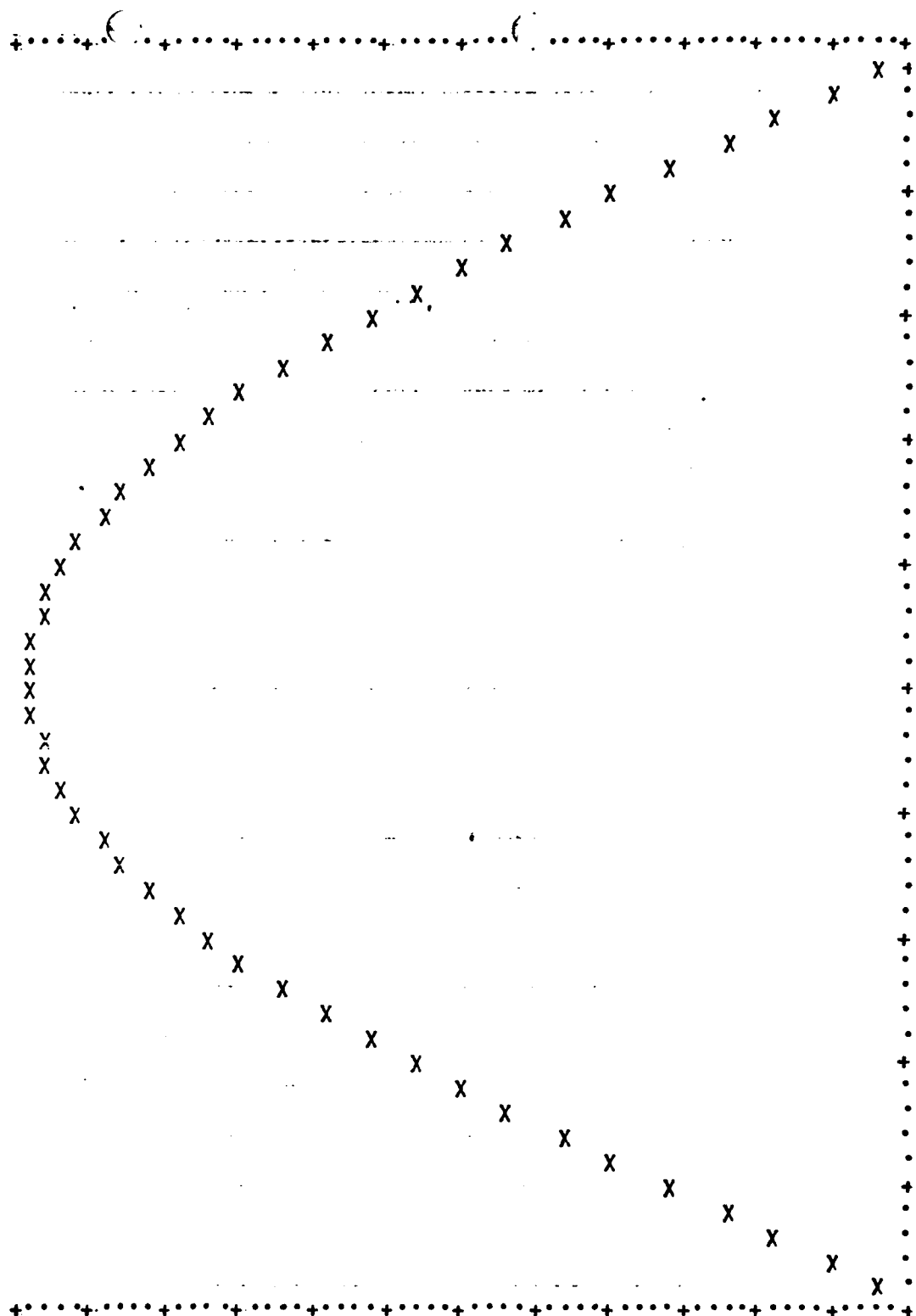
RIGHT-TO-LEFT
Normalized Mode = .55



RIGHT-TO-LEFT
Normalized Mode = .50

HIGH VARIANCE
Uncertainty Coefficient = .75

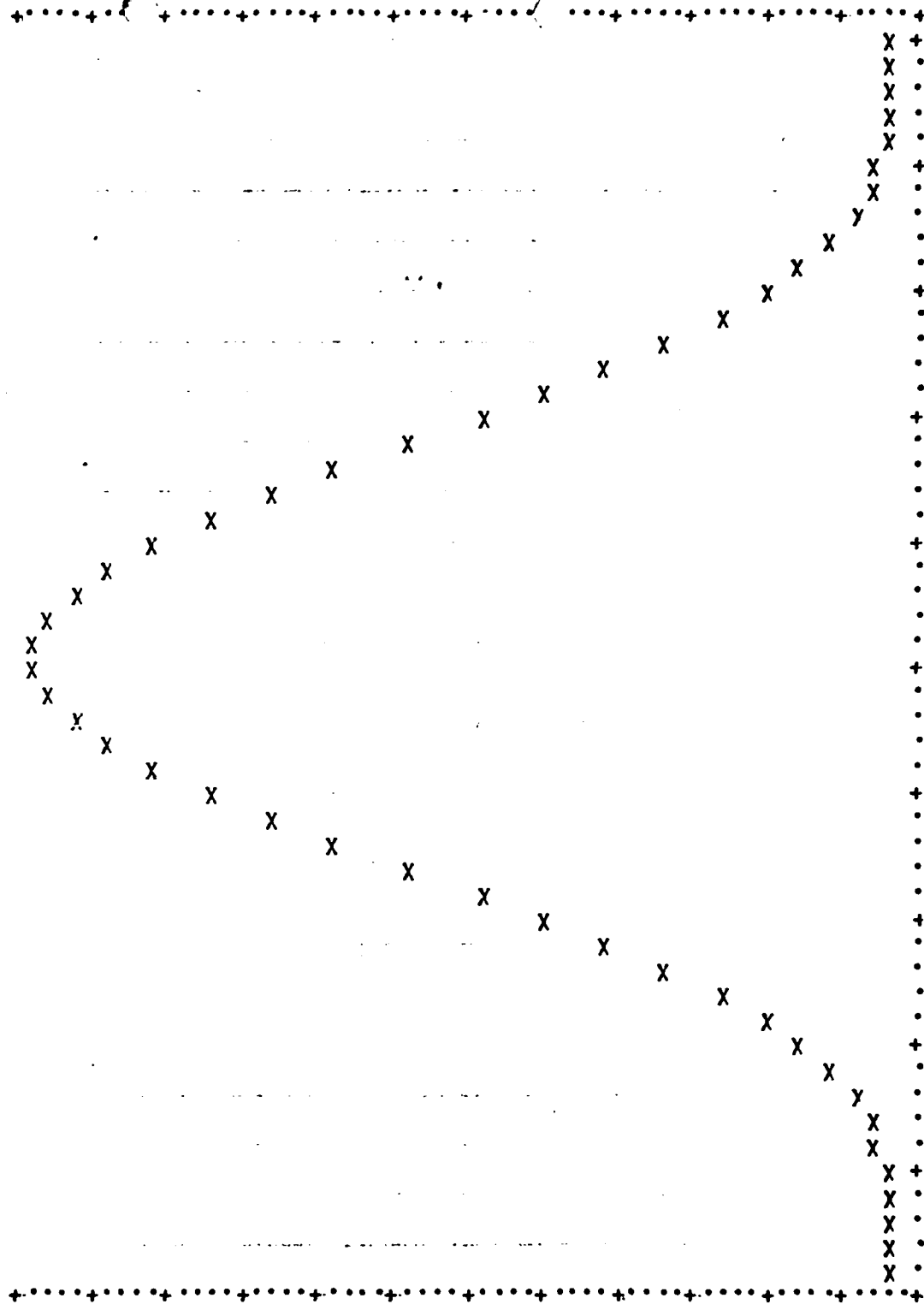
LEFT-TO-RIGHT
Normalized Mode = .50



LEFT-TO-RIGHT
Normalized Mode = .50

MEDIUM VARIANCE
Uncertainty Coefficient = .50

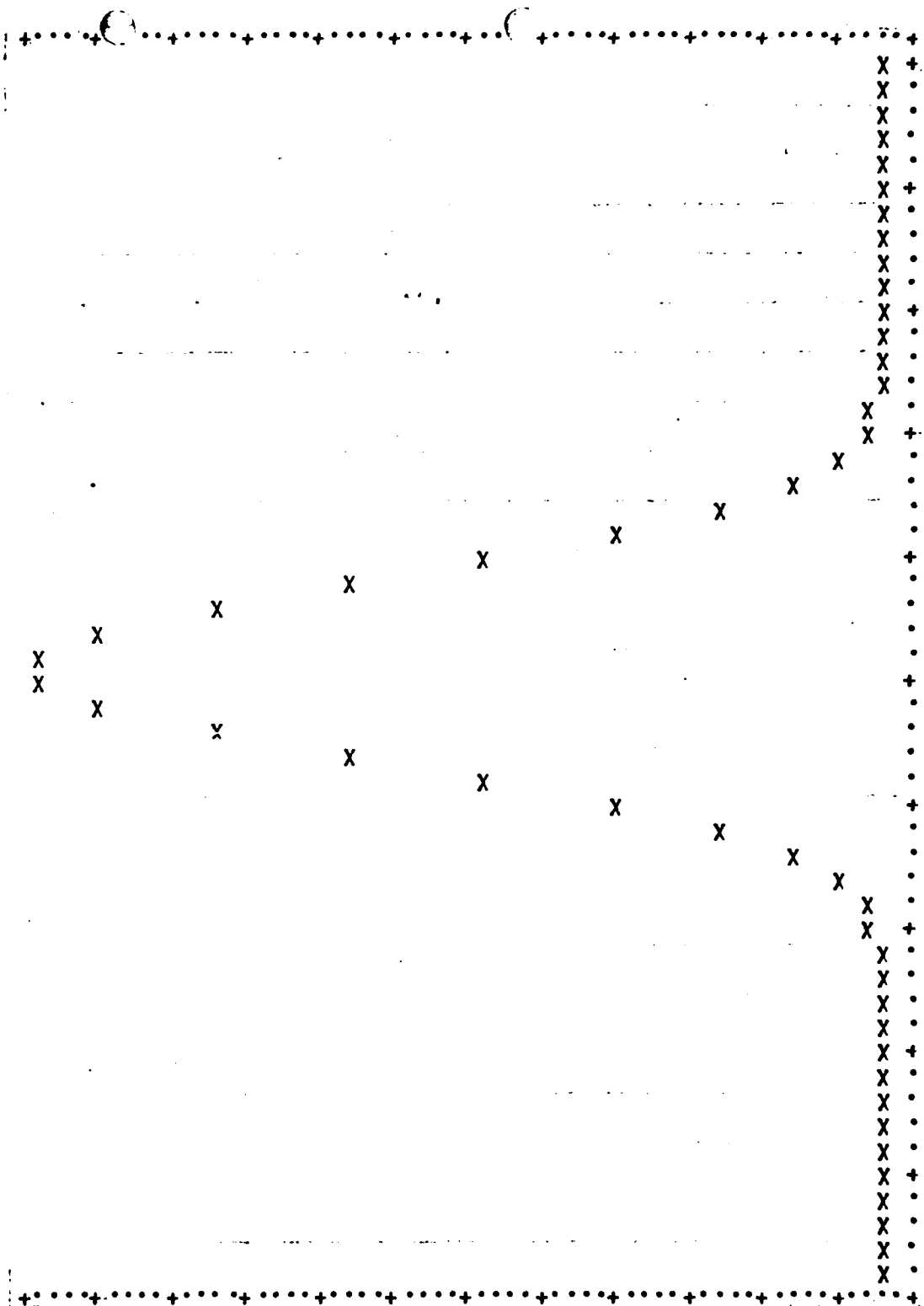
RIGHT-TO-LEFT
Normalized Mode = .50



LEFT-TO-RIGHT
Normalized Mode = .50

LOW VARIANCE
Uncertainty Coefficient = .25

RIGHT-TO-LEFT
Normalized Mode = .50



APPENDIX C

COMPUTER PROGRAM SPET

Description

Program SPET is written in FORTRAN IV for use in conjunction with a Univac 1108 computer in an interactive time-sharing mode. It is designed to approximate the frequency distribution of the sum of up to fifty independent beta or other random variables by means of Monte Carlo simulation. This is accomplished in the following manner:

(1) The user specifies the beta variables by entering the lowest (L), most likely (M), and highest (H) value as well as the uncertainty coefficient (U) of each variable. He specifies other variables by entering the lowest, most likely, and highest values, as well as the mean, variance, and cumulative density function (cdf) of each variable. He enters only those values of the cdf $F(x)$ corresponding to $x = L + i(H - L)/10$, $i = 1, \dots, 10$.

(2) For all beta variables the computer:

(a) Computes beta parameters a and b from M and U by first converting U to V with equation [6] and then simultaneously solving equations [3] and [4].

(b) Computes the discrete cdf, $F(x)$, for $x = i/10$, $i = 1, \dots, 10$, using Subroutine DQG32.^{1/}

(3) Next the class intervals for the distribution of the sum of the beta and other variables are computed. Adding the L values and throughput (TPUT)^{2/} establishes the lower limit of the range; adding the H values and TPUT establishes the upper limit. The range is then divided into 15 intervals of equal width.^{3/}

(4) Four frequency distributions of the sum are generated. The first distribution results from the assumption that the distributions making up the sum are statis-

^{1/} DQG-32 uses the 32 point Gaussian quadrature method of integration. It is taken from Convolution of Inverse Beta Distributions by a Sampling Technique (Bethesda, MD: Mathematica, Inc. 1971).

^{2/} Throughput means constant, and is usually the cost of a subsystem that is known with certainty.

^{3/} The number of intervals can be varied by assigning the desired number to KK in line 11 of the main program.

tically independent. It is generated as follows:

(a) Obtain a random number lying between zero and one from a uniform random number generator.^{4/}

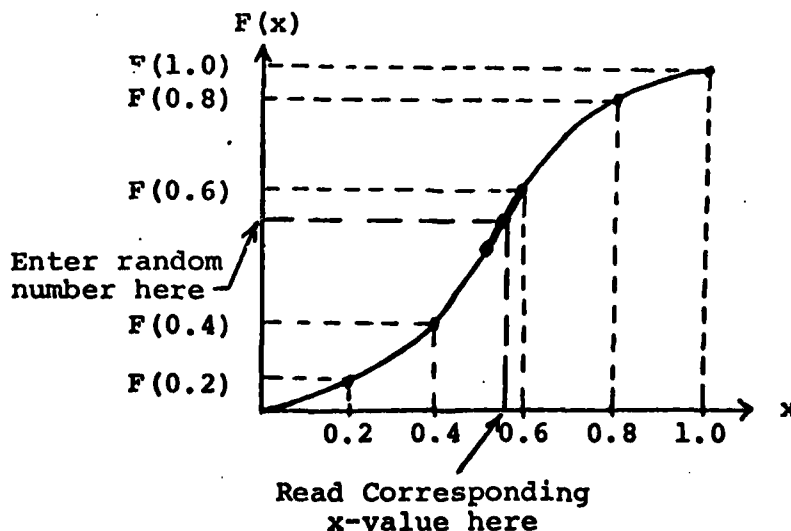
(b) Compare the generated number with the values of the discrete cdf, $F(x)$, for one of the variables and note the interval $[F(x_i), F(x_j)]$, $x_i < x_j$, into which it falls.

(c) Find the x-value, $x_i < x < x_j$, corresponding to the random number by means of linear interpolation.^{5/}

(d) Transform the x-value from its normalized (0,1) value to its standard (L,H) value by means of the transformation $x^* = L + (H - L)x$. [C1]

4/ An unsuccessful attempt was made to find a machine-independent random number generator for inclusion in the program. Therefore, Program SPET requires the use of a user-supplied generator. Make the appropriate changes in line 150 of the main program to accommodate the generator (it may be necessary to also change lines 68-69, 155, 158, 162, 175, 178, 182, 262, and 370-372).

5/ Steps (b) and (c) can be illustrated for a normalized random variable x as follows:



(e) Repeat steps (a) - (d) for every variable.

(f) Compute the observation,

$$X_j = \text{TPUT} + \sum_{i=1}^N x_i^*$$

where N denotes the number of random variables.

(g) Find the class interval [see (3) above] in which X_j occurs and register one occurrence in that interval.

(h) Repeat steps (a) - (g) for $j = 1, \dots, M$ observations, where M = number of desired observations. When step (h) is terminated, one has a frequency distribution of the sum of independent random variables.

(5) The second distribution of the sum is generated in the same way as the first except for step (e), which should now read, "Repeat steps (b) - (d) for every variable." This means that the same random number is used to obtain an observation on each component of the sum rather than a new number as done when constructing the first distribution. The procedure of using only one random number introduces a correlation among the component variables because when the value of one variable is known, it in turn maps uniquely to the values of all other variables. This correlation is positive because the cdfs are positive monotonic functions.

(6) The third distribution is simply the uniform distribution over the range of the first distribution. Both the second and third distributions serve as indicators of how a violation of the independence assumption could affect the distribution of the sum and the summary statistics.

(7) The fourth distribution is the same as the first distribution except the component variables are all uniform random variables. This distribution serves as an indicator of the relative sensitivity of the distribution of the sum to the degree of uncertainty in the component variables.

(8) Along with the four frequency distributions just described, the computer generates the mean, mode, standard deviation, variance, 90% confidence interval about the

mean, and the probability of exceeding the mean for each of the distributions. In addition, the user can specify any confidence interval about the mean and the probability of exceeding any number within the range of the distribution and the computer will generate it for him.

User Instructions

The following example illustrates the use of Program SPET:

ISPET

```

TITLE
**EXAMPLE**
NUMBER OF BETA DISTRIBUTIONS
4
DATA FILE(1) OR TERMINAL(2) INPUT
2
LOWEST,MODE,HIGHEST,UNCERTAINTY COEFF
DATA
100,250,900,.8
DATA
300,350,390,.2
DATA
200,400,675,.5
DATA
150,330,800,.4
NUMBER OF OTHER DISTRIBUTIONS
1
LOWEST,MODE,HIGHEST,MEAN,VARIANCE,DISCRETE CDF
DATA
50,75,100,75,208.3,.1,.2,.3,.4,.5,.6,.7,.8,.9,1.0
THROUGHPUT
35
NUMBER OF OBSERVATIONS
2000
SEED RANDOM NUMBER GENERATOR
87654321

```

EXAMPLE

INPUTS

```

OBSERVATIONS = 2000
SEED = 87654321

```

BETA VBLE	LOWEST	MODE	HIGHEST	U COEFF	NMODE	ALPHA	BETA	MEAN	VARIANCE
1	100.0	250.0	900.0	.80	.19	.3	1.1	396.4	33806.0
2	300.0	350.0	390.0	.20	.56	39.6	31.6	349.9	27.0
3	200.0	400.0	675.0	.50	.42	3.7	5.1	406.9	4697.4
4	150.0	330.0	800.0	.40	.28	3.6	9.4	349.3	5614.6
OTHER WBLE									
5	50.0	75.0	100.0					75.0	208.3
THROUGHPUT	35.0		35.0					35.0	
SUNS	835.0		2900.0					1612.5	44353.2

OTHER WBLE

CDF

	.1000	.2000	.3000	.4000	.5000	.6000	.7000	.8000	.9000	1.0000
5										

OUTPUT

INTERVAL	RANGE	INDEPENDENT BETA/OTHER		DEPENDENT BETA/OTHER		TOTAL UNIFORM		INDEPENDENT UNIFORM		
		PDF	COF	PDF	COF	PDF	COF	PDF	COF	
1	835.0 -	972.7	0.	0.	.0110	.0110	.0667	.0667	.0010	.0010
2	972.7 -	1110.3	.0010	.0010	.0545	.0655	.0667	.1333	.0025	.0035
3	1110.3 -	1248.0	.0240	.0250	.1040	.1695	.0667	.2000	.0215	.0250
4	1248.0 -	1385.7	.1140	.1390	.1245	.2940	.0667	.2667	.0435	.0645
5	1385.7 -	1523.3	.2270	.3660	.1355	.4295	.0667	.3333	.0805	.1490
6	1523.3 -	1661.0	.2360	.6020	.1365	.5660	.0667	.4000	.1140	.2630
7	1661.0 -	1798.7	.1945	.7965	.1170	.6830	.0667	.4667	.1575	.4205
8	1798.7 -	1936.3	.1240	.9205	.1115	.7945	.0667	.5333	.1615	.5420
9	1936.3 -	2074.0	.0600	.9805	.0875	.8820	.0667	.6000	.1525	.7345
10	2074.0 -	2211.7	.0185	.9990	.0540	.9360	.0667	.6667	.1070	.8415
11	2211.7 -	2349.3	.0010	1.0000	.0385	.9745	.0667	.7333	.0445	.9310
12	2349.3 -	2487.0	0.	1.0000	.0175	.9920	.0667	.8000	.0415	.9725
13	2487.0 -	2624.7	0.	1.0000	.0065	.9985	.0667	.8667	.0210	.9935
14	2624.7 -	2762.3	0.	1.0000	.0015	1.0000	.0667	.9333	.0060	.9995
15	2762.3 -	2900.0	0.	1.0000	0.	1.0000	.0667	1.0000	.0005	1.0000

MEAN 1615.5 1621.1 1867.5 1869.2
 MODE 1592.2 1592.2 1867.5 1867.5
 VARIANCE 44843.7 128488.0 355352.1 103015.0
 STD DEVIATION 211.8 358.5 596.1 321.0
 90% CONFIDENCE INTERVAL 1307.2, 1124.8, 938.3, 1337.5,
 2059.2 2364.7 2794.8 2423.2
 PROB EXCEED MEAN .48 .47 .50 .50

ANOTHER CONFIDENCE INTERVAL?
 95 1277.0, 1071.2, 886.6, 1258.3,
 95% CONFIDENCE INTERVAL 2218.5 2900.0 2849.4 2524.7

ANOTHER CONFIDENCE INTERVAL?
 0

PROB EXCEED SOME VALUE?
 1450 .75 .64 .70 .89
 PROB EXCEED 1450.0

PROB EXCEED SOME VALUE?
 0

ADDITIONAL OBSERVATIONS?
 0

ANOTHER SEED?
 0
 STOP
 SEU'S114.7

- (1) The user enters a title up to 60 characters long.
- (2) If he is using beta distributions to represent some or all of his component distributions, the user enters the number of distributions to be so represented and proceeds to step (a) below. However, if he is not using the beta, he enters the number "0" and proceeds to step (3).
 - (a) The user specifies whether he will enter the four-tuples L, M, H, U defining each beta variable directly from the terminal or from a data file stored in the computer by entering the number "1" for data file input or the number "2" for terminal input.
 - (b) If he chooses the terminal input, his next step is to enter one four-tuple L, M, H, U for each beta variable. If he chooses the data file input, he merely enters the name of the data file.
- (3) If he is using distributions other than the beta to represent some or all of his component distributions, the user enters the number of distributions to be so represented, followed by the L, M, H, mean, variance, and the discrete cdf (see page C-1) of the distributions he has chosen. If he is not using other distributions, he enters the number "0".
- (4) If there is a throughput (constant), the user enters it now.
- (5) He then enters the number of observations (sample size) he desires, followed by a seed for the random number generator.
- (6) The computer prints the user's inputs, followed by the output.
- (7) The computer then queries the user if he desires another confidence interval. The user responds with the confidence interval he desires, or types the number "0" if he desires none.
- (8) The computer asks the user if he desires the probability that some value within the range of the distribution will be exceeded. The user responds with that number, or the number "0" if he desires none.
- (9) The computer inquires to see if the user desires additional observations. The user responds with the num-

ber of additional observations he desires (the number "0" if none). If he desires additional observations, the computer repeats steps (6) - (9).

(10) The computer asks the user if he desires to use another seed for the random number generator. The user responds with the seed if he does, or with the number "0" if he does not. If he enters another seed, the computer repeats steps (6) - (10).

Program Listing

```

1  DOUBLE PRECISION A(50),H(50),XL,XU,Y,ALPHA,BETA
2  REAL L(50),M(50),MODE(50),NM(50),U(50),C(50),MEAN(50),VAR(50)
3  REAL TITLE(12),VALINT(10,50),W(51),P1(15),P2(15),P3(15),P4(15)
4  REAL C1(15),C2(15),C3(15),C4(15)
5  REAL MODE1,MEAN1,MODE2,MEAN2,MEAN3,MODE4,MEAN4,MSUM,L1,L2,L3,L4
6  REAL NAME(4)
7  INTEGER F1(15),F2(15),F4(15)
8  DATA NAME/'!EQU','ATE ','3 ',' '
9  II=6
10 JJ=5
11 KK=15
12 %
13 % INPUT
14 %
15 WRITE(II,5)
16 5 FORMAT(//6H TITLE)
17 READ(JJ,7) TITLE
18 7 FORMAT(15A4)
19 WRITE(II,10)
20 10 FORMAT(1H ,28HNUMBER OF BETA DISTRIBUTIONS)
21 READ(JJ,*) N1
22 IF(N1.EQ.0) GO TO 60
23 WRITE(II,11)
24 11 FORMAT(1H ,31HDATA FILE(1) OR TERMINAL(2) INPUT)
25 READ(JJ,*) INPUT
26 IF(INPUT.EQ.2) GO TO 16
27 WRITE(II,13)
28 13 FORMAT(1H ,17HNAME OF DATA FILE)
29 READ(JJ,14) NAME(4)
30 14 FORMAT(A4)
31 CALL OBEY(NAME,4)
32 DO 15 I=1,N1
33 READ(3,*) L(I),MODE(I),H(I),U(I)
34 15 CONTINUE
35 GO TO 60
36 16 CONTINUE
37 WRITE(II,30)
38 30 FORMAT(1H ,37HLOWEST,MODE,HIGHEST,UNCERTAINTY COEFF)
39 DO 40 I=1,N1
40 WRITE(II,50)
41 50 FORMAT(1H 4HDATA)
42 READ(JJ,*) L(I),MODE(I),H(I),U(I)
43 40 CONTINUE
44 60 CONTINUE
45 WRITE(II,70)
46 70 FORMAT(1H ,29HNUMBER OF OTHER DISTRIBUTIONS)
47 READ(JJ,*) N2
48 IF(N2.EQ.0) GO TO 78
49 NOME=N1+1
50 NTWO=N1+N2
51 WRITE(II,75)
52 75 FORMAT(1H ,46HLOWEST,MODE,HIGHEST,MEAN,VARIANCE,DISCRETE CDF)
53 DO 76 I=NOME,NTWO
54 WRITE(II,50)
55 READ(JJ,*) L(I),MODE(I),H(I),MEAN(I),VAR(I),(VALINT(J,I),J=1,10)
56 76 CONTINUE
57 GO TO 79
58 78 CONTINUE
59 NTWO=N1
60 79 WRITE(II,72)
61 72 FORMAT(1H ,10HTHROUGHPUT)
62 READ(JJ,*) TPUT
63 WRITE(II,77)
64 77 FORMAT(1H ,20HNUMBER OF OBSERVATIONS)
65 READ(JJ,*) M
66 WRITE(II,80)
67 80 FORMAT(1H ,28HSEED RANDOM NUMBER GENERATOR)
68 READ(JJ,*) KSD
69 KSD2=KSD
70 IF(N1.EQ.0) GO TO 105
71 %
72 % COMPUTE A AND B
73 %
74 DO 85 I=1,N1
75 N=0

```

```

76 NH(I)=(MODE(I)-L(I))/(H(I)-L(I))
77 V=(1.28467508*U(I))*2.
78 B(I)=1.0
79 86 A(I)=H(I)*NM(I)/(1.-NM(I))
80 VBETA=((A(I)+1.)*(B(I)+1.))/((A(I)+B(I)+2.)*2.)*(A(I)+B(I)+3.)
81 Z=V-VBETA
82 IF(N.EQ.1) GO TO 87
83 IF(Z) 3,85,1
84 87 IF(Z) 4,85,85
85 3 B(I)=B(I)+1.
86 GO TO 86
87 1 B(I)=B(I)-1.
88 N=N-1
89 GO TO 86
90 4 B(I)=B(I)+.05
91 GO TO 86
92 85 CONTINUE
93 *
94 * COMPUTE DISCRETE CDFS
95 *
96 XL=0.00
97 DO 90 J=1,N1
98 XU=1.00
99 ALPHA=A(J)
100 BETA=B(J)
101 CALL DGG32(XL,XU,Y,ALPHA,BETA)
102 C(J)=1./Y
103 XU=0.00
104 VALINT(10,J)=1.
105 DO 100 I=1,9
106 XU=XU+.100
107 CALL DGG32(XL,XU,Y,ALPHA,BETA)
108 VALINT(I,J)=C(J)*Y
109 100 CONTINUE
110 90 CONTINUE
111 *
112 * COMPUTE CLASS INTERVALS
113 *
114 105 CONTINUE
115 XLOW=TPUT
116 XHIGH=TPUT
117 DO 110 I=1,NTWO
118 XLOW=XLOW+L(I)
119 XHIGH=XHIGH+H(I)
120 110 CONTINUE
121 RANGE=XHIGH-XLOW
122 WIDTH=RANGE/KK
123 W(1)=XLOW
124 W(KK+1)=XHIGH
125 DO 120 I=2,KK
126 W(I)=W(I-1)+WIDTH
127 120 CONTINUE
128 *
129 * GENERATE HISTOGRAM
130 *
131 140 CONTINUE
132 TOTAL1=0.
133 TOTAL2=0.
134 TOTAL4=0.
135 SUM1=0.
136 SUM2=0.
137 SUM4=0.
138 SUMSQ1=0.
139 SUMSQ2=0.
140 SUMSQ4=0.
141 MSTAR=M
142 DO 150 I=1,KK
143 F1(I)=0
144 F2(I)=0
145 F4(I)=0
146 150 CONTINUE
147 160 CONTINUE
148 DO 180 K=1,MSTAR
149 DO 190 J=1,NTWO
150 RAND=UORNRT(KS0)

```

```

151 IF(J.GT.1)GO TO 195
152 DO 200 LL=1,NTWO
153 DO 205 I=1,10
154 M=I
155 IF(RAND.LE.VALINT(I,LL)) GO TO 210
156 205 CONTINUE
157 210 IF(M.EQ.1) GO TO 215
158 T2=(RAND-VALINT(M-1,LL))/(VALINT(M,LL)-VALINT(M-1,LL))
159 S=(M-1)/10.
160 SMALL2=.1+T2+S
161 GO TO 220
162 215 SMALL2=.1+(RAND/VALINT(I,LL))
163 220 TOTAL2=TOTAL2+SMALL2*(H(LL)-L(LL))+L(LL)
164 200 CONTINUE
165 TOTAL2=TOTAL2+TPUT
166 SUM2=SUM2+TOTAL2
167 SUMSQ2=SUMSQ2+TOTAL2**2.
168 COST2=(TOTAL2-XLO#)/RANGE
169 J2=COST2*KK+1.
170 IF(J2.EQ.KK+1) J2=KK
171 F2(J2)=F2(J2)+1.
172 TOTAL2=0.
173 195 DO 205 I=1,10
174 M=I
175 IF(RAND.LE.VALINT(I,J)) GO TO 227
176 225 CONTINUE
177 227 IF(M.EQ.1) GO TO 231
178 T1=(RAND-VALINT(M-1,J))/(VALINT(M,J)-VALINT(M-1,J))
179 S=(M-1)/10.
180 SMALL1=.1+T1+S
181 GO TO 233
182 231 SMALL1=.1+(RAND/VALINT(I,J))
183 233 TOTAL1=TOTAL1+SMALL1*(H(J)-L(J))+L(J)
184 TOTAL4=TOTAL4+RAND*(H(J)-L(J))+L(J)
185 190 CONTINUE
186 TOTAL1=TOTAL1+TPUT
187 TOTAL4=TOTAL4+TPUT
188 SUM1=SUM1+TOTAL1
189 SUM4=SUM4+TOTAL4
190 SUMSQ1=SUMSQ1+TOTAL1**2.
191 SUMSQ4=SUMSQ4+TOTAL4**2.
192 COST1=(TOTAL1-XLO#)/RANGE
193 J1=COST1*KK+1.
194 IF(J1.EQ.KK+1) J1=KK
195 F1(J1)=F1(J1)+1.
196 COST4=(TOTAL4-XLO#)/RANGE
197 J4=COST4*KK+1.
198 IF(J4.EQ.KK+1) J4=KK
199 F4(J4)=F4(J4)+1.
200 TOTAL1=0.
201 TOTAL4=0.
202 180 CONTINUE
203 *
204 * COMPUTE STATISTICS
205 *
206 IC=90
207 DO 235 I=1,KK
208 P1(I)=FLOAT(F1(I))/M
209 P2(I)=FLOAT(F2(I))/M
210 P4(I)=FLOAT(F4(I))/M
211 235 CONTINUE
212 CALL STAT(SUM1,M,KK,P1,W,SUMSQ1,MEAN1,MODE1,VAR1,STD1)
213 CALL CI(MEAN1,IC,P1,KK,W,WIDTH,L1,U1)
214 CALL CDF(P1,MEAN1,KK,W,C1,PX1)
215 CALL STAT(SUM2,M,KK,P2,W,SUMSQ2,MEAN2,MODE2,VAR2,STD2)
216 CALL CI(MEAN2,IC,P2,KK,W,WIDTH,L2,U2)
217 CALL CDF(P2,MEAN2,KK,W,C2,PX2)
218 CALL STAT(SUM4,M,KK,P4,W,SUMSQ4,MEAN4,MODE4,VAR4,STD4)
219 CALL CI(MEAN4,IC,P4,KK,W,WIDTH,L4,U4)
220 CALL CDF(P4,MEAN4,KK,W,C4,PX4)
221 IF(N1.EQ.0) GO TO 252
222 DO 250 I=1,N1
223 MEAN(I)=(A(I)+B(I)+R(I)+L(I)+H(I)+L(I))/(A(I)+B(I)+2.)
224 VAR(I)=(A(I)+1.)*(B(I)+1.)*(H(I)-L(I)+2.)/((A(I)+B(I)+2.)*(A(I)+B(I)+3.))
225 250 CONTINUE

```

```

226      252 CONTINUE
227      MSUM=TPUT
228      VSUM=0.
229      DO 251 I=1,NTWO
230      MSUM=MSUM+MEAN(I)
231      VSUM=VSUM+VAR(I)
232      251 CONTINUE
233      %
234      MEAN3=(XHIGH+XLOW)/2.
235      VAR3=(1./12.)*RANGE**2.
236      STD3=SQRT(VAR3)
237      UNI=1./KK
238      DO 290 I=1,KK
239      P3(I)=UNI
240      290 CONTINUE
241      CALL C1(MEAN3,IC,P3,KK,W,WIDTH,L3,U3)
242      CALL CDF(P3,MEAN3,KK,W,C3,PX3)
243      %
244      % PRINT INPUT
245      %
246      WRITE(II,295)
247      295 FORMAT(1H )
248      WRITE(II,300)
249      300 FORMAT(/)
250      WRITE(II,310)
251      310 FORMAT(/)
252      WRITE(II,320)
253      320 FORMAT(////)
254      WRITE(II,330) TITLE
255      330 FORMAT(1H ,55X,15A4)
256      WRITE(II,320)
257      WRITE(II,335)
258      335 FORMAT(1H ,55X,11H I N P U T S)
259      WRITE(II,320)
260      WRITE(II,390) M
261      390 FORMAT(1H ,15H OBSERVATIONS = ,15)
262      WRITE(II,400) SEED
263      400 FORMAT(1H ,7H SEED = ,5X,18)
264      WRITE(II,310)
265      WRITE(II,340)
266      340 FORMAT(1H ,9H BETA VRLE,7X,6H LOWEST,6X,4H MODE,5X,7H HIGHEST,
267      5X,7H COEFF,5X,5H MODE,5X,5H ALPHA,6X,4H BETA,7X,4H MEAN,5X,6H VARIANCE)
268      WRITE(II,295)
269      IF(N1.EQ.0) GO TO 351
270      WRITE(II,350) (I,L(I),MODE(I),H(I),U(I),NM(I),A(I),B(I),MEAN(I),VAR(I),I=1,N1)
271      350 FORMAT(1H ,3X,12,8X,F9.1,1X,F9.1,3X,F9.1,6X,F4.2,7X,F4.2,
272      6X,F5.1,5X,F5.1,2X,F9.1,1X,F12.1)
273      IF(N2.EQ.0) GO TO 352
274      WRITE(II,300)
275      351 CONTINUE
276      WRITE(II,353)
277      353 FORMAT(1H ,10H OTHER VRLE)
278      WRITE(II,354) (I,L(I),MODE(I),H(I),MEAN(I),VAR(I),I=NONE,NTWO)
279      354 FORMAT(1H ,3X,12,8X,F9.1,1X,F9.1,3X,F9.1,4X,F9.1,1X,F12.1)
280      352 CONTINUE
281      WRITE(II,295)
282      WRITE(II,355) TPUT,TPUT,TPUT
283      355 FORMAT(1H ,10H THROUGHPUT,3X,F9.1,13X,F9.1,44X,F9.1)
284      WRITE(II,300)
285      WRITE(II,360) XLOW,XHIGH,MSUM,VSUM
286      360 FORMAT(1H ,8H SUMS,5X,F9.1,13X,F9.1,44X,F9.1,3X,F10.1)
287      IF(N2.EQ.0) GO TO 410
288      WRITE(II,300)
289      WRITE(II,320)
290      WRITE(II,370)
291      370 FORMAT(1H ,10H OTHER VRLE,50X,3H CDF)
292      WRITE(II,300)
293      WRITE(II,380) (I,(VALINT(J,I),J=1,10),I=NONE,NTWO)
294      380 FORMAT(1H ,3X,12,20X,10(F6.4,3X))
295      WRITE(II,320)
296      %
297      % PRINT OUTPUT
298      %
299      410 CONTINUE
300      WRITE(II,320)

```

```

301 WRITE(11,490)
302 490 FORMAT(1H ,5X,11H O U T P U T)
303 WRITE(11,320)
304 WRITE(11,500)
305 500 FORMAT(1H ,AHINTERVAL,10X,5H RANGE,6X,22H INDEPENDENT BETA/OTHER,3X,*
306 20H DEPENDENT BETA/OTHER,3X,13H TOTAL UNIFORM,5X,19H INDEPENDENT UNIFORM)
307 WRITE(11,295)
308 WRITE(11,505)
309 505 FORMAT(1H ,36X,3HPDF,4X,3HCDF,12X,3HPDF,4X,3HCDF,10X,*
310 3HPDF,4X,3HCDF,12X,3HPDF,4X,3HCDF)
311 WRITE(11,300)
312 WRITE(11,510) (I,W(I),W(I+1),P1(I),C1(I),P2(I),C2(I),P3(I),C3(I),P4(I),C4(I),I=1,KK)
313 510 FORMAT(1H ,2X,12,4X,F9.1,3H - ,F9.1,5X,F6.4,1X,F6.4,9X,F6.4,1X,F6.4,7X,*
314 F6.4,2X,F6.4,8X,F6.4,2X,F6.4)
315 WRITE(11,310)
316 WRITE(11,520) MEAN1,MEAN2,MEAN3,MEAN4
317 520 FORMAT(1H ,AH*,MEAN*,27X,F9.1,13X,F9.1,12X,F9.1,13X,F9.1)
318 WRITE(11,530) MODE1,MODE2,MODE4
319 530 FORMAT(1H ,AH*,MODE*,27X,F9.1,13X,F9.1,34X,F9.1)
320 WRITE(11,540) VAR1,VAR2,VAR3,VAR4
321 540 FORMAT(1H ,12H*,VARIANCE*,20X,F12.1,10X,F12.1,9X,F12.1,10X,F12.1)
322 WRITE(11,545) STD1,STD2,STD3,STD4
323 545 FORMAT(1H ,17H*,STD DEVIATION*,17X,F10.1,12X,F10.1,11X,F10.1,12X,F10.1)
324 WRITE(11,295)
325 620 CONTINUE
326 WRITE(11,550) IC,L1,L2,L3,L4,U1,U2,U3,U4
327 550 FORMAT(1H ,2H*,12,23H*,CONFIDENCE INTERVAL*,AX,F9.1,1H,,12X,F9.1,1H,,
328 11X,F9.1,1H,,12X,F9.1,1H,,/36X,F9.1,13X,F9.1,12X,F9.1,13X,F9.1)
329 IF(IC.NE.90) GO TO 570
330 WRITE(11,295)
331 WRITE(11,560) PX1,PX2,PX3,PX4
332 560 FORMAT(1H ,20H*,PROB EXCEED MEAN*,22X,F3.2,19X,F3.2,18X,F3.2,19X,F3.2)
333 WRITE(11,300)
334 570 CONTINUE
335 WRITE(11,310)
336 WRITE(11,600)
337 600 FORMAT(1H ,20H ANOTHER CONFIDENCE INTERVAL?)
338 READ(JJ,*) IC
339 IF(IC.EQ.0) GO TO 610
340 CALL CI(MEAN1,IC,P1,KK,W,WIDTH,L1,U1)
341 CALL CI(MEAN2,IC,P2,KK,W,WIDTH,L2,U2)
342 CALL CI(MEAN3,IC,P3,KK,W,WIDTH,L3,U3)
343 CALL CI(MEAN4,IC,P4,KK,W,WIDTH,L4,U4)
344 GO TO 620
345 610 CONTINUE
346 WRITE(11,310)
347 WRITE(11,620)
348 622 FORMAT(1H ,23H PROB EXCEED SOME VALUE?)
349 READ(JJ,*) Z
350 IF(Z.EQ.0) GO TO 623
351 CALL CDF(P1,Z,KK,W,C1,PX1)
352 CALL CDF(P2,Z,KK,W,C2,PX2)
353 CALL CDF(P3,Z,KK,W,C3,PX3)
354 CALL CDF(P4,Z,KK,W,C4,PX4)
355 WRITE(11,624) Z,PX1,PX2,PX3,PX4
356 624 FORMAT(1H ,13H*,PROB EXCEED,F10.1,2H*,17X,F3.2,19X,F3.2,18X,F3.2,19X,F3.2)
357 GO TO 610
358 623 CONTINUE
359 WRITE(11,310)
360 WRITE(11,630)
361 630 FORMAT(1H ,24H ADDITIONAL OBSERVATIONS?)
362 READ(JJ,*) MSTAR
363 IF(MSTAR.EQ.0) GO TO 640
364 M=M+MSTAR
365 GO TO 160
366 640 CONTINUE
367 WRITE(11,310)
368 WRITE(11,650)
369 650 FORMAT(1H ,13H ANOTHER SEED?)
370 READ(JJ,*) KSD
371 IF(KSD.EQ.0) STOP
372 KSD2=KSD
373 WRITE(11,77)
374 READ(JJ,*) M
375 GO TO 140

```

```

1  SUBROUTINE CI(MEAN,IC,P,KK,W,WIDTH,LBU,UBD)
2  REAL W(1),MEAN,P(1),LWT1,LPROB,LBD,LWT2
3  DO 250 I=2,KK
4    J=I
5    IF(W(I).GE.MEAN) GO TO 260
6    250 CONTINUE
7    260 MM=0
8    CON=.5*(IC/100.)
9    LWT1=(MEAN-W(J-1))/(W(J)-W(J-1))
10   270 LPROB=LWT1*P(J-1)
11   K=1
12   IF(LPROB.GE.CON) GO TO 290
13   J1=J-1
14   DO 280 I=2,J1
15     K=I
16     LPROB=LPROB+P(J-I)
17     IF(LPROB.GE.CON) GO TO 290
18     280 CONTINUE
19     CON=2.*CON-LPROB
20     LBD=W(1)
21     GO TO 300
22     290 LWT2=(LPROB-CON)/P(J-K)
23     LBD=W(J-K)+LWT2*WIDTH
24     IF(MM.EQ.1) GO TO 300
25     *
26     300 RWT1=1.-LWT1
27     RPROB=RWT1*P(J-1)
28     LL=1
29     IF(RPROB.GE.CON) GO TO 320
30     J2=KK-J+2
31     DO 310 I=2,J2
32       LL=I
33       RPROB=RPROB+P(J+I-2)
34       IF(RPROB.GE.CON) GO TO 320
35     310 CONTINUE
36     CON=2.*CON-RPROB
37     UBU=W(KK+1)
38     MM=1
39     GO TO 270
40     320 RWT2=(RPROB-CON)/P(J+LL-2)
41     UBD=W(J+LL-1)-RWT2*WIDTH
42   330 RETURN

1  SUBROUTINE CDF(P,Z,KK,W,C,PROBX)
2  REAL P(1),W(1),C(1)
3  C(1)=P(1)
4  DO 10 I=2,KK
5    C(I)=C(I-1)+P(I)
6  10 CONTINUE
7  N=KK+1
8  DO 20 I=2,N
9    J=I
10   IF(W(I).GT.Z) GO TO 30
11   20 CONTINUE
12   30 A=(Z-W(J-1))/(W(J)-W(J-1))
13   PROBX=(1.-A)*P(J-1)
14   DO 40 I=J,KK
15     PROBX=PROBX+P(I)
16   40 CONTINUE
17   RETURN

```

```

1  SUBROUTINE DGG32(XL,XU,Y,ALPHA,BETA)
2  DOUBLE PRECISION XL,XU,Y,A,B,C,FCT,ALPHA,BETA
3  FCT(X)=X**ALPHA*(1.-X)**BETA
4  A=.5D0*(XL+XU)
5  B=XU-XL
6  C=.49863193092474078D0*B
7  Y=(.35093050047350483D-2)*(FCT(A+C)+FCT(A-C))
8  C=.49280575 7/2634170D*B
9  Y=Y+(.6137197365452835D-2)*(FCT(A+C)+FCT(A-C))
10 C=.48238112779375322D0*B
11 Y=Y+(.12696032654631030D-1)*(FCT(A+C)+FCT(A-C))
12 C=.46745303796860984D0*B
13 Y=Y+(.17136931456510717D-1)*(FCT(A+C)+FCT(A-C))
14 C=.44816057783302606D0*B
15 Y=Y+(.214179490111 334D-1)*(FCT(A+C)+FCT(A-C))
16 C=.4246838068602849 D0*B
17 Y=Y+(.2549 02963118 08D-1)*(FCT(A+C)+FCT(A-C))
18 C=.39724189798397120D0*B
19 Y=Y+(.29342046739267774D-1)*(FCT(A+C)+FCT(A-C))
20 C=.36009105937014484D0*B
21 Y=Y+(.3291111138210923D-1)*(FCT(A+C)+FCT(A-C))
22 C=.33152213346510760D0*B
23 Y=Y+(.36172897054-24253D-1)*(FCT(A+C)+FCT(A-C))
24 C=.29385787862038116D0*B
25 Y=Y+(.39096947893535153D-1)*(FCT(A+C)+FCT(A-C))
26 C=.25344985446611470D0*B
27 Y=Y+(.41655962113473378D-1)*(FCT(A+C)+FCT(A-C))
28 C=.2106756100531767D0*B
29 Y=Y+(.43826046502201906D-1)*(FCT(A+C)+FCT(A-C))
30 C=.16593430114106382D0*B
31 Y=Y+(.458869392478-1942D-1)*(FCT(A+C)+FCT(A-C))
32 C=.11964368112606854D0*B
33 Y=Y+(.46922197540402283D-1)*(FCT(A+C)+FCT(A-C))
34 C=.7223598079139825D-1*B
35 Y=Y+(.47819360139637430D-1)*(FCT(A+C)+FCT(A-C))
36 C=.24153832843869158D-1*B
37 Y=B*(Y+(.4827004425736390D-1)*(FCT(A+C)+FCT(A-C)))
38 RETURN
39 END

```

```

1  SUBROUTINE STAT(SUM,M,KK,P,W,SUMSQ,MEAN,MODE,VAR,STD)
2  REAL MEAN,MODE,P(1),W(1)
3  MEAN=SUM/M
4  MODE=0.
5  DO 10 I=1,KK
6  IF(P(I).GT.MODE) MODE=P(I)
7  IF(P(I).EQ.MODE) IMODE=I
8  10 CONTINUE
9  MODE=.5*(W(IMODE)+W(IMODE+1))
10 VAR=(SUMSQ-M*MEAN**2.)/(M-1)
11 STD=SQRT(VAR)
12 RETURN

```

Optimal Sample Size

In most sampling procedures the larger the sample the closer the sample distribution approximates the true distribution. But larger samples are more expensive to generate than smaller ones. The simple experiments described below represent an attempt to determine an optimal sample size for Program SPET -- the smallest sample size that will ensure "reasonable" accuracy in the sampling procedure.

A statistic called the K statistic in this study is used in the search for the optimal sample size. It is defined as

$$K = \sum_{i=1}^N (o_i - e_i)^2$$

where N = the number of class intervals

o_i = the number of observations occurring in the i th interval $\div M$.

M = the total number of observations

e_i = the number of observations expected in the i th interval (given that the process generating the observations is following a particular statistical distribution) $\div M$.

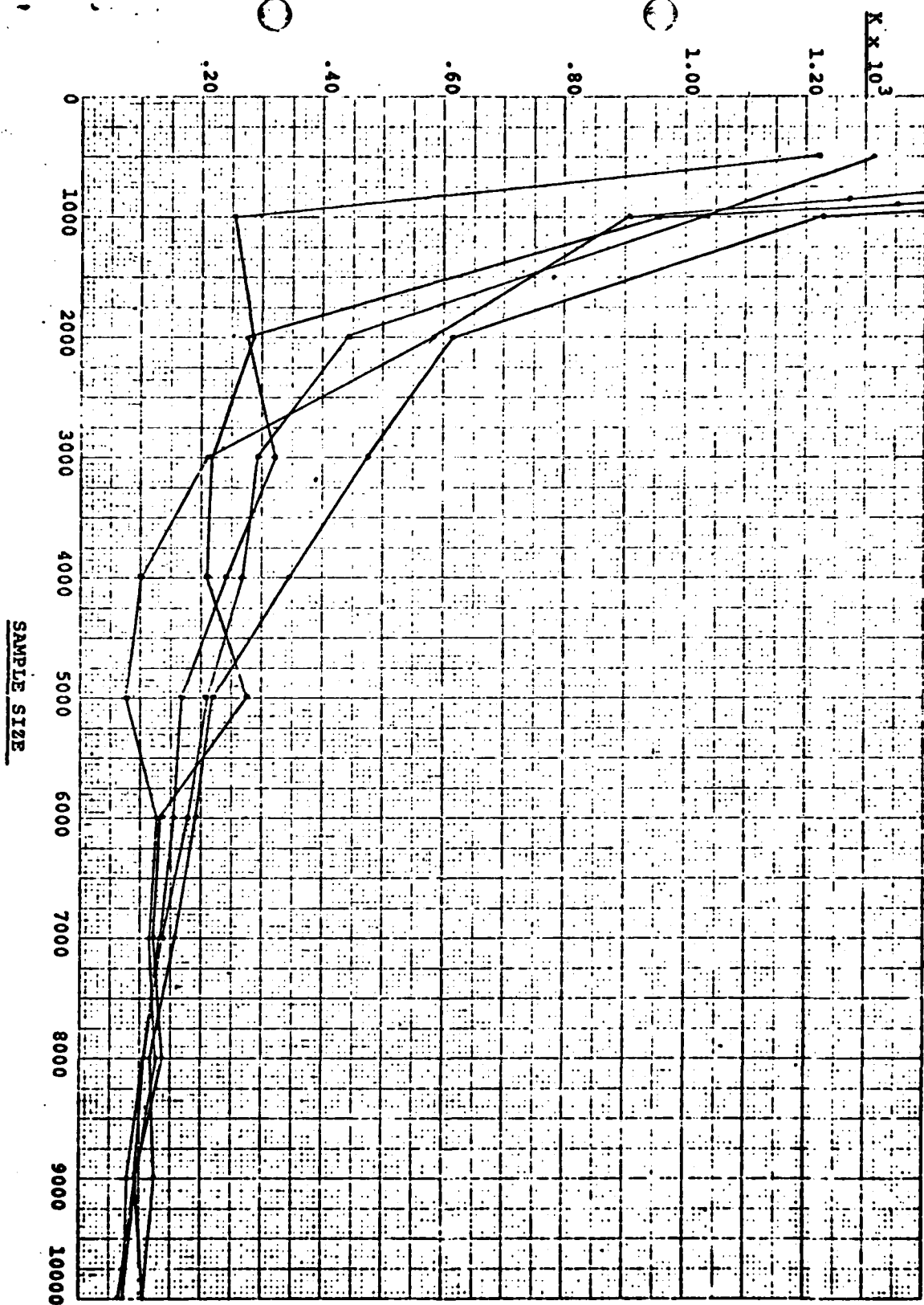
The K statistic is the sum of squared deviations of the observed probabilities from the expected probabilities of each class interval. As the size of a randomly drawn sample is increased, the K statistic decreases in value until $\lim_{M \rightarrow \infty} K = 0$.

$M \rightarrow \infty$

The first experiment consists of using five randomly selected seeds with the uniform random number generator used by Program SPET to generate five sequences of K statistics. Each sequence contains a K statistic for sample sizes 500, 1000, 2000, ..., 10000. These K statistics are plotted in Figure C-1 on page C-16. Note how the sequences converge at sample size 6000. It appears that this might be the optimal sample size. Can one expect a sample size of 6000 to ensure "reasonable" accuracy in the sampling procedure?

FIGURE C-1

FIVE SEQUENCES OF K STATISTICS AS A FUNCTION OF SAMPLE SIZE



The accuracy of the Monte Carlo sampling procedure used in Program SPET, for purposes of this inquiry, is measured in terms of the percent deviation of certain statistics from their true values. The second experiment is an attempt to measure the accuracy of Program SPET for various sample sizes. It consists of using the same five seeds selected in the first experiment to draw five sequences of samples from a uniform distribution. Each sequence contains samples of size 500, 1000, 2000, 3000, 6000, and 9000. The mean and the lower and upper confidence limits of the 90% confidence interval are noted from the output of Program SPET and the percent deviation of these statistics from their population values is computed. Then the maximum of the absolute value of the deviations is selected for each statistic in every sample size and plotted in Figure C-2 on page C-18. Note that the rate of decrease in the error (maximum percent deviation) of these statistics is rapid in the range of the sample sizes 500 to 2000, slowing somewhat after sample size 2000.

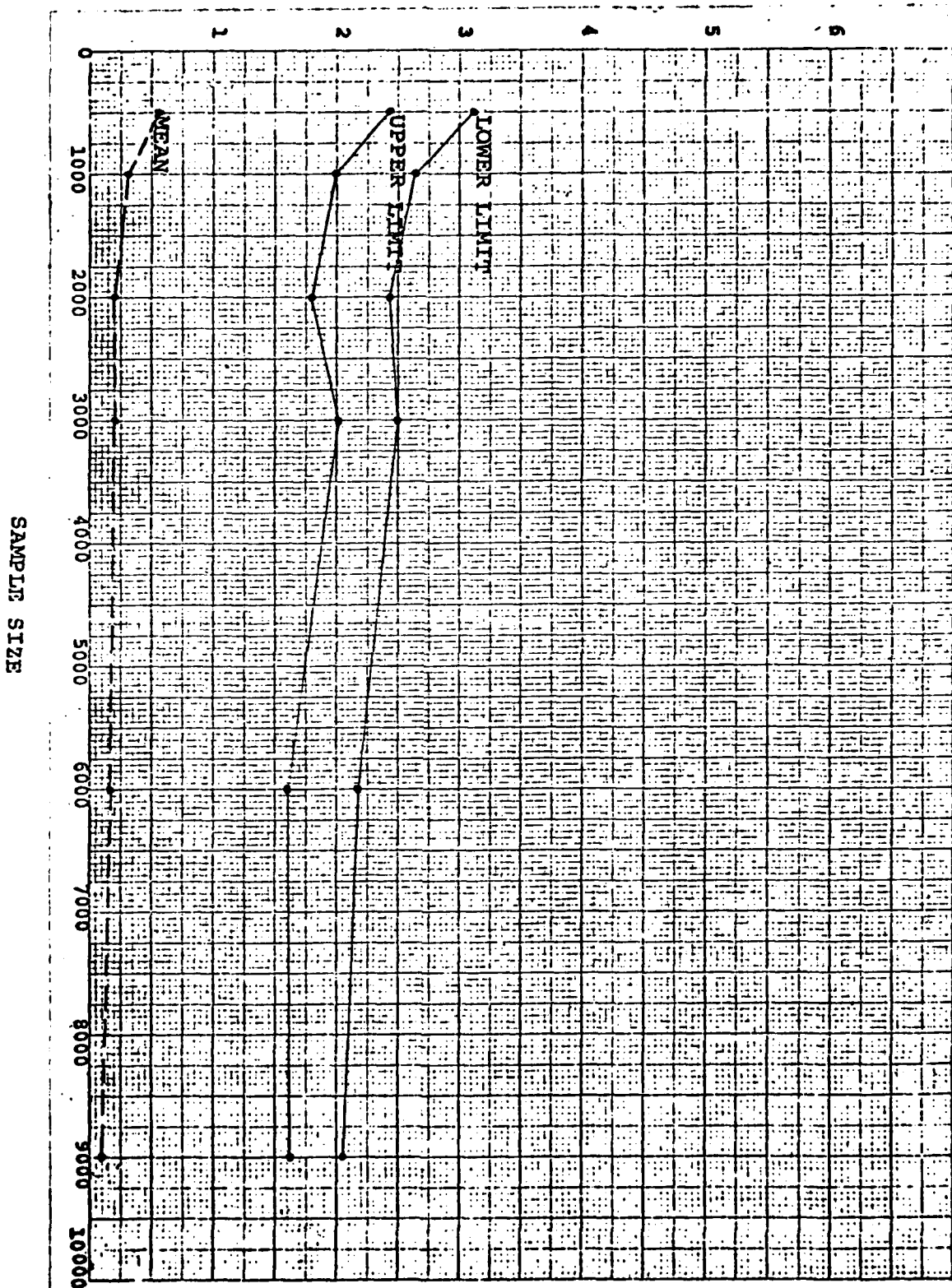
Consider the error in sample size 2000. Would the expectation of a deviation of at most .21 percent in the mean and 2.44 and 1.81 percent in the lower and upper limits of the confidence interval respectively, be "reasonable?" The authors would answer affirmatively. Reasonableness is subjective. It is felt that the accuracy of sample size 6000 is not enough better than that of sample size 2000 to warrant incurring the increased cost of generating an additional 4000 observations.

Much greater confidence could be placed in these tentative observations if, instead of five sequences, 30, 40 or more sequences had been generated.^{4/} But even the results of the five sequences permit a more confident choice of sample size than no experimentation at all.

^{4/} Along with a greater number of sequences one might have repeated experiment two using one or two representative beta distributions in addition to the uniform distribution used above.

MAXIMUM PERCENT DEVIATION (ERROR)

MAXIMUM PERCENT DEVIATION (ERROR) OF CERTAIN STATISTICS
FROM THEIR POPULATION VALUES AS A FUNCTION OF SAMPLE SIZE



REFERENCES

1. Convolution of Inverse Beta Distributions by a Sampling Technique. Bethesda, MD: Mathematica, Inc., 1971.
2. Dienemann, Paul F., Estimating Uncertainty Using Monte Carlo Techniques. Santa Monica, CA: The RAND Corp., RM-4854-PR, 1966.
3. Hillier, F. S. and Lieberman, G. J., Introduction to Operations Research. San Francisco, CA: Holden-Day, Inc., 1967.
4. Kamméner, J. T., ASW Force Level Study - Equipment Readiness: Models, Computer Simulation and Results. Washington, DC: Office of the Chief of Naval Operations, 1968.
5. Larson, Harold J., Introduction to Probability Theory and Statistical Inference. New York: John Wiley and Sons, Inc., 1969.
6. Lindgren, Bernard W., Statistical Theory. 2nd ed. New York: The MacMillan Co., 1968.
7. MacCrimmon, K. R. and Ryavec, C. A., An Analytical Study of the PERT Assumptions. Santa Monica, CA: The RAND Corp., RM-3408-PR, 1962.
8. Schaefer, Donald F., et. al., A Monte Carlo Simulation Approach to Cost-Uncertainty Analysis. McLean, VA: Research Analysis Corp., RAC-TP-349, 1969.
9. Sobel, S., A Computerized Technique to Express Uncertainty in Advanced System Cost Estimates. Bedford, MA: The MITRE Corp., TM-3728, 1963.
10. Sutherland, William H., Fundamentals of Cost Uncertainty Analysis. McLean, VA: Research Analysis Corp., RAC-CR-4, 1971.
11. _____, A Method for Combining Asymmetric Three-Value Predictions of Time or Cost. McLean, VA: Research Analysis Corp., RAC-P-65, 1972.