

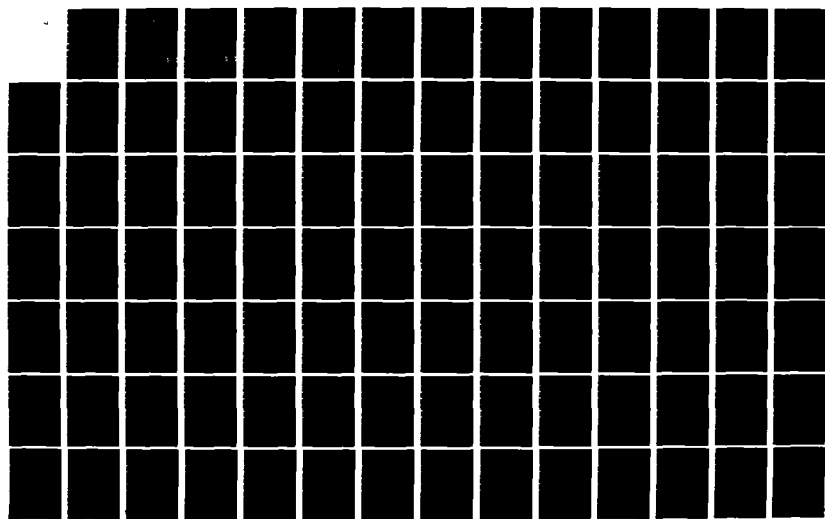
AD-A138 098

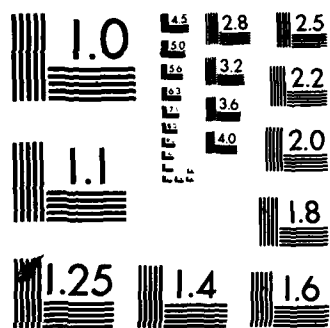
MODIFIED KOLMOGOROV-SMIRNOV ANDERSON-DARLING AND
CRAMER-VON MISES TESTS F. (U) AIR FORCE INST OF TECH
WRIGHT-PATTERSON AFB OH SCHOOL OF ENGI... J D YODER
16 DEC 83 AFIT/GOR/ENC/83D-7 F/G 12/1

1/2

UNCLASSIFIED

NL



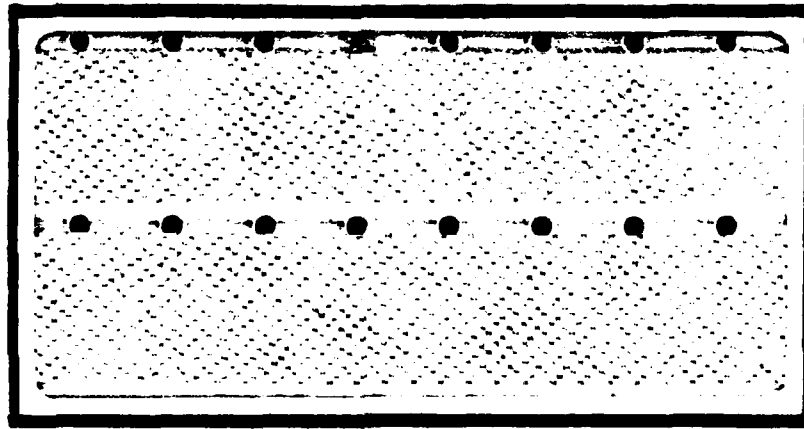


MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS-1963-A

AD A138098



①



DISTRIBUTION STATEMENT A

Approved for public release
Distribution Unlimited

DTIC
ELECTE
FEB 22 1984

S

B

DEPARTMENT OF THE AIR FORCE
AIR UNIVERSITY

AIR FORCE INSTITUTE OF TECHNOLOGY

Wright-Patterson Air Force Base, Ohio

DTIC FILE COPY

84 02 21 108

AFIT/GOR/ENC/83D-7

MODIFIED KOLMOGOROV-SMIRNOV, ANDERSON-DARLING
AND CRAMER-VON MISES TESTS FOR THE
LOGISTIC DISTRIBUTION WITH UNKNOWN
LOCATION AND SCALE PARAMETERS

THESIS

John D. Yoder
First Lieutenant, USAF

AFIT/GOR/ENC/83D-7

DTIC
ELECTE
S FEB 22 1984 **D**
B

Approved for public release; distribution unlimited

AFIT/GOR/ENC/83D-7

MODIFIED KOLMOGOROV-SMIRNOV, ANDERSON-DARLING
AND CRAMER-VON MISES TESTS FOR THE
LOGISTIC DISTRIBUTION WITH UNKNOWN
LOCATION AND SCALE PARAMETERS

THESIS

Presented to the Faculty of the School of Engineering
of the Air Force Institute of Technology
Air University
In Partial Fulfillment of the
Requirements for the Degree of
Master of Science in Operations Research

John D. Yoder, B.S.
First Lieutenant, USAF

December 1983

Approved for public release; distribution unlimited

Preface

Goodness-of-fit tests are developed for the Logistic distribution when the location and scale parameters are unknown. The Kolmogorov-Smirnov, Anderson-Darling and Cramer-von Mises statistics are used to develop tables of critical values to be used in goodness-of-fit hypothesis testing. A power study is conducted to compare the Kolmogorov-Smirnov, Anderson-Darling and Cramer-von Mises goodness-of-fit tests.

I would like to thank my advisor, Capt. Brian Woodruff, whose continual help and encouragement were instrumental to the successful completion of my thesis.

In addition, I would like to thank my readers, Dr. Albert H. Moore and Dr. James Dunne; their advice and guidance was very helpful.

Finally, I would like to express my appreciation to my classmates in general, and Lt. Jim Keffer in particular, for their help and encouragement during the preparation of my thesis.

John D. Yoder



Accession For	
NTIS GRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By	
Distribution/	
Availability Codes	
Dist	Avail and/or Special
A-1	

Contents

	page
Preface	ii
List of Figures	vi
List of Tables	vii
Abstract	viii
I. Introduction	1
Background	2
Empirical Distribution Function	3
Unknown Parameters	5
The Kolmogorov-Smirnov Statistic	5
The Anderson-Darling Statistic	6
The Cramer-von Mises Statistic	7
Problem Statement	8
Objectives	9
II. The Logistic Distribution	10
History and Development	10
The Logistic Distribution	12
Application of the Logistic Density Function	14
III. Methodology	17
Steps in the Monte Carlo method	17
Generation of random logistic deviates	19
Maximum likelihood estimates of the logistic parameters	20
The Secant Method	23

	page
Ordering the deviants	26
The Hypothesized Distribution Function	27
Determining the critical values of the goodness-of-fit tests	27
Power Comparison	32
IV. Use of Tables	38
Example	39
V. Discussion of Results	42
Presentation of the Table of Critical Values	42
Analysis of Critical Value Tables	44
Validation of Computer Program	49
Sensitivity Analysis of Program	52
Presentation of the Power Study	53
VI. Conclusions and Recommendations	57
Conclusions	57
Recommendations	57
Bibliography	59
APPENDIX A: Table of the Kolmogorov-Smirnov Critical Values for the Logistic Distribution with unknown location and scale parameters	64
APPENDIX B: Table of the Anderson-Darling Critical Values for the Logistic Distribution with unknown location and scale parameters	66

APPENDIX C:	Table of the Cramer-von Mises Critical Values for the Logistic Distribution with unknown location and scale parameters	68
APPENDIX D:	Tables of Power Comparison Between the Kolmogorov-Smirnov, Anderson-Darling and Cramer-von Mises goodness-of-fit tests for sample sizes of 5,10,15,20,25 and 30 at the .05 significance level	70
APPENDIX E:	Computer programs	77
	Critical Value program	78
	Power Study program	89

List of Figures

<u>Figure</u>	<u>Page</u>
1. Typical hazard function in the area of life testing	15
2. Flow chart of methodology	18
3. Plotting position vs. Statistics	31
4a. Gamma distribution: $K=3,6,9$	34
4b. Gamma distribution: $K=15,30$	35
5. Weibull distribution: $K=3$	37
6. Kolmogorov-Smirnov Critical values by sample size . .	45
7. Anderson-Darling Critical values by sample size . . .	46
8. Cramer-von Mises Critical values by sample size . . .	47
9. General distribution shape of test statistics	48

List of Tables

<u>Table</u>	<u>page</u>
I. Example calculations	40
II. Kolmogorov-Smirnov Comparison	
with Stephens values	50
III. Anderson-Darling Comparison	
with Stephens values	51
IV. Cramer-von Mises Comparison	
with Stephens values	51
V. Critical values for the Kolmogorov-Smirnov statistic .	65
VI. Critical values for the Anderson-Darling statistic .	67
VII. Critical values for the Cramer-von Mises statistic .	69
VIII. Power comparison, sample size 5	71
IX. Power comparison, sample size 10	72
X. Power comparison, sample size 15	73
XI. Power comparison, sample size 20	74
XII. Power comparison, sample size 25	75
XIII. Power comparison, sample size 30	76

Abstract

The method of maximum likelihood is used to determine invariant estimates of the unknown location and scale parameters of a sample from the Logistic distribution. The partial derivatives of the likelihood function can not be solved explicitly, therefore the Secant method is used to iteratively determine the roots of the partial derivatives. Using these estimates, modified Kolmogorov-Smirnov, Anderson-Darling and Cramer-von Mises statistics are calculated for a given sample. This procedure is repeated 5000 times for sample sizes of $n = 5(5)30$. The 80th, 85th, 90th, 95th and 99th percentiles of the distribution of each statistic, for each sample size, is then calculated. These values are then used to generate tables of critical values for the Logistic distribution with unknown location and scale parameters. A power comparison between the three tests is performed using samples from various distributions.

The Secant method requires "good" initial estimates of the parameters in order to converge. This thesis uses the sample mean and standard deviation as initial estimates. In four of the total 30,000 samples used, these initial estimates did not allow convergence. While discarding these samples biases the theoretical results, it was determined that discarding these samples would not bias the numerical results. This does however place a constraint on using the Secant

method with respect to obtaining maximum likelihood estimates of the parameters. The power of these tests for non-symmetrically convex distributions is very good. However, for symmetrically convex distributions, the power ranges from moderate to only slightly more than the significance level.

MODIFIED KOLMOGOROV-SMIRNOV, ANDERSON-DARLING AND
CRAMER-VON MISES TESTS FOR THE LOGISTIC DISTRIBUTION
WITH UNKNOWN LOCATION AND SCALE PARAMETERS

I. INTRODUCTION

The Air Force is highly interested in the reliability and maintainability of proposed systems. When no historical data exists, statistical and probabilistic measures are often the only approaches possible in gathering meaningful information to aid the decision maker. The mean time to failure and the failure rate of a proposed system are usually unknown but important considerations in the decision making process.

Various methods can be used to collect, for example, time to failure data from a experimental or prototype system. The data can then be compared to a theoretical probability distribution. How well the distribution of the experimental data matches the theoretical distribution is known as a "goodness-of-fit test". If such tests show that the distribution of the experimental data "fits" the theoretical distribution well, the hypothesized theoretical distribution can be used in simulation modeling, for example, to predict the failure rate of the proposed system. The Gamma distribution has often been used in such studies, and its hazard function approaches a constant value. The hazard function of the Logistic distribution also approaches a constant value, as time approaches infinity, and is therefore a useful alternative in some reliability and life-testing situations.

BACKGROUND

There are several classical goodness-of-fit tests, such as the Chi-Square test, the Kolmogorov-Smirnov test (KS), the Anderson-Darling test (A^2), and the Cramer-von Mises test (W^2). Of these tests, the two most popular are the Chi-Square and the Kolmogorov-Smirnov tests. The Chi-square test compares the observed frequencies of the empirical distribution to the expected frequencies of the hypothesized distribution. However, because it groups observations (with a minimum of five observations per cell) it is restricted to larger samples ($n \geq 73$). The KS test, on the other hand, has no such restriction. This test uses as a measure of fit the absolute difference between the empirical distribution and the hypothesized distribution. Because of this, no grouping of data is required and smaller sample sizes can be accommodated. The classical KS goodness-of-fit test is valid to test whether a set of observations comes from a completely specified distribution.

However, in practical applications the distribution is seldom fully specified. In cases where the parameters must be estimated from the sample, the Chi-square test is easily adjusted by reducing the number of degrees of freedom by the number of parameters estimated. The KS test can also be modified to consider the case where the parameters are estimated from the sample data. H.W. Lilliefors developed a modified KS goodness-of-fit test for the Normal distribution (30) in 1967 and for

the Exponential distribution (31) in 1969. R. Cortes developed a modified KS test in 1980 (16) to be used with the Gamma and Weibull distributions. J.G. Bush, in 1981, developed modified A^2 and W^2 tests to be used with the Weibull distribution (13). In 1982 P.J. Viviano developed modified KS, A^2 and W^2 tests for the Gamma distribution (42). The KS test, as well as the A^2 and W^2 tests, were modified for the Uniform, Normal, Laplace, Exponential and Cauchy distributions by Green and Hegazy (18) in 1976. In 1981, Koutrouvelis and Kellermeier developed a goodness-of-fit test based on the empirical characteristic function when the characteristic function is a member of a specified parametric family of such functions (28). Masaaki, Hiroshi and Shigeo, in 1980, developed a goodness-of-fit test for the extreme value distribution based on the entropy of the sample data (32).

EMPIRICAL DISTRIBUTION FUNCTION

A class of statistics based on a comparison between the theoretical cumulative distribution function $F(x)$ and the sample cumulative distribution function $S(x)$ is generally called empirical distribution function (EDF) statistics. Historically, EDF statistics are used in cases where parameters are estimated from the sample observations. The EDF of a random sample is defined as

$$S(x) = \frac{\text{number of values } \leq x}{\text{total number of values}} \quad (1)$$

When there are n observations in the sample, $S(x)$ is a step function with $1/n$ jumps at each order statistic of the sample (19:73). When the n observations are arranged in ascending order, $S(x)$ is defined by Eq (2)

$$S(x) = \begin{cases} 0 & x \leq x_1 \\ i/n & x_i \leq x \leq x_{i+1} \quad i=1,2,\dots,n-1 \\ 1 & x \geq x_n \end{cases} \quad (2)$$

Since $S(x)$ yields the proportion of the sample less than or equal to x , it is a good estimate of the hypothesized cumulative distribution function $F(x)$. It should be noted that because the cumulative distribution is not fully specified this thesis uses a modified form of the EDF statistic. An estimated distribution function is used whose parameters are derived from the observed sample.

UNKNOWN PARAMETERS

In general, EDF tests are valid whenever the hypothesized distribution is fully specified. The probability integral transformation converts the values of a completely specified cumulative distribution to ordered values from a uniform distribution over the interval zero to one (17:184). However in practice the distribution is seldom fully specified. Hence, the probability integral transform, by itself, is not enough help. David and Johnson (17) have however shown that when location and scale are the parameters being estimated, the cumulative distribution of EDF statistics depends on the functional form of the distribution rather than the estimated parameters. This allows the probability integral transform to remain valid in certain cases when the distribution is not completely specified. It is this quality, coupled with invariant estimates of the parameters, that allows the generation of valid critical value tables dependent only on n and the significance level (α).

THE KOLMOGOROV-SMIRNOV STATISTIC

The Kolmogorov-Smirnov statistic is defined as the absolute difference between $F(x)$ and $S(x)$. Yet, since we are interested in the greatest discrepancy between the distributions, our test statistic is

$$KS = \sup_x | F(x) - S(x) | \quad (4)$$

In effect, this measures the superior of the absolute vertical discrepancy between the hypothesized distribution and the empirical distribution (15:346).

THE ANDERSON-DARLING STATISTIC

Goodness-of-fit tests based on the difference between empirical and hypothesized distributions almost always have smaller discrepancies in the tails of the distribution (40). To overcome this, several things can be done. The most popular is to weight the squared differences between distributions. The A^2 statistic is an example of such an approach. It is based on a nonnegative, weighted average of the squared discrepancy. That is

$$A^2 = n \int_{-\infty}^{\infty} [S(x) - F(x)]^2 \theta[F(x)] dx \quad (5)$$

where

$$\theta[F(x)] = [(F(x))(1-F(x))]^{-1}$$

In computational form

$$A^2 = -n - \frac{1}{n} \sum_{j=1}^n (2j-1) [\ln F(x_j) + \ln(1-F(x_{n+1-j}))] \quad (6)$$

In effect, this accentuates the difference between $F(x)$ and $S(x)$ in the tails of the distribution (40).

THE CRAMER-VON MISES STATISTIC

Another example of this approach is the Cramer-von Mises statistic. In this case the weighting function $\theta[F(x)]$ equals one. This defines the statistic as

$$W^2 = \int_{-\infty}^{\infty} [S(x) - F(x)]^2 dx \quad (7)$$

with a computational form of

$$W^2 = \frac{1}{12n} + \sum_{j=1}^n [F(x_j) - ((2j-1)/2n)]^2 \quad (8)$$

In effect, this equally accentuates the differences between $F(x)$ and $S(x)$ (40).

PROBLEM STATEMENT

Application of a particular distribution to a problem often is hampered in one of two ways. Either there is a lack of knowledge about the parameters of a specific distribution or there is not a method to easily test if a set of observations are a random sample from a specific distribution. The former problem is overcome by theoretical investigations of the parameters and characteristics of the distribution. The latter problem is overcome, generally, by development of a general test statistic or a table of critical values used in goodness-of-fit hypothesis testing for various sample sizes and specific parameters.

Numerous investigations of the parameters of the Logistic distribution have been accomplished. These investigations have used various methods. For example, least squares, maximum likelihood and best linear unbiased estimates of the parameters have all been accomplished. Based on asymptotic distribution theory, Stephens has developed critical values to apply the logistic distribution to goodness-of-fit hypothesis testing (41). However, his test statistics are highly modified for the asymptotic theory to hold true. No known effort has been done, based on finite distribution theory, to easily apply the Logistic distribution to goodness-of-fit hypothesis testing. Because of its applicability in reliability, there is a need to develop such a set of critical values tables for various sample sizes when the location and scale parameters are esti-

mated from the sample.

OBJECTIVES

This thesis has the following objectives:

1. Based on finite distribution theory, generate and document modified Kolmogorov-Smirnov, Anderson-Darling, and Cramer-von Mises rejection tables for the two parameter Logistic distribution where the location and scale parameters are unknown.
2. Conduct a power comparison between the Kolmogorov-Smirnov, Anderson-Darling and Cramer-von Mises goodness-of-fit tests.

II. THE LOGISTIC DISTRIBUTION

History and Development

Although the name might imply it, the logistic distribution does not have any special association with the fields of supply and maintenance or the function of logistical support. The distribution initially was used in a study of population growth by Verhulst in 1845. In 1920, Pearl and Reed renewed interest in the distribution with their use of it in a study of population growth in the United States (33). In the mid-1940's Dr. Berkson began using the distribution with respect to studies in the biological sciences (5). He established applications for the distribution in the study of autocatalysis, electro-chemical reactions and biochemistry. Reed and Berkson used the logistic distribution in physiochemical phenomena studies in 1929 (36). Bioassay applications were established by Wilson and Worchester (43) in 1943 and by Berkson between 1944 and 1957 (5,6,7). In 1958, Birnbaum used the distribution in a study of mental ability (10). Recently, in 1981, Leach (29) used the logistic distribution in a study on the original subject of population growth.

The logistic distribution was introduced to the reliability and maintainability field when Plackett used it in a study of life-test data (34). Shah, in 1965, showed its applicability in psychometrics, a field of interest to the reliability engineer (39). Bain (4) showed that, like the gamma distribution, the hazard function of the logistic dis-

tribution approaches a constant and is therefore a useful alternative in reliability and life testing situations.

Much work has also been done with respect to estimation of the parameters of the logistic distribution. Pearl and Reed (33), Schultz (38) and Berkson (5) all obtained least square estimates of the parameters. Maximum likelihood estimation of the parameters has been done by Wilson and Worchester (43), Berkson (7), and Harter and Moore (24). Berkson and Hodges developed a minimax estimator (8). The minimum chi-square technique was also used by Berkson to estimate the parameters. Plackett (34), Kjelsberg (27) and Gupta, Qureishi and Shah (21) have all developed best linear unbiased estimates for complete and censored samples, while Plackett developed linear estimates from censored data (35). Beyer investigated the conditional estimation of the scale parameter using selected order statistics in 1966 (9). Also using selected order statistics, Richardson investigated simultaneous linear estimation of the location and scale parameters (37).

Exact moments of the order statistics of the logistic distribution have also been developed. Birnbaum (10), Birnbaum and Dudman (11), Plackett (34), and Gupta and Shah (22) have all worked on the exact moments of the distribution. Shah has tabled variances and covariances of logistic order statistics for sample sizes up to 10 (39). Gupta, Qureishi and Shah extended this through a sample size of 25 (21). Harter and Moore established the asymptotic variances and covariances of

the maximum likelihood estimators in 1967 (24) and Bain (4) presented tabled percentage points for the maximum likelihood estimators.

The Logistic Distribution

The shape of the logistic distribution is convex and symmetric about the mean and similar to that of the normal distribution. The notable difference is that the tails are relatively thick, more like that of the exponential distribution. The location parameter or mean (μ) relates the point of symmetry, the median and the mode of the logistic density function. The standard deviation (σ) measures the relative dispersion of the distribution along the axis of the independent variable.

The logistic distribution with mean μ and standard deviation σ is defined as

$$F(x) = [1 + \exp[-\pi(x-\mu)/\sigma\sqrt{3}]]^{-1} \quad (9)$$

with the scale parameter defined as $\sigma\sqrt{3}/\pi$.

The density function is expressed as

$$f(x) = \frac{\pi \exp[-\pi(x-\mu)/\sigma\sqrt{3}]}{\sigma\sqrt{3}[1 + \exp[-\pi(x-\mu)/\sigma\sqrt{3}]]^2} \quad (10)$$

for $-\infty < x < \infty$
 $-\infty < \mu < \infty$
 $0 < \sigma < \infty$

The location parameter μ indicates the value of x at which failures are most likely to begin occurring.

The moments of the logistic distribution can be found using

$$E(x^k) = \int_{-\infty}^{\infty} x^k \pi \exp[-\pi(x-\mu)/\sigma\sqrt{3}] \\ [\sigma\sqrt{3}[1+\exp[-\pi(x-\mu)/\sigma\sqrt{3}]]^{-2}]^{-1} dx \quad (11)$$

or more easily by the moment generating function (mgf)

$$E(\exp[xt]) = \int_{-\infty}^{\infty} \exp[xt] \pi \exp[-\pi(x-\mu)/\sigma\sqrt{3}] \\ [\sigma\sqrt{3}[1+\exp[-\pi(x-\mu)/\sigma\sqrt{3}]]^{-2}]^{-1} dx \quad (12)$$

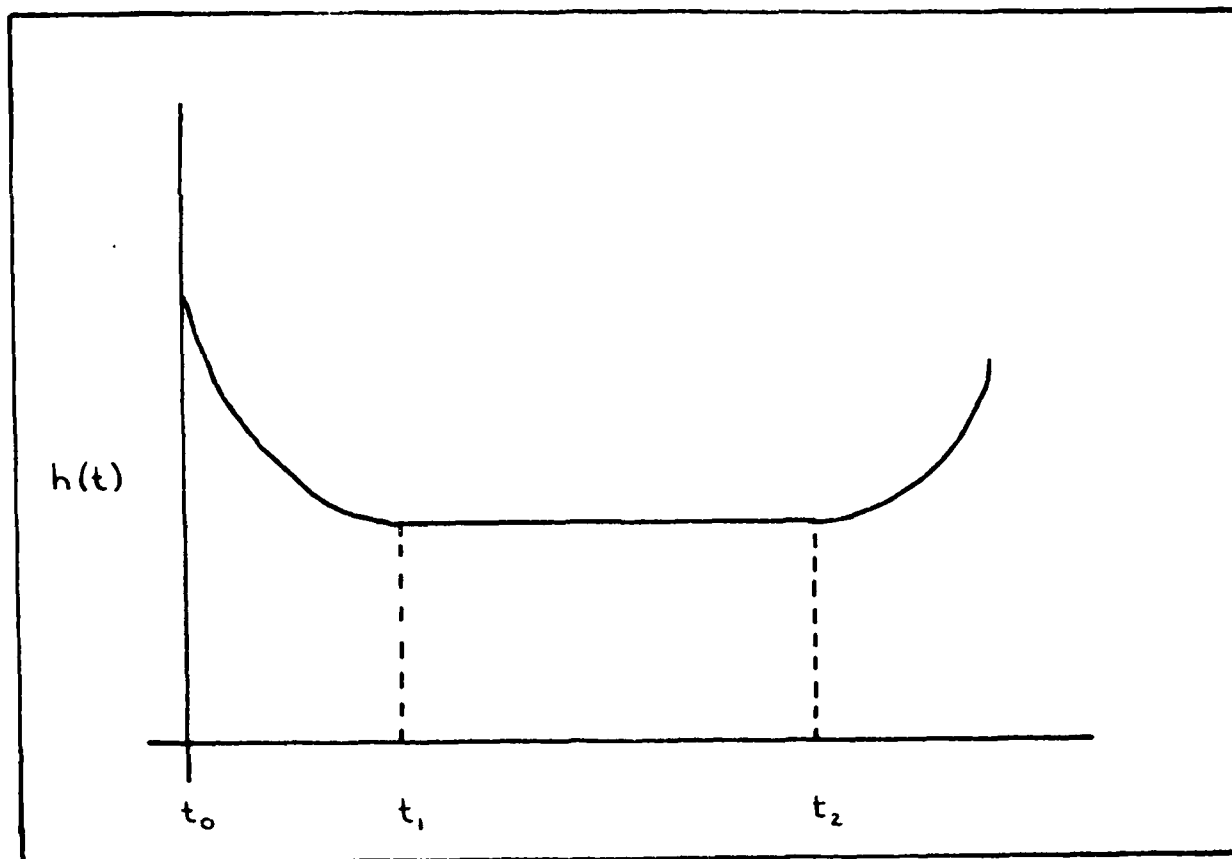
From the reduced variant, $y = (x-\mu)/\sigma$, Gumbel (20) has shown that the mgf can be expressed as

$$M_x(t) = \Gamma(1+t/c)\Gamma(1-t/c) \quad (13)$$

which is a Beta function where $c = \pi / \sqrt{3}$.

APPLICATION OF THE LOGISTIC DENSITY FUNCTION

If a random variable represents the lifetime or time to failure of a unit, then the study of that variable is said to be in the area of life-testing or reliability theory. The probability that a unit survives until time x is called the reliability of that unit at time x and denoted as $R(x) = 1 - F(x)$. On the other hand, the hazard function of the unit may be interpreted as the instantaneous failure rate of the unit (4:42). It is often more informative to consider the hazard function of a model than to look at the shape of the pdf or cdf directly. A typical hazard function in the area of life testing is a U-shaped or bathtub shaped curve (Fig 1).



Typical hazard function

Figure 1

Normally, when a unit is placed in service it first goes through a period where the frequency of failures is decreasing. That is, as manufacturing defects are overcome, the reliability of the unit improves with age. The unit then goes through a period where failures are more or less random at a constant rate. As the unit begins to wear out or deteriorate the failures become more frequent and the failure rate increases.

The hazard function for the logistic distribution is

expressed as

$$h(x) = \frac{f(x)}{1 - F(x)} = \frac{\pi F(x)}{\sigma\sqrt{3}} \quad (14)$$

$h(x)$ is an increasing function of x and it is easy to see that $h(x)$ approaches $\pi/\sigma\sqrt{3}$ as $x \rightarrow \infty$. This property may be well suited to the analysis of certain systems. For example, the hazard function of a system may not be characterized by the normal bath-tub curve. When placed in operation, the system may almost immediately begin to wear out with an increasing failure rate. The failure rate may then approach a constant, and continue thus for quite some time. In situations such as this the logistic function can be used to advantage.

III. METHODOLOGY

This thesis presents tabled critical values of the modified Kolmogorov-Smirnov, Anderson-Darling and Cramer-von Mises goodness-of-fit tests for the logistic distribution with unknown location and scale parameters. This chapter discusses the procedures used in this thesis. First, an outline and flow chart of the Monte Carlo method is presented. Next, there is a description of the steps followed in this method. Finally, the chapter concludes with a discussion of the power study between these test statistics.

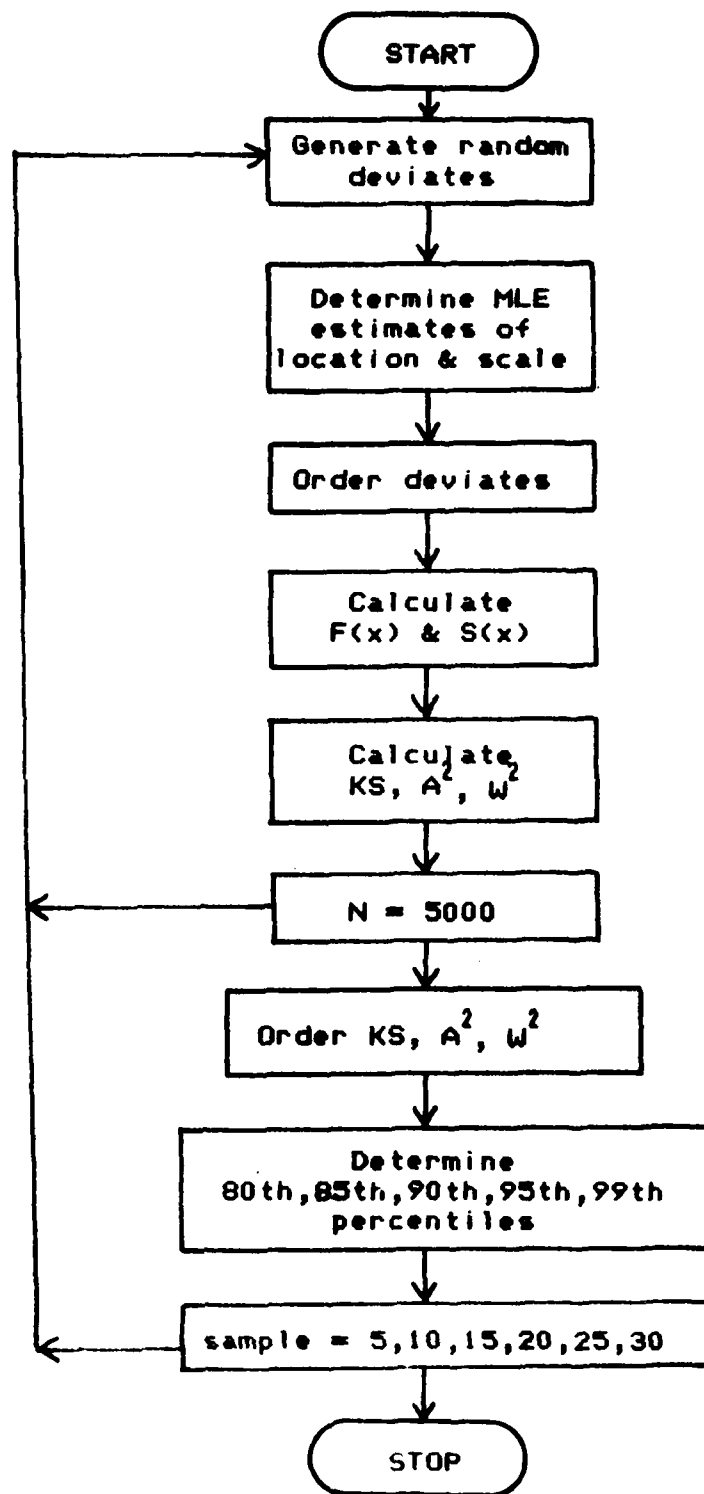
STEPS IN THE MONTE CARLO METHOD

The following eight steps outline the Monte Carlo method used to calculate the critical values for the modified KS, A^2 and W^2 goodness-of-fit tests. The flow chart in Figure 2 illustrates this method.

1. For a fixed sample size n , random deviates from the logistic distribution are generated using a computer subroutine.

2. Estimates of the location and scale parameters are calculated using the method of maximum likelihood. The MLE estimates are calculated iteratively using the Secant method.

3. The random deviates obtained in step one are arranged in ascending order using a computer subroutine.



4. The estimated parameters are used to calculate the hypothesized distribution $F(x)$.

5. The KS statistic is calculated using Eq (4), the A^2 statistic using Eq (6) and the W^2 using Eq (8).

6. The above steps are repeated 5000 times. As a result, 5000 independent KS, A^2 and W^2 statistics are generated.

7. Each group of 5000 statistics are ordered. Using plotting positions, discussed later, the 80th, 85th, 90th, 95th, and 99th percentiles of the distribution of each statistic are calculated by linear interpolation.

8. Steps 1 to 7 are repeated for samples sizes equal to 5, 10, 15, 20, 25 and 30.

Generation of random logistic deviates

There is no known available computer routine to generate random deviates from the logistic distribution. However, the probability integral transform insures that a random variable from any distribution can be transformed to a random variable distributed uniformly on the interval zero to one. Further, this transformation is independent of the location and scale parameters of the distribution (17). Because of this, the concept can be reversed to generate random deviates for the logistic distribution by transforming random deviates distributed uniformly on the interval zero to one. That is, if RN is a pseudo-random deviate distributed uniformly $(0,1)$, it can

be set equal to the logistic cdf

$$RN = [1 + \exp[-\kappa(x - \mu)/\sigma\sqrt{3}]]^{-1} \quad (15)$$

and solving for x yields

$$x = \mu - [\sigma\sqrt{3}(\ln[(1-RN)/RN])]/\kappa \quad (16)$$

However, this still leaves the problem of generating pseudo-random deviates distributed uniformly $(0,1)$. The random uniform deviates for this thesis are obtained on the Control Data Corporation (CDC) 6600 computer using the International Mathematical and Statistics Library (IMSL) subroutine GGUBFS (25:Ch G).

Maximum likelihood estimates of the logistic parameters

A procedure to derive the maximum likelihood estimates (MLE) for location and scale (μ_{E} and σ_{E}) of the logistic distribution was developed by Harter and Moore (23). This procedure iteratively solved the first partial derivatives of the likelihood function after initial estimates have been chosen.

Iterative linear interpolation was applied at each step until the results of successive steps agree to within some assigned tolerance. A procedure similar to this is used in this thesis.

The maximum likelihood method of estimation is one of the most commonly used methods of parameter estimation when using EDF statistics. Some of the reasons for this are that MLE's are consistent and, more importantly, they are invariant (24:239). The method of maximum likelihood, in essence, selects as estimates of the unknown parameters those values for which the observed sample would have most likely occurred (3:83). That is, this method selects as estimates those values that maximize the probability of the occurrence of the sample results (24:236).

The likelihood function is defined in the following way, if $x_1, x_2, \dots, x_n$ are sample observations on X , then the likelihood function (L) is defined to be the joint density function evaluated at $x_1, x_2, \dots, x_n$. The likelihood function can be represented by

$$L = \prod_{i=1}^n f(x_i; \mu, \sigma) \quad (17)$$

The procedure to determine the MLEs μ_{ξ} and σ_{ξ} is:

1. Take the partial derivatives of L with respect to each parameter.

2. Set these equations equal to zero and solve simultaneously for the values of the parameter estimates.

Quite often this procedure is easier to implement if the natural logarithms of L are taken first. Doing this and using Eqs (10) and (17) the likelihood function for the logistic distribution can be written as

$$L = \prod_{i=1}^n \frac{\pi \exp[-\pi(x_i - \mu)/\sigma\sqrt{3}]}{\sigma\sqrt{3}(1 + \exp[-\pi(x_i - \mu)/\sigma\sqrt{3}])^2} \quad (18)$$

and if $z_i = \pi(x_i - \mu)/\sigma\sqrt{3}$ then

$$L = \left(\frac{\pi}{\sigma\sqrt{3}} \right)^n \prod_{i=1}^n \frac{\exp[-z_i]}{(1 + \exp[-z_i])^2} \quad (19)$$

but since there are $n!$ permutations of the realizations on the random variables we have

$$L = n! \left(\frac{\pi}{\sigma\sqrt{3}} \right)^n \prod_{i=1}^n \frac{\exp[-z_i]}{(1 + \exp[-z_i])^2} \quad (20)$$

The natural logarithm of L is then

$$\ln(L) = \ln(n!) + n \ln(\kappa/\sigma\sqrt{3}) - \sum_{i=1}^n z_i + 2 \sum_{i=1}^n \ln(1/\{1+\exp[-z_i]\}) \quad (21)$$

Taking the partial derivatives of $\ln(L)$ with respect to μ and σ yields

$$\frac{\partial \ln(L)}{\partial \mu} = \frac{\kappa}{\sigma\sqrt{3}} - 2 \sum_{i=1}^n \frac{\exp[-z_i]}{1+\exp[-z_i]} \quad (22)$$

$$\frac{\partial \ln(L)}{\partial \sigma} = \frac{1}{\sigma} \sum_{i=1}^n z_i - 2 \sum_{i=1}^n \frac{z_i \exp(-z_i)}{1+\exp[-z_i]} - n \quad (23)$$

These equations, when set equal to zero, can not be solved explicitly. However, the roots of these equations can be found using iterative methods on a computer.

The Secant Method

The Secant method for finding a root of $f(x)=0$ is a slight variation of Newton's method (12:39). Therefore, a brief discussion of Newton's method will precede that of the Secant method.

Suppose that the function f is twice continuously differentiable on the interval $[a,b]$. Let x' be an element of the set of real numbers bracketed by $[a,b]$. Then x' is an approximation of the root p of the equation $f(x)=0$ if $f'(x')$ is not equal to zero and if the absolute difference between x' and p is "small". Consider the second-degree Taylor polynomial for $f(x)$ expanded about x'

$$f(x) = f(x') + (x-x')f'(x') + [(x-x')^2/2]f''(\beta(x)) \quad (24)$$

where $\beta(x)$ lies between x and x'

Since $f(p)=0$ if $x=p$ Eq. (24) becomes

$$0 = f(x') + (p-x')f'(x') + [(p-x')^2/2]f''(\beta(x)) \quad (25)$$

Since $|p-x'|$ is assumed to be small, $(p-x')^2$ is smaller. If $(p-x')^2$ is considered negligible then the third term in Eq (25) can be dropped. Solving for p then yields approximately

$$p = x' - [f(x')/f'(x')] \quad (26)$$

However, with respect to the logistic function, we want to determine the roots of the first derivative instead of the likelihood function itself. That is, Eq (26) becomes

$$p = x' - [L'(x')/L''(x')] \quad (27)$$

and this would require iteratively evaluating both the first and second derivatives of the likelihood function. To circumvent the cumbersome calculations of the second derivative the Secant method can be used. It is a variation of Newton's method and derived as follows:

By definition

$$f'(p_{n-1}) = \lim_{x \rightarrow p_{n-1}} \frac{f(x) - f(p_{n-1})}{x - p_{n-1}} \quad (28)$$

if $x = p_{n-2}$ then

$$f'(p_{n-1}) = \frac{f(p_{n-2}) - f(p_{n-1})}{p_{n-2} - p_{n-1}} \quad (29)$$

and Newton's formula becomes

$$p_n = p_{n-1} - \frac{f(p_{n-1})(p_{n-1} - p_{n-2})}{f(p_{n-1}) - f(p_{n-2})} \quad (30)$$

or in the case of the likelihood function for the logistic distribution

$$p_n = p_{n-1} - \frac{L'(p_{n-1})(p_{n-1} - p_{n-2})}{L'(p_{n-1}) - L'(p_{n-2})} \quad (31)$$

thus eliminating the need for second derivatives.

One requirement of both Newton's and the Secant method is that a good initial approximation be chosen. Otherwise the method may diverge. In this thesis the sample mean and standard deviation will be used as initial approximations.

Ordering the Deviates

The random logistic deviates are arranged in ascending order using the IMSL subroutine VSRTA (25:Ch V).

The Hypothesized Distribution function

The location and standard deviation maximum likelihood estimates (μ_E and σ_E) and the n ordered logistic deviates (x_i) are used to calculate the hypothesized distribution function $F(x)$ by

$$F(x_i) = 1 / (1 + \exp[-\pi(x_i - \mu_E) / \sigma_E \sqrt{3}]) \quad (32)$$

Using the hypothesized and sample distributions the KS, A^2 and W^2 statistics are calculated. This is done 5000 times, once for each of the 5000 samples of size n .

Determining the critical values of the goodness-of-fit tests

The 5000 statistics for each of the three tests, for each sample size, are ordered before determining the critical values. Using plotting positions and linear interpolation the 80th, 85th, 90th, 95th, and 99th percentiles of the distribution of each test statistic are found. These percentiles are the critical values for the KS, A^2 and W^2 goodness-of-fit tests.

Given a series of ordered values, the plotting position of each event is its cumulative probability (22:5). Therefore, the use of plotting positions to determine critical values requires a careful enumeration of the cumulative prob-

abilities on the ordinate axis. This enumeration is complicated by the lack of probabilistic values for the end points of the order statistics. For example, it can not be said that the first order statistic occurs with probability zero. Likewise, the probability of the last order statistic is not one since one more realization of the random variable may yield a higher valued statistic. Because of this, it is necessary to consider the ordinate value of each statistic as a function of its relative position.

Three plotting positions for each point are considered. These are the middle of the interval between the i/n th and $(i-1)/n$ th points, the median ranks plotting position and the mean position. The first plotting position can be expressed as

$$(i-.5)/n \quad (33)$$

This is the value midway through the jump from $(i-1)/n$ to i/n and was first proposed by Hazen in 1914 (22:1). The second plotting position can be very closely approximated by

$$(i-0.3) / (n-0.4) \quad (34)$$

This value is essentially the median value of the distribution of the i th order statistic and was proposed by Johnson in 1951 (22:1). The third plotting position is the mean position and is expressed as

$$i/(n+1) \quad (35)$$

This value is the expected value of the cdf population at the i th order statistic. It was proposed by Weibull in 1939 (22:1).

The KS , A^2 and W^2 ordered statistics form the basis of the abscissa axis. The associated ordered plotting positions form the basis of the ordinate axis. The last entry is added to these ordered values making up the 5001 values on the abscissa and ordinate axes. The 5000 ordered KS (A^2 or W^2) statistics form the 1st to 5000th positions on the abscissa axis. The last position is calculated by linearly extrapolating from the 4999th and 5000th entries. The calculation of the 5001st entry is not subject to a maximum value. The 5000 ordered plotting positions form the 1st to 5000th positions on the ordinate axis. The interval is completed by entering a one in the 5001st position. The addition of these "extra" values on each axis allows an easier determination of the i th and $(i-1)$ th points.

The critical values are calculated by linear interpolation between the ordered statistics and the corresponding

plotting positions. For example, the plotting position just larger and just smaller than .90 are determined. These become the i th and $(i-1)$ th points. The ratio of the difference between the desired percentage point and the $(i-1)$ st point is then calculated. This ratio is then applied to the difference between the i th and $(i-1)$ th test statistics to yield the desired percentile. Figure 3 shows this procedure graphically. At 5000 repetitions there is no difference in the third significant digit in calculating the KS, A^2 or W^2 statistics for the three plotting positions. For simplicity, the plotting position described by Eq (33) is used.

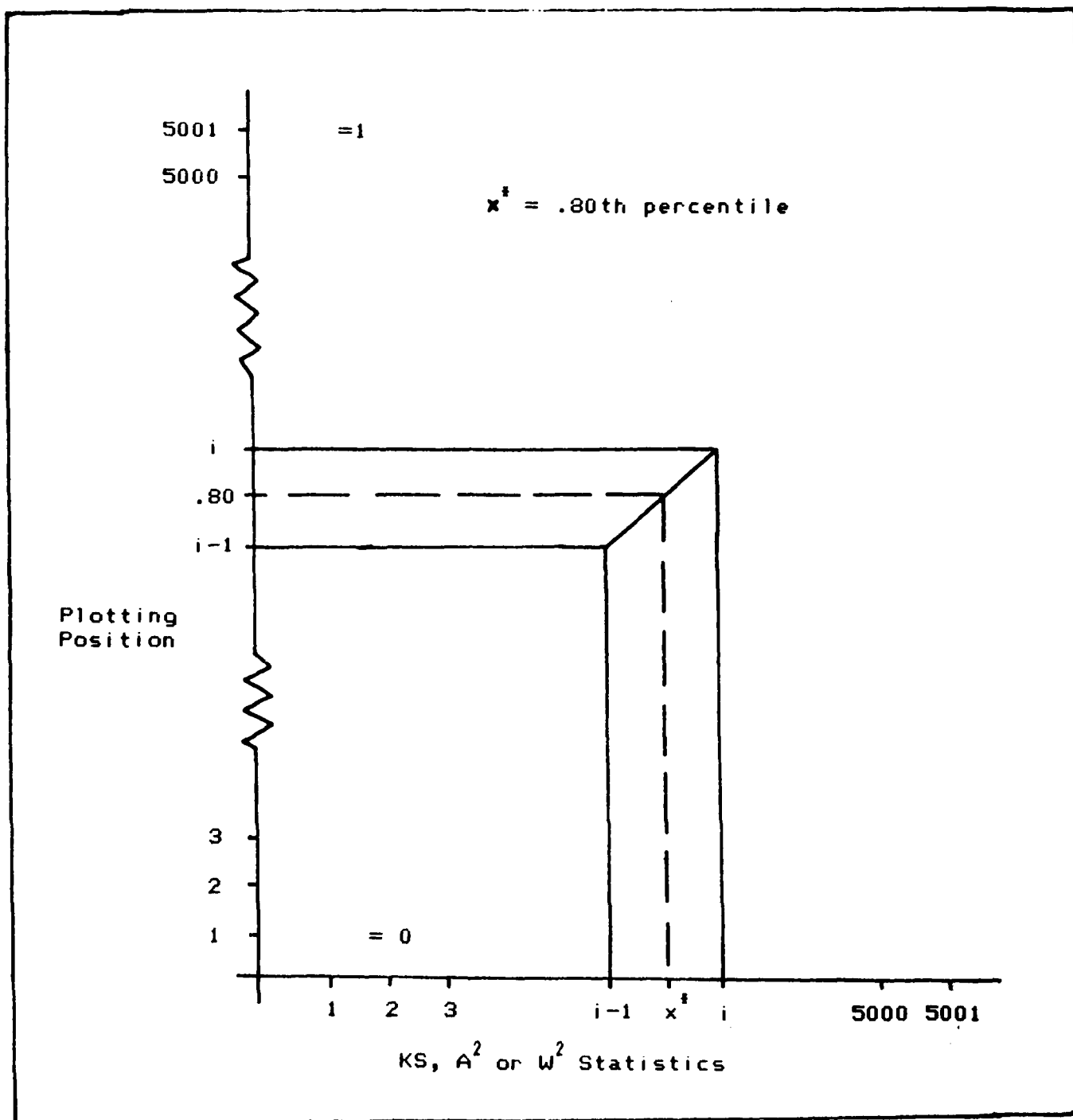


Figure 3
Plotting Position vs. Statistics

Power comparison

In this thesis a comparison of the power of the modified KS, A^2 and W^2 tests is made. The power of each test is compared for several different distributions. The hypothesis being tested is

H_0 : the sample values follow a logistic distribution.

H_A : the sample values do not follow a logistic distribution.

Using IMSL subroutines on the CDC 6600, random deviates from different distributions of sample size n are generated. The test statistics KS, A^2 and W^2 are calculated under the null hypothesis. These test statistics are then compared to the respective critical values developed in this thesis. This procedure is repeated 1000 times for each sample size. The number of times the test statistic exceeds the critical value is counted. Exceeding the critical value results in a rejection of H_0 . The power of the test for a given sample size is the number of rejections divided by the total number of tests, 1000.

The different distributions considered in this power study are:

1. Uniform
2. Exponential
3. Weibull(shape=3)
4. Gamma (shape=3)

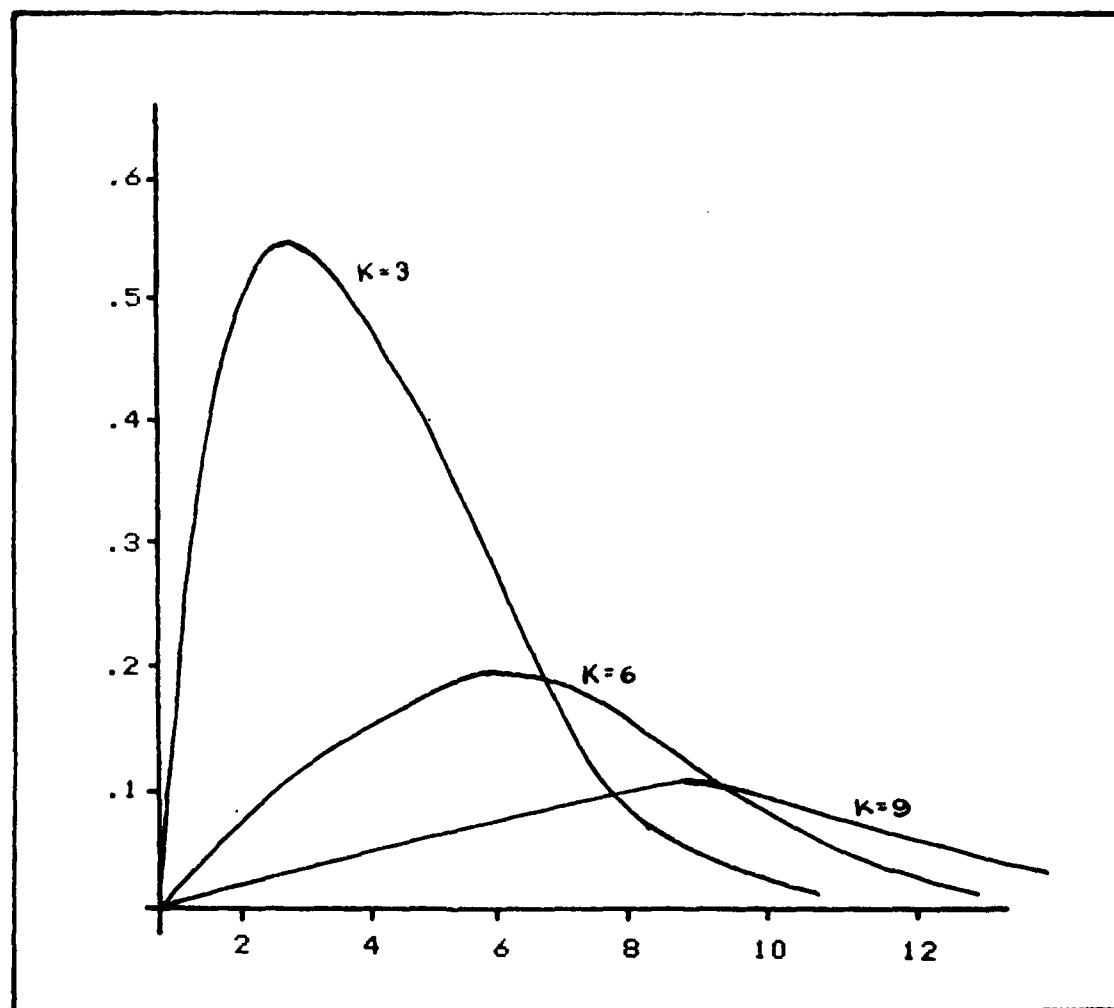
5. Gamma (shape=6)
6. Gamma (shape=9)
7. Gamma (shape=15)
8. Gamma (shape=30)
9. Normal
10. Logistic

The gamma distribution is often used in reliability and maintainability theory. Its density function is expressed as

$$f(x) = \frac{x^{K-1} \exp[-x]}{\Gamma(K)} \quad (36)$$

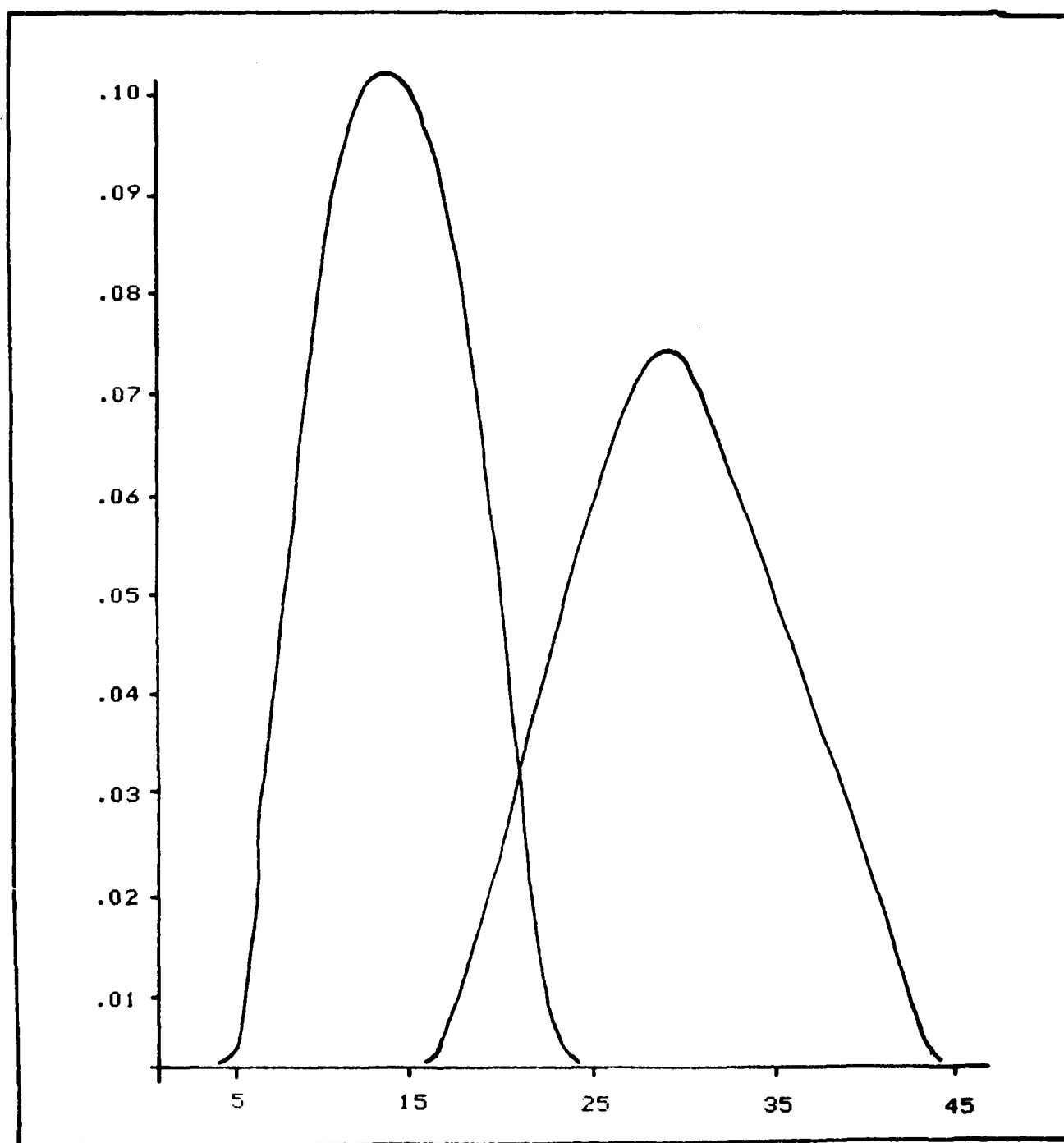
where K = shape parameter

A graph of the gamma distribution for various values of K is shown in Figures 4a and 4b.



Gamma distribution: $K=3,6,9$

Figure 4a



Gamma distribution: $K=15,30$

Figure 4b

The Weibull distribution is also often used in conjunction with reliability and maintainability studies. Its density function is expressed as

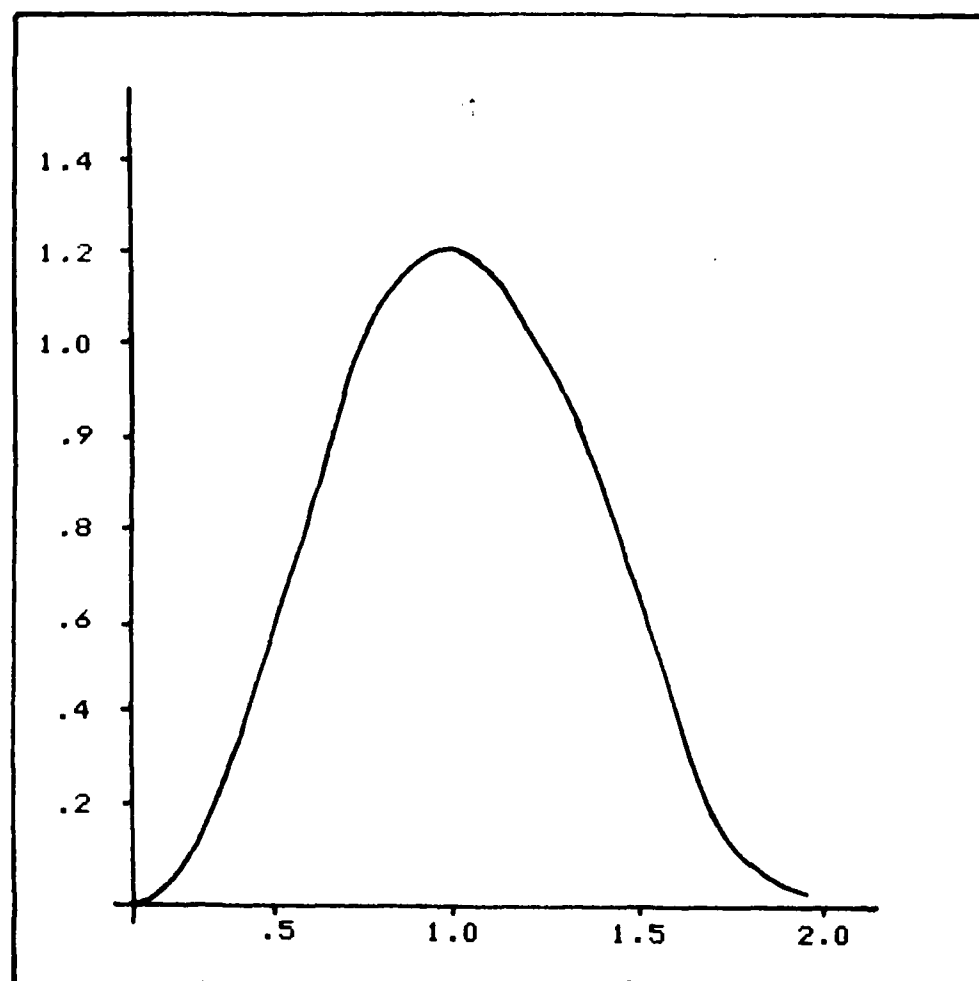
$$f(x) = Kx^{K-1} \exp[-x^K] \quad (37)$$

where K holds the same notation as in Eq (27). A graph of the Weibull distribution is shown in Figure 5.

As can be seen, the gamma and weibull distributions are generally symmetrically convex for the correct parameter values. The normal and exponential distributions are well known and need not be discussed here.

Computer programs

Computer programs used in this thesis are presented in Appendix E.



Weibull distribution: $K=3$

Figure 5

IV. USE OF TABLES

This chapter discusses the use of the tabled critical values of the two parameter Logistic distribution for the modified KS, A^2 and W^2 statistics generated in this thesis. An example follows, with an explanation of the basic procedure to utilize the table.

For all three goodness-of-fit tests, a theoretical distribution $F(x)$ is compared to an empirical, observed distribution $S(x)$. The KS, A^2 , or W^2 statistic is calculated using the appropriate equation (Eq (4), Eq (6) or Eq (8)). If the value of the statistic exceeds the tabled critical value, the theoretical distribution is rejected. The steps in applying this procedure are:

1. Determine the sample size n and the level of significance. The significance level is the probability of rejecting the null hypothesis that the sample is from the Logistic distribution when the null hypothesis is true.
2. Select, in a random manner, the n observations from the population to be tested and order them from smallest to largest.
3. Estimate the unknown location and scale parameters from the observed sample using the method of maximum likelihood.
4. Completely specify the hypothesized distribution $F(x)$

using the estimated location and scale parameters. Determine the values of the empirical distribution just prior to and after each of the $1/n$ jumps.

5. Determine the value of the KS, A^2 , or W^2 test statistics by using Eq (4), Eq (6) and Eq (8).

6. Find the intersection of the significance column and the sample size row in the table of values. This is the critical value to be compared to the test statistic.

7. Reject the null hypothesis if the value of the test statistic exceeds the critical value. If the test statistic does not exceed the critical value, we fail to reject the null hypothesis and conclude there is insufficient evidence to say the observed sample does not follow the logistic distribution.

EXAMPLE

The following example illustrates this procedure for the modified Kolmogorov-Smirnov test. For a sample size of five the following numbers are obtained: 104.9829, 81.1517, 87.2204, 113.5512, and 61.5415. For this sample, The maximum likelihood estimate subroutine used in this thesis yields location and standard deviation parameter estimates of: $\mu_E = 90.0986$ and $\sigma_E = 20.1107$. The scale parameter is calculated using $\sigma_E \sqrt{3/\pi}$, yielding a value of 11.0876. Using these values the hypothesized distribution is calculated for each sample value using Eq (9). The significance level is .05. The hypothesis tested is:

H_0 : sample data is from a Logistic distribution

H_A : sample data is not from a Logistic distribution

The calculations for this test are shown in Table I.

Table I Example Calculations				
x	F(x)	$S(x^-)$	$S(x^+)$	KS
64.5415	.071	0	.2	.129
81.1517	.309	.2	.4	.109
87.2204	.436	.4	.6	.164
104.9829	.793	.6	.8	.193
113.5512	.892	.8	1.0	.108

where $S(x^-)$ is just prior to the jump in the empirical distribution

$S(x^+)$ is just after the jump in the empirical distribution

Using Eq (4), the KS statistic equals .193. The critical value from Table V with a significance level of .05 and $n = 5$

is .309. Since .193 is less than .309, we fail to reject the null hypothesis that the sample comes from a Logistic distribution.

V. DISCUSSION OF RESULTS

This chapter discusses the results obtained in this thesis as they pertain to the objectives set forth in Chapter I.

PRESENTATION OF THE TABLE OF CRITICAL VALUES

The tabled critical values (for sample sizes of 5, 10, 15, 20, 25, and 30) for the modified KS, A^2 , and W^2 goodness-of-fit tests are presented in Appendancies A, B and C respectively. Each table is valid for the Logistic distribution when the location and scale parameters are estimated from the sample using the method of maximum likelihood presented in this thesis.

However, the Secant method used in this thesis has certain properties that require explanation. As stated earlier, the Secant method converges only when the initial estimates of the parameters are "good". In this thesis the sample mean and standard deviation are used as initial estimates of the unknown parameters. In all but four of the 30,000 samples used in this thesis, the sample mean and standard deviation allowed convergence using the Secant method. The four cases of non-convergence occurred when the sample was very tightly grouped, thus yielding a very small sample standard deviation. This occurred two times in a sample size of five and once each

in the sample sizes of 10 and 15. Discarding these non-convergent samples and obtaining a new sample theoretically biases the resulting calculations. However, it was felt that such a small number of samples could be discarded without meaningfully biasing the numerical results of this thesis.

This conclusion is reached for the following reason. Discarding samples, within a single sample size, changes the relative positions of some of the KS, A^2 and W^2 statistics calculated for that sample size. That is, if five samples are taken and five KS statistics obtained, they will have a specific ascending order. However, if the third sample is discarded and a "new" sample and KS statistic obtained, the order of some of the five KS statistics might be changed. The "new" KS statistic may not fit in the third ordered position.

The most severe case of discarding samples in this thesis is two out of 5,000. This occurred with a sample size of five. This effectively means that any calculated statistic would be, at most, two "positions" out of order. The statistics calculated in this thesis, when ordered, generally increment in the fourth digit. Additionally, a statistical value two positions out of order would cause an error of $\pm 0.04\%$. That is, for 5000 repetitions, if a statistical value is two positions out of order, then a critical value based on this ordering is at most $\pm 0.04\%$ in error. For sample sizes of 10 and 15, the critical values reported would be in error by at most $\pm 0.02\%$. It seems unlikely that any change in the three significant digits reported occurred due to the

policy chosen.

However, this property of non-convergence in some cases represents a constraint on using the Secant method when finding the maximum likelihood estimates of the unknown parameters. It is not unreasonable to expect some cases of live data to truly have a small standard deviation. The program used for this thesis to estimate the unknown parameters of the Logistic distribution may therefore be inappropriate in some cases. A gradient search method, such as that presented by Wingo (44:91), may overcome this problem.

ANALYSIS OF CRITICAL VALUE TABLES

Analysis of the tabled critical values generated by this thesis and the distributions of the KS, A^2 and W^2 statistics reveal the following results.

In the case of the KS tabled critical values, analysis relative to the significance levels and sample size reveal that a concave pattern is relatively constant. The change in significance level simply shifts this curve upward (Fig 6).

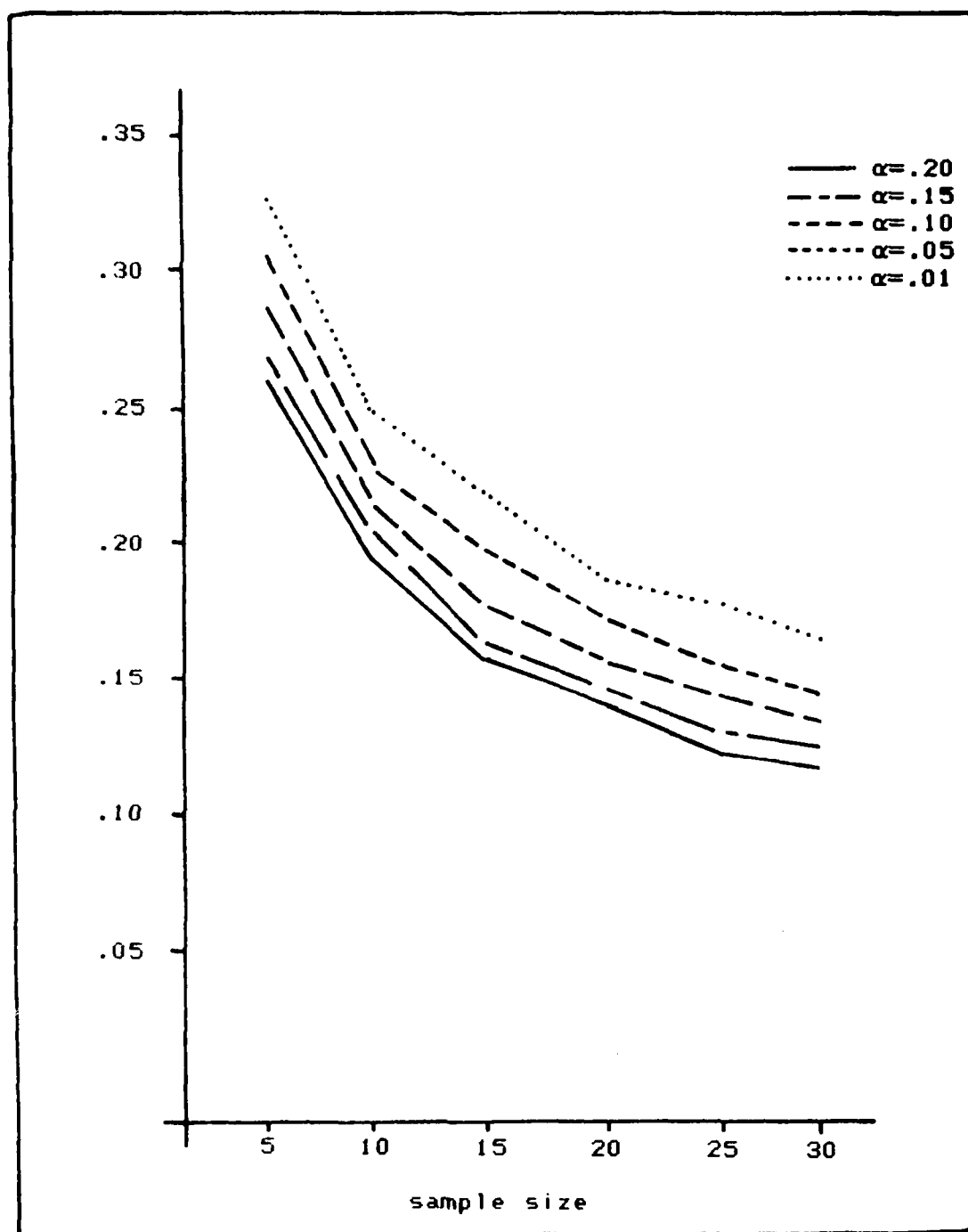


Figure 6: KS critical values by sample size

For the A^2 tabled critical values, a similar analysis shows that the first three significance levels have a smooth, slightly increasing, monotonic trend. The last two significance levels show an increasing trend, but the transition between sample sizes is not as smooth or as constant (Fig 7).

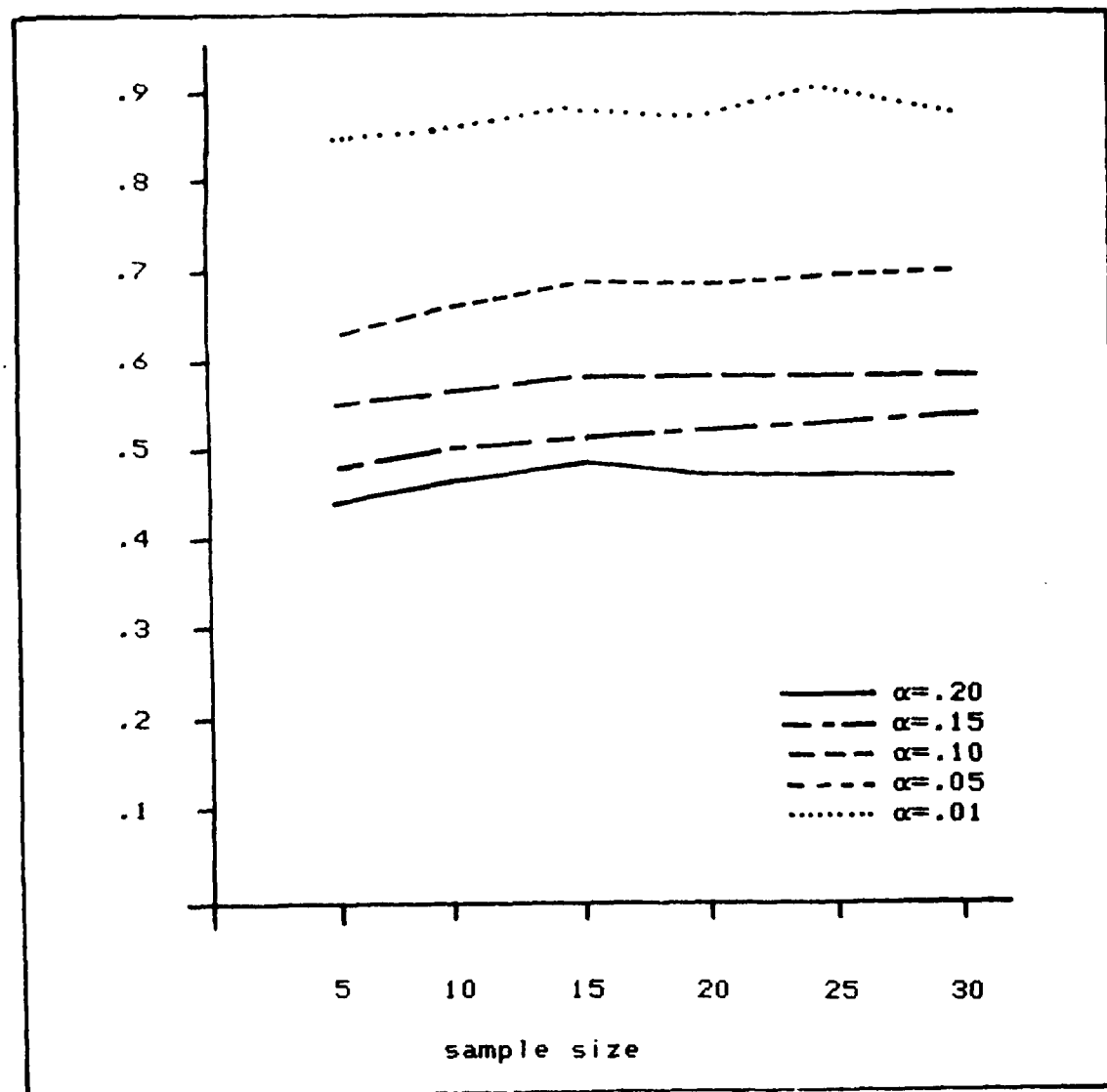


Figure 7: A^2 critical values by samples size

Results of the analysis of the W^2 tabled critical values are very similar to those of the A^2 statistic. However, instead of a smooth, slightly increasing, trend for the first three significance levels there appears to be a generally constant horizontal trend (Fig 8). At the .05 significance level the change between sample sizes becomes more erratic and possibly increasing. At the .01 significance level the change between samples is generally increasing with an upward jump between sample sizes of 10 and 15.

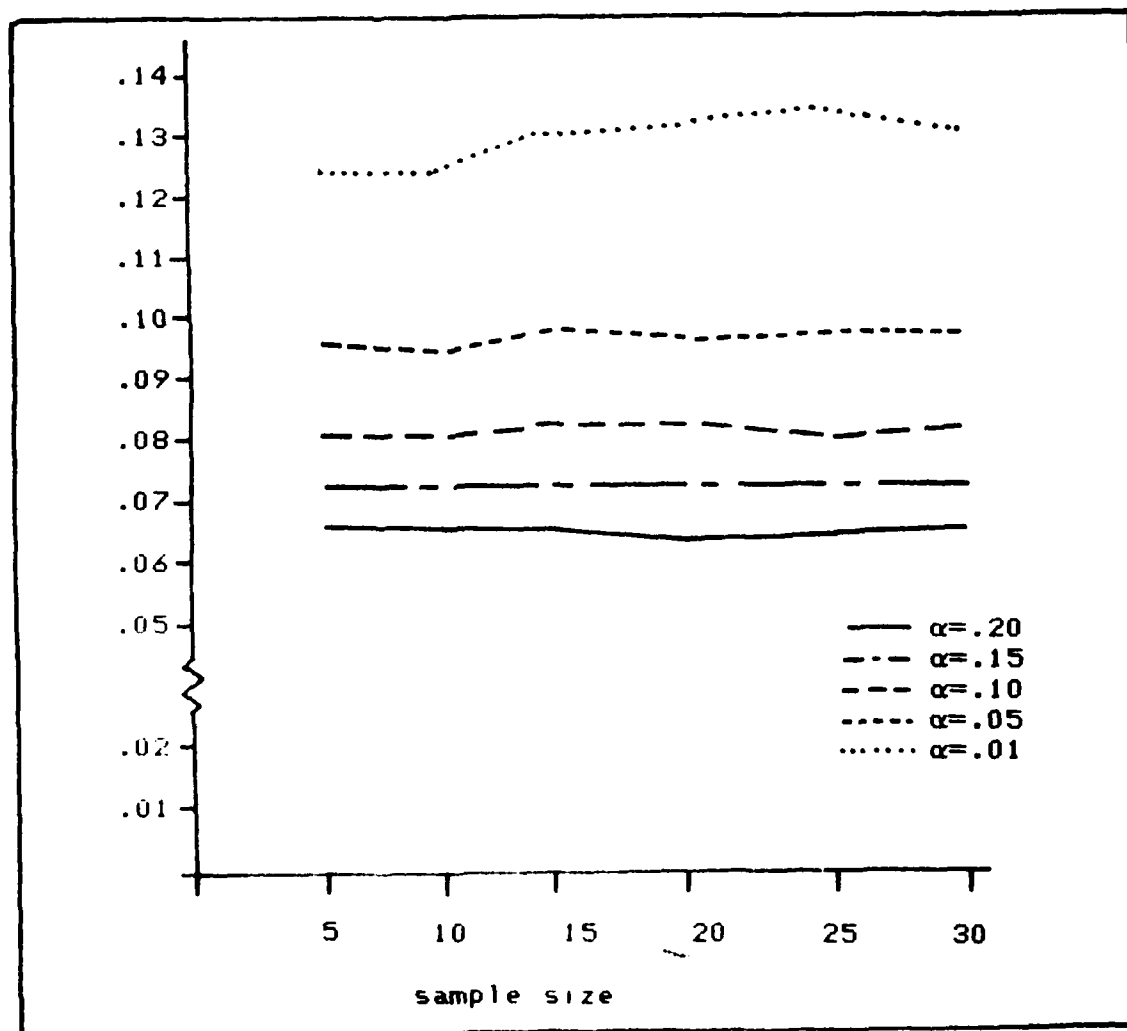
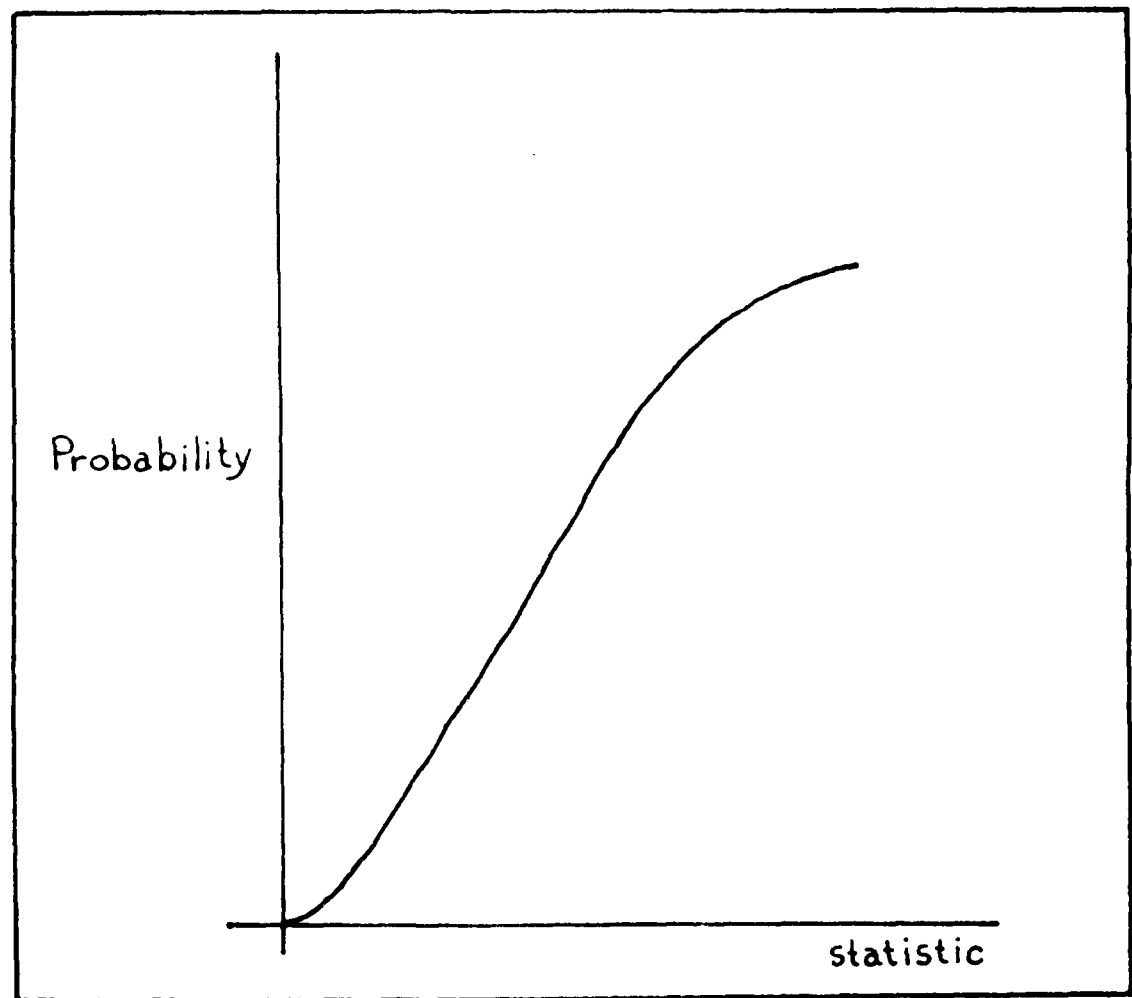


Figure 8: W^2 critical values by sample size

Analysis of the distributions of the KS , A^2 and W^2 statistics developed in this thesis reveal the following result. For all sample sizes, the distribution of all three statistics approximately follows a Gompertz curve (Fig 9). For each statistic, this curve is slightly modified by sample size.



general distribution shape of test statistics

Figure 9

VALIDATION OF COMPUTER PROGRAM

The computer program used to generate the critical values obtained in this thesis is validated by comparing the critical values obtained in this thesis to those obtained by M.A. Stephens and by hand calculations. The hand calculations validate that the computer program generates numbers correctly, while the comparison to Stephens values validates that the values of this thesis are acceptable. In 1979 Stephens calculated modified KS, A^2 and W^2 statistics for the Logistic distribution based on asymptotic distribution theory. This thesis calculated KS, A^2 and W^2 statistics based on finite distribution theory. If the two sets of statistics are comparable this implies the computer program generates valid results.

This comparison was done in the following manner. The critical values obtained in this thesis are modified according to the appropriate formula presented in Stephens paper (14:14,16). The result is then compared to the tabled values presented by Stephens. For example, with a sample size of 20 and an alpha level of .05, the A^2 statistic obtained in this thesis is .655. This value was then modified by equation 38.

$$A^2(1.0+(0.25/n)) \quad (38)$$

The result (.663) is then compared to 0.660, the value presented

by Stephens. Such comparisons for all values obtained in this thesis are generally good. Tables II, III and IV show these comparisons.

Table II Kolmogorov-Smirnov (KS) Comparison						
1- α	KS $\sqrt{n}$			Stephens		
	5	10	20	5	10	20
.90	.633	.677	.698	.643	.679	.698
.95	.691	.727	.760	.679	.730	.755
.99	.754	.813	.836	.751	.823	.854

Table III Anderson-Darling (A^2) Comparison				
$1-\alpha$	$\frac{A^2(1+.25/n)}{n}$			Stephens Critical values
	5	10	20	
.90	.570	.562	.564	.563
.95	.651	.656	.663	.660
.99	.891	.874	.887	.906

Table IV Cramer-Von Mises (W^2) Comparison				
$1-\alpha$	$\frac{(nW^2-0.08)/(n-1)}{n}$			Stephens Critical values
	5	10	20	
.90	.081	.081	.081	.081
.95	.099	.097	.096	.098
.99	.135	.129	.133	.136

An assumption requiring validation is with respect to the MLE estimates generated. Maximum likelihood theory specifies that MLE estimates are invariant. The invariant property is critical with respect to this thesis since the validity of the tabled critical values depends on it. This assumption is examined with the MLE estimates of a logistic(1,4) and a logistic(4,16). In order to be invariant the MLE estimates of the logistic(1,4) should be equal to the MLE estimates of the logistic(4,16) after multiplication by a factor of four. These estimates were obtained on separate computer runs for sample sizes of 5, 15 and 25. Hand calculations show that when multiplied by a factor of four the logistic(1,4) MLEs differed from those of the logistic(4,16) by generally 0.000001. This small difference is attributed to computer round-off.

SENSITIVITY ANALYSIS OF PROGRAM

A condition required by the Secant method is that for convergence the initial estimates must be "good". Since in a few cases the Secant method used in this thesis did not converge, an examination of what "good" means is conducted. This is done by systematically changing the population standard deviation used to generate the logistic deviates, and generating 2000 samples for each sample size. Each time the Secant method used in this thesis does not converge, the sample mean and standard deviation

are printed out. By comparing results it becomes apparent that when the sample is tightly grouped about the population mean, the sample standard deviation becomes very small. When the sample standard deviation is generally less than 25-30% of the population standard deviation, the method fails to converge.

PRESENTATION OF THE POWER STUDY

A comparison of the relative power of the modified Kolmogorov-Smirnov, Anderson-Darling and Cramer-Von Mises goodness-of-fit tests developed in this thesis is made. For distributions with non-symmetrical convex patterns the power of the KS, A^2 and W^2 tests is very good. For distributions with symmetric convex patterns, the power is much lower. The distributions used in this power study are listed in Chapter III. For each of these distributions 1000 samples, for each sample size, are obtained. For each sample a KS, A^2 , and W^2 statistic is calculated using Eq (4), Eq (6) or Eq (8). Each statistic is compared to the appropriate tabled critical value in table V, VI or VII. This comparison is done at the 0.05 significance level. For each sample size, the number of times the calculated statistic exceeds the tabled critical value, an index variable is incremented by one. This index variable is then divided by 1000 to yield the percent of time the null hypothesis (that the sample came from a logistic distribution) is rejected.

When the sample comes from a uniform or an exponential

distribution the power of the modified KS, A^2 and W^2 tests is very good (Tables IIX through XIII). The shape of the pdf of these distributions is decidedly not symmetrically convex. This undoubtedly is the reason the tests presented in this thesis are so powerful when testing samples from these distributions.

However, when the sample comes from a distribution with a symmetric convex pattern the power of the tests presented in this thesis is much lower. This is due primarily to the difficulty of distinguishing between similar, yet different, symmetrical convex patterns. The distributions tested in this category are the Normal, Weibull, and various Gamma distributions. In the case of the Weibull and Normal distributions, as well as the Gamma with higher shape parameter values, the power of the modified KS, A^2 and W^2 tests is little more than the significance level. Only when the symmetric convex pattern of the Gamma distribution becomes skewed does the power of these tests begin to noticeably exceed the significance level.

The modified tests presented in this thesis generally reject the null hypothesis of a logistic distribution when the sample comes from a Gamma distribution. Yet, they have only moderate power for lower shape parameter values, when the symmetric convex pattern is skewed. For higher shape parameter values their rejection power is only slightly above the significance level. This shows how difficult it is for these tests to distinguish between other symmetric convex patterns and the symmetric convex pattern of the logistic distribution.

The Normal distribution is symmetrically convex and very close in shape to that of the Logistic distribution. All three modified tests have considerably more difficulty rejecting the null hypothesis in this case. Each test, generally, reports a different power for each sample size. The power is sometimes slightly above the significance level and at other times slightly below it. There does not appear to be any consistent pattern in rejecting the null hypothesis.

When applied to the Logistic distribution, all tests generally failed to reject the null hypothesis that the sample came from a logistic distribution. Yet, in this case, the rejection percentage is not always equal to the significance level as might be expected. This variability is primarily a function of the Monte Carlo method and the 1000 repetitions upon which the study is based. For example, the expected number of times the power of these tests should be between .04 and .06 based on 1000 repetitions is approximately three. In actuality, this occurred four times. In most cases, the rejection percentage, when the sample comes from the Logistic distribution, is between .048 and .052.

In summary, all three tests are better at rejecting the null hypothesis of a logistic distribution when the sample distribution is not from a Normal distribution, than when it is from a Normal distribution. When the sample distribution is not a Normal distribution, the W^2 test is slightly better than the A^2 or KS tests. When the sample is from a Normal distribution the W^2 test is generally better than

the KS test and A^2 tests. Yet, rejecting the null hypothesis of a logistic distribution when the sample has a symmetric convex pattern is generally not possible. All percentages used in these comparisons are presented in Appendix D.

VI. CONCLUSIONS AND RECOMMENDATIONS

Conclusions

Based on results obtained in this thesis, the following conclusions are noted:

1. The Kolmogorov-Smirnov, Anderson-Darling and Cramer-von Mises critical values for the two parameter logistic distribution are valid.

2. The power comparison study based on the ten different distributions listed in Chapter III shows that for non-symmetrical convex distributions all three tests exceed the claimed level of significance. For skewed symmetrical convex distributions all three tests approximate or exceed the level of significance claimed. For symmetric convex distributions all three tests very closely approximate the claimed level of significance. This indicates a goodness-of-fit test for a sample from a symmetrical convex distribution is not practical. The W^2 test is generally more powerful than the A^2 or KS tests.

Recommendations

The following recommendations are suggested for further investigation:

1. Determine if another method for determining initial

estimates significantly affects the critical values and power presented in this thesis.

2. Reduce the variability in the significance level when the null hypothesis of a logistic distribution is true by employing a method other than the Monte Carlo method, or increase the number of repetitions used in the power study.

3. Employ some other method to determine the critical values for the logistic distribution.

4. Using Stephens values and formulae perform a power study and compare those results to the results of this thesis to determine which method is more powerful.

5. Develop a test to distinguish between symmetrical convex distributions.

BIBLIOGRAPHY

1. Anderson T. and D. Darling, "Asymptotic theory of certain 'goodness-of-fit' criteria based on stochastic processes", Annals of Mathematical Statistics, 23: 193-212, 1952.
2. _____, "A Test of Goodness of Fit", Journal of American Statistical Association, 49, 765-769, 1954.
3. Armstradter, B., Reliability Mathematics, New York: McGraw-Hill Book Company, 1971.
4. Bain, Lee J., Statistical Analysis of Reliability and Life-Testing Models, New York, Marcell Dekker Inc., 1978.
5. Berkson, Joseph, "Application of the Logistic Function to Bioassay", Journal of the American Statistical Association, 39: 357-365, 1944.
6. _____, "Why I prefer Logits to Probits", Biometrics, 7: 327-339, 1951.
7. _____, "Tables for the Maximum Likelihood Estimate of the Logistic Function", Biometrics, 13: 28-34, 1957.
8. Berkson, J. and J. Hodges Jr., "A Minimax Estimator of the Logistic Function", Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability, 4, 77-86, University of California Press, Berkeley, 1961.
9. Beyer, J. N., Linear Estimation of the Scale Parameter of the Logistic Distribution by the Use of Order Statistics, Unpublished MS Thesis, Air Force Institute of Technology, Wright-Patterson Air Force Base, Ohio, 1966.
10. Birnbaum, A., "On Logistic Order Statistics", Annals of Mathematic Statistics, 29: 1285, 1958, Abstract.
11. Birnbaum, A. and J. Dudman, "Logistic Order Statistics",

Annals of Mathematic Statistics, 34: 658-663, 1963.

12. Burden, Faires and Reynolds, Numerical Analysis (second ed), Boston: Prindle, Weber & Schmidt, 1981.
13. Bush, J., "A Modified Cramer-von Mises and Anderson-Darling Test for the Weibull Distribution with Unknown Location and Scale Parameters", Unpublished MS Thesis, Air Force Institute of Technology, Wright-Patterson AFB, Ohio, 1981.
14. Caplen, R., A Practical Approach to Reliability, London: Business Books, Ltd., 1972.
15. Conover, W., Practical Nonparametric Statistics, 2ed, New York: John Wiley & Sons, Inc., 1980.
16. Cortes, R., "A Modified Kolmogorov-Smirnov Test for the Gamma and Weibull Distributions with Unknown Location and Scale Parameters", Unpublished MS Thesis, Air Force Institute of Technology, Wright-Patterson AFB, Ohio, 1980.
17. David, F. and N. Johnson, "The Probability Integral Transformation when Parameters are Estimated from the Sample", Biometrika, 35: 182-190, 1948.
18. Green, J. and Y. Hegazy, "Powerful Modified-EDF goodness-of-fit Tests", Journal of the American Statistical Association, 71: 204-209, 1965.
19. Gibbons, J., Nonparametric Statistical Inference, New York: McGraw-Hill Book Company, 1971.
20. Gumbel, E., "Ranges and Midranges", Annals of Mathematic Statistics, 15: 420, 1958.
21. Gupta, S., A. Qureishi, and B. Shah, Best Linear Unbiased Estimators of the Parameters of the Logistic Distribution Using Order Statistics, Dept. of Statistics, Purdue University, Mimeograph Series Number 52, 1965, AD 621150.
22. Gupta, S. and B. Shah, "Exact Moments and Percentage Points of the Order Statistics and the Distribution of the Range

- from the Logistic Distribution", Annals of Mathematical Statistics, 36: 907-920, 1965.
23. Harter, H., "Another look at Plotting Positions", Unpublished Paper, Air Force Institute of Technology, Wright-Patterson AFB, Ohio and Wright State University, Fairborn, Ohio, 1983.
 24. Harter, H. and Moore, Albert H., "Maximum Likelihood Estimation, From Censored Samples, of the Parameters of a Logistic Distribution", Journal of the American Statistical Association, 318: 675-684, 1967.
 25. Hines, W. and D. Montgomery, Probability and Statistics in Engineering and Management Science (second ed), New York: John Wiley & Sons, 1980.
 26. International Mathematical and Statistic Library Reference Manual - 0008. Houston: IMSL, 1980.
 27. Kjelsberg, M., Estimation of the Logistic Distribution Under Truncation and Censoring, Unpublished MS Thesis, University of Minnesota, 1962.
 28. Koutrouvelis, I. and J. Kellermeier, "A Goodness-of-fit Test Based on the Empirical Characteristic Function Where Parameters must be Estimated", Journal of the Royal Statistical Society, 43: 173-176, 1981.
 29. Leach, D., "Re-evaluation of the logistic curve for human populations", Journal fo Royal Statistical Society, 144: 94-103, 1981.
 30. Lilliefors, H., "On the Kolmogorov-Smirnov Test for Normality with Mean and Variance Unknown", Journal of the American Statistical Association, 62: 399-402, 1967.
 31. _____, "On the Kolmogorov-Smirnov Test for the Exponential Distribution with Mean Unknown", Journal of the American Statistical Association, 64: 387-399, 1979.
 32. Masaaki, T., O. Hiroshi, and K. Shigeo, "Goodness-of-fit test for extreme value distribution", IEEE Transactions of Reliability, R-29: 151-153, 1980.

33. Pearl, R. and L.J. Reed, "On the Rate of Growth of the Population of the United States Since 1790 and Its Mathematical Representation", Proceedings National Academy of Sciences, 6: 275-288, 1920.
34. Plackett, R. L., "The Analysis of Life Test Data", Technometrics, 1: 9-19, 1959.
35. _____, "Linear Estimation from Censored Data", Annals of Mathematical Statistics, 29: 131-142, 1958.
36. Reed, L. and J. Berkson, "The Application of the Logistic Function to Experimental Data", Journal of Physical Chemistry, 33: 760-779, 1929.
37. Richardson, E., Simultaneous Linear Estimation of the Location and Scale parameters of the Extreme-Value and Logistic Distributions by the Use of Selected Order Statistics, Unpublished MS Thesis, Air Force Institute of Technology, Wright-Patterson Air Force Base, Ohio, 1966.
38. Schultz, H., "The Standard Error of a Forecast From a Curve", Journal of the American Statistical Association, 25: 139-185, 1930.
39. Shah, B., On the Bivariate Moments of Order Statistics from a Logistic Distribution and Applications, Dept. of Statistics, Purdue University, Mimeograph Series Number 48, 1965, AD 619501.
40. Stephens, M., "The Anderson-Darling Statistic", Unpublished Technical Report No. 39, for the U.S. Army Research Office, Stanford University, Stanford, California, 1979.
41. _____, "EDF Tests of Fit for the Logistic Distribution", Unpublished Technical Report No. 36, for the U.S. Army Research Office, Stanford University, Stanford, California, 1979.
42. Viviano, P., A Modified Kolmogorov-Smirnov, Anderson-Darling,

Cramer-von Mises Test for the Gamma Distribution with Unknown Location and Scale parameters, Unpublished MS Thesis, Air Force Institute of Technology, Wright-Patterson AFB, Ohio, 1982.

43. Wilson, E. and J. Worcester, "The Determination of L. D. 50 and Its Sampling Error in Bio-assay", Proceedings of the National Academy of Sciences, 29, 79ff, 1943.
44. Wingo, D.R., "Maximum Likelihood Estimation of the Parameters of the Weibull distribution by Modified Quasilinearization", IEEE Transactions on Reliability, R-21, 89-93, May 1972.

Appendix A

TABLE V

**CRITICAL VALUES OF THE KOLMOGOROV-SMIRNOV
STATISTIC FOR THE LOGISTIC DISTRIBUTION
WITH UNKNOWN LOCATION AND SCALE PARAMETERS**

Sample size n	Level of Significance				
	.20	.15	.10	.05	.01
5	.262	.272	.283	.309	.337
10	.195	.203	.214	.230	.257
15	.163	.169	.178	.191	.220
20	.143	.148	.156	.170	.187
25	.128	.136	.141	.153	.175
30	.118	.124	.131	.141	.160

Appendix B

TABLE VI

**CRITICAL VALUES OF THE ANDERSON-DARLING
STATISTIC FOR THE LOGISTIC DISTRIBUTION
WITH UNKNOWN LOCATION AND SCALE PARAMETERS**

Sample size n	Level of Significance				
	.20	.15	.10	.05	.01
5	.443	.484	.543	.620	.849
10	.456	.494	.549	.640	.853
15	.461	.498	.554	.665	.882
20	.456	.499	.557	.658	.876
25	.453	.493	.555	.663	.918
30	.456	.502	.552	.649	.855

Appendix C

TABLE VII

CRITICAL VALUES OF THE CRAMER-VON MISES
STATISTIC FOR THE LOGISTIC DISTRIBUTION
WITH UNKNOWN LOCATION AND SCALE PARAMETERS

Sample size n	Level of Significance				
	.20	.15	.10	.05	.01
5	.066	.072	.081	.096	.124
10	.065	.072	.081	.095	.124
15	.065	.072	.082	.098	.130
20	.064	.072	.082	.097	.131
25	.065	.072	.081	.098	.135
30	.065	.072	.082	.098	.132

Appendix D

TABLE IIX			
COMPARISON OF KOLMOGOROV-SMIRNOV, ANDERSON-DARLING AND CRAMER-VON MISES STATISTICS FOR SAMPLE SIZE OF 5, AND SIGNIFICANCE LEVEL OF .05			
PERCENT OF TIME NULL HYPOTHESIS REJECTED			
	KS	A^2	W^2
Uniform	.208	.119	.168
Exponential	.257	.119	.178
Weibull(3)	.023	.027	.039
Gamma(3)	.069	.109	.069
Gamma(6)	.050	.079	.059
Gamma(9)	.040	.069	.069
Gamma(15)	.099	.079	.089
Gamma(30)	.059	.069	.059
Normal	.059	.050	.059
Logistic	.051	.049	.049

TABLE IX			
COMPARISON OF KOLMOGOROV-SMIRNOV, ANDERSON-DARLING AND CRAMER-VON MISES STATISTICS FOR SAMPLE SIZE OF 10, AND SIGNIFICANCE LEVEL OF .05			
PERCENT OF TIME NULL HYPOTHESIS REJECTED			
	KS	A^2	w^2
Uniform	.287	.198	.267
Exponential	.535	.495	.564
Weibull(3)	.046	.049	.057
Gamma(3)	.089	.129	.109
Gamma(6)	.050	.059	.040
Gamma(9)	.079	.069	.069
Gamma(15)	.050	.050	.079
Gamma(30)	.047	.050	.053
Normal	.040	.059	.089
Logistic	.049	.052	.061

TABLE X			
COMPARISON OF KOLMOGOROV-SMIRNOV, ANDERSON-DARLING AND CRAMER-VON MISES STATISTICS FOR SAMPLE SIZE OF 15, AND SIGNIFICANCE LEVEL OF .05			
PERCENT OF TIME NULL HYPOTHESIS REJECTED			
	KS	A^2	W^2
Unifor	.317	.277	.337
Exponent	.633	.703	.782
Weibull(3)	.049	.051	.059
Gamma(3)	.129	.218	.158
Gamma(6)	.119	.129	.139
Gamma(9)	.079	.089	.089
Gamma(15)	.079	.069	.069
Gamma(30)	.049	.040	.049
Normal	.050	.020	.030
Logistic	.050	.052	.061

TABLE XI

COMPARISON OF KOLMOGOROV-SMIRNOV, ANDERSON-DARLING
AND CRAMER-VON MISES STATISTICS
FOR SAMPLE SIZE OF 20, AND SIGNIFICANCE LEVEL OF .05

PERCENT OF TIME NULL HYPOTHESIS REJECTED

	KS	A^2	W^2
Uniform	.465	.528	.528
Exponential	.812	.851	.881
Weibull(3)	.058	.059	.059
Gamma(3)	.099	.178	.149
Gamma(6)	.099	.129	.149
Gamma(9)	.069	.129	.109
Gamma(15)	.059	.089	.089
Gamma(30)	.069	.069	.069
Normal	.118	.069	.109
Logistic	.041	.048	.050

TABLE XII

COMPARISON OF KOLMOGOROV-SMIRNOV, ANDERSON-DARLING
AND CRAMER-VON MISES STATISTICS
FOR SAMPLE SIZE OF 25, AND SIGNIFICANCE LEVEL OF .05

PERCENT OF TIME NULL HYPOTHESIS REJECTED

	KS	A^2	W^2
Uniform	.545	.584	.634
Exponential	.832	.891	.891
Weibull(3)	.058	.049	.078
Gamma(3)	.099	.208	.139
Gamma(6)	.129	.139	.129
Gamma(9)	.109	.129	.129
Gamma(15)	.089	.139	.109
Gamma(30)	.089	.069	.089
Normal	.040	.020	.040
Logistic	.050	.039	.039

TABLE XIII

COMPARISON OF KOLMOGOROV-SMIRNOV, ANDERSON-DARLING
AND CRAMER-VON MISES STATISTICS
FOR SAMPLE SIZE OF 30, AND SIGNIFICANCE LEVEL OF .05

PERCENT OF TIME NULL HYPOTHESIS REJECTED			
	KS	A^2	W^2
Uniform	.604	.673	.624
Exponential	.931	.980	.980
Weibull(3)	.068	.066	.079
Gamma(3)	.218	.386	.317
Gamma(6)	.109	.248	.198
Gamma(9)	.099	.119	.109
Gamma(15)	.069	.109	.089
Gamma(30)	.119	.079	.107
Normal	.040	.050	.059
Logistic	.050	.050	.051

Appendix E

PROGRAM VALUES

```

*****
*
*   DEFINITION OF VARIABLES IN MAIN PROGRAM
*
*   PURPOSE: GENERATE A SAMPLE FROM THE LOGISTIC DISTRIBUTION,
*             OBTAIN THE MLE ESTIMATES OF THE LOCATION AND SCALE,
*             DETERMINE THE HYPOTHETICAL AND EMPIRICAL DISTRIBUTION
*             OF THE SAMPLE, AND CALCULATE THE KOLMOGOROV-SMIRNOV,
*             ANDERSON-DARLING & CRAMER-VON MISES STATISTICS FOR THE
*             SAMPLE. REPEAT THIS OPERATION 5000 TIMES FOR EACH SAMPLE
*             SIZE OF 5,10,15,20,25, & 30
*
*   MM=      SAMPLE MEAN
*   SS=      SAMPLE STANDARD DEVIATION
*   M2=      ITERATIVE EST OF MEAN RETURNED BY MLE
*   SD2=     ITERATIVE EST OF STANDARD DEVIATION RETURNED BY MLE
*   N=       # IN SAMPLE
*   X(J)=    ARRAY OF LOGISTIC DEVIATES
*   Z(J)=    ARRAY OF STANDARDIZED DEVIATES
*   PIE=     ARITHMATIC VALUE
*   CONST=   ARITHMATIC VALUE
*   SEED=    INITIAL VALUE FOR RANDOM NUMBER GENERATOR
*   GGUBFS=  PSUEDORANDOM NUMBER GENERATOR (IMSL)
*   VSRTA=   SORT ROUTINE (IMSL)
*   FX=      HYPOTHETICAL SAMPLE DISTRIBUTION
*   SX=      EMPIRICAL SAMPLE DISTRIBUTION
*
*****
*
*   REAL KS,A2,W2,AA,WW,BB,SSQ,X(30),Z(30),FX(30),SX(30)
*   REAL XXKS(5001),XXA2(5001),XXW2(5001),YY(5001)
*   REAL M2,SD2,MM,SS,SUM1,SUM2,PIE,CONST,DIFF(30)
*   INTEGER N,IER,LOOP,REP
*   DOUBLE PRECISION SEED
*   EXTERNAL GGUBFS, VSRTA
*   COMMON N,CONST,PIE,LOOP
*
*****
*   ***** INITIALIZATION *****
*
*   OPEN(3,FILE='VALOUT',STATUS='NEW',FORM='UNFORMATTED')
*   WRITE(3,101)
101  FORMAT(2X,' CRITICAL VALUES FOR LOGISTIC(100,625) FOR JOHN5 ')
*
*   SEED=453689621.D0
*   CONST=3.**0.5
*   PIE=3.14159265358
*   N=0
9995 N=N+5
*   IF(N.GT.30)GOTO 9999
*   WRITE(3,511)N
511  FORMAT(////2X,' SAMPLE SIZE= ',15)
*   DO 5 I=1,N
*       X(I)=0.
5     CONTINUE

```

```

*
***** START MAJOR DO LOOP OF 5000 REPS *****
*
      REP=5000
      DO 4 I=1,REP
          XXKS(I)=0.
          XXA2(I)=0.
          XXW2(I)=0.
4      CONTINUE
*
      DO 10 LOOP=1,REP
*
***** START MINOR DO LOOPS BASED ON N *****
*
***** CALCULATE LOGISTIC DEVIATES
*
          RN=0.
          AA=0.
          COUNT=0
2000      DO 20 J=1,N
              RN=GGUBFS(SEED)
              AA=LOG((1.-RN)/RN)
              X(J)=100.0-(((25.0*CONST)*AA)/PIE)
20      CONTINUE
*
***** INITIAL ESTIMATES USING SAMPLE MEAN & STD DEV
*
          SUM1=0.
          SUM2=0.
          DO 30 J=1,N
              SUM1=SUM1+X(J)
30      CONTINUE
*
          MM=0.
          MM=SUM1/N
*
          BB=0.
          DO 40 J=1,N
              BB=(X(J)-MM)**2.
              SUM2=SUM2+BB
40      CONTINUE
*
          SSQ=0.
          SS=0.
          SSQ=SUM2/(N-1)
          SS=SSQ**0.5
*
***** CALL MLE SUBROUTINE
*
          CALL MLE(X,MM,SS,M2,SD2,IER)
509      FORMAT(/2X,' REP= ',I5)
          IF(IER.EQ.1)THEN
              COUNT=COUNT+1
              WRITE(3,509)LOOP

```

```

        WRITE(3,520)
520      FORMAT(2X,' MLE COULD NOT CONVERGE, TRYING NEW SAMPLE ')
        GOTO 2000
      ENDIF
*
      COUNT1=0
      IF( IER.EQ.3) THEN
        COUNT1=COUNT1+1
        WRITE(3,509) LOOP
        WRITE(3,522)
522      FORMAT(2X,' MLE EST OF S.D. LE 0.1 ')
        GOTO 2000
      ENDIF
*
***** CALCULATE F(X) AND S(X)
*
      CALL VSRTA(X,N)
*
      DO 88 K=1,N
        FX(K)=0.
        Z(K)=0.
88      CONTINUE
*
      DO 50 J=1,N
        Z(J)=(PIE*(X(J)-M2))/(SD2*CONST)
        FX(J)=(1./(1.+EXP(-Z(J))))
50      CONTINUE
*
***** CALCULATE KS, A(SQ) & W(SQ)
*
      CALL STATS(DIFF,FX,A2,W2,KS)
      XXKS(LOOP)=KS
      XXA2(LOOP)=A2
      XXW2(LOOP)=W2
*
***** END OF MAJOR LOOP
*
10      CONTINUE
*
      CALL VSRTA(XXKS,REP)
      CALL VSRTA(XXA2,REP)
      CALL VSRTA(XXW2,REP)
*
      YY(REP+1)=1.0
      DO 90 K=1,REP
        YY(K)=(K-.5)/REP
90      CONTINUE
*
      CALL ENDPT(REP,XXKS,YY,POINT)
      XXKS(REP+1)=POINT
*
      CALL ENDPT(REP,XXA2,YY,POINT)
      XXA2(REP+1)=POINT
*
      CALL ENDPT(REP,XXW2,YY,POINT)

```

```
XXW2(REP+1)=POINT
*
  CALL CV(REP,XXKS,YY)
  CALL CV(REP,XXA2,YY)
  CALL CV(REP,XXW2,YY)
*
  GOTO 9995
9999 ENDFILE(3)
      CLOSE(3)
      END
```

```

*****
*****
*****
*****
*
*   DEFINITION OF VARIABLES IN SUBROUTINE MLE
*
*   PURPOSE: TO FIND A SOLUTION TO  $L'(X)=0$  GIVEN INITIAL
*   APPROXIMATIONS M2 & SD2
*
*   MM=      SAMPLE MEAN
*   SS=      SAMPLE STD DEV
*   M1=      OLD ITERATIVE EST MEAN
*   M2=      NEW ITERATIVE EST MEAN
*   M3=      CANDIDATE EST MEAN
*   SD1=     OLD ITERATIVE EST STD DEV
*   SD2=     NEW ITERATIVE EST STD DEV
*   SD3=     CANDIDATE EST STD DEV
*   PIE=     ARITHMATIC VALUE
*   CONST=   ARITHMATIC VALUE
*   X(N)=    ARRAY OF DEVIATES
*   TOL=     TOLERANCE FOR 5 SIGNIFICANT DIGITS
*   Q(1)=    L'(M1) SD2 KNOWN
*   Q(2)=    L'(M2) SD2 KNOWN
*   Q(3)=    L'(SD1) M2 KNOWN
*   Q(4)=    L'(SD2) M2 KNOWN
*   AE1&2=   ABSOLUTE ERRORS
*   N=       SAMPLE SIZE
*   IER=     CONVERGENCE ERROR
*
*****
*
*   SUBROUTINE MLE(X,MM,SS,M2,SD2,IER)
*   REAL MM,SS,M1,M2,M3,SD1,SD2,SD3,PIE,X(N),TOL,CONST
*   REAL CC,DD,EE,GG,HH,Q(4),SUM(6),AE1,AE2
*   INTEGER N,IER
*   COMMON N,CONST,PIE,LOOP
*
*   ***** INITIAL VALUES
*
*       TOL=0.000005
*       M1=0.
*       M2=MM
*       M3=0.
*       SD3=0.
*       SD2=SS
*       SD1=SD2+0.5
*
*   ***** START SECANT MTHD
*
*       DO 5 K=1,50
*
*       SUM(1)=0.
*       SUM(2)=0.

```

```

DD=0.
EE=0.
*
DO 10 I=1,N
    DD=((X(I)-M1)*PIE)/(SD2*CONST)
    EE=((X(I)-M2)*PIE)/(SD2*CONST)
    SUM(1)=SUM(1)+(EXP(-DD)/(1+EXP(-DD)))
    SUM(2)=SUM(2)+(EXP(-EE)/(1+EXP(-EE)))
10 CONTINUE
*
***** 1ST DERIVATIVE VALUES FOR MEAN
*
CC=0.
CC=PIE/(SD2*CONST)
*
Q(1)=0.
Q(2)=0.
Q(1)=CC*(N-(2.0*SUM(1)))
Q(2)=CC*(N-(2.0*SUM(2)))
*
M3=M2-(Q(2)*((M2-M1)/(Q(2)-Q(1))))
AE1=ABS(M3-M2)
*
M1=M2
M2=M3
*
SUM(3)=0.
SUM(4)=0.
SUM(5)=0.
SUM(6)=0.
GG=0.
HH=0.
*
DO 20 J=1,N
    GG=((X(J)-M2)*PIE)/(SD1*CONST)
    HH=((X(J)-M2)*PIE)/(SD2*CONST)
    SUM(3)=SUM(3)+GG
    SUM(4)=SUM(4)+HH
    SUM(5)=SUM(5)+((GG*EXP(-GG))/(1.+EXP(-GG)))
    SUM(6)=SUM(6)+((HH*EXP(-HH))/(1.+EXP(-HH)))
20 CONTINUE
*
***** 1ST DERIVATIVE VALUES FOR STD DEV
*
Q(3)=0.
Q(4)=0.
Q(3)=(1./SD1)*(SUM(3)-(2.0*SUM(5))-N)
Q(4)=(1./SD2)*(SUM(4)-(2.0*SUM(6))-N)
*
SD3=SD2-(Q(4)*((SD2-SD1)/(Q(4)-Q(3))))
AE2=ABS(SD3-SD2)
*
*
IF(SD3.LE.0.1)THEN
    IER=3

```

```

        GOTO 1200
ENDIF
*
SD1=SD2
SD2=SD3
*
IF(AE1.LT.TOL.AND.AE2.LT.TOL)GOTO 1000
5 CONTINUE
*
IF(K.GE.50.AND.AE1.GT.TOL.OR.AE2.GT.TOL)THEN
    IER=1
    GOTO 1200
ENDIF
1000 IER=0
*
1200 RETURN
*
END

```

```

*****
*****
*****
*
*   DEFINITION OF VARIABLES IN SUBROUTINE STATS
*
*   PURPOSE: TO CALCULATE KOLMOGOROV-SMIRNOV, ANDERSON-
*   DARLING, & CRAMER-VON MISES STATISTICS FOR A GIVEN SAMPLE
*
*   N=      # IN SAMPLE
*   FX=     HYPOTHETICAL DISB OF SAMPLE
*   SX=     EMPIRICAL DISB OF SAMPLE
*   KS=     KOLMOGOROV-SMIRNOV STATISTIC
*   A2=     ANDERSON-DARLING STATISTIC
*   W2=     CRAMER-VON MISES STATISTIC
*
*****
*
*   SUBROUTINE STATS(DIFF,FX,A2,W2,KS)
*   REAL A2,KS,W2,FX(N),DIFF(N),ONE,TWO
*   INTEGER N
*   COMMON N,CONST,PIE,LOOP
*
*   ***** CALCULATE KS STATISTIC
*
*   DO 2 I=1,N
*       DIFF(I)=0.
*   2   CONTINUE
*
*   DO 10 I=1,N
*       ONE=0.
*       TWO=0.
*       Q=I*1
*       QQ=(Q-1)/N
*       QQQ=Q/N
*       ONE=ABS(FX(I)-QQ)
*       TWO=ABS(FX(I)-QQQ)
*       IF(ONE.GT.TWO)THEN
*           DIFF(I)=ONE
*       ELSE
*           DIFF(I)=TWO
*       ENDIF
*   10  CONTINUE
*
*   KS=0.
*   DO 20 I=1,N
*       IF(DIFF(I).GT.KS)KS=DIFF(I)
*   20  CONTINUE
*
*   ***** CALCULATE A(SQ) STATISTIC
*
*   SUM=0.
*   DO 30 I=1,N
*       SUM=SUM+((2.*I)-1.)*((LOG(FX(I)))+(LOG(1.-FX(N+1-I))))
*   30  CONTINUE

```

AD-A138 098

MODIFIED KOLMOGOROV-SMIRNOV ANDERSON-DARLING AND
CRAMER-VON MISES TESTS F... (U) AIR FORCE INST OF TECH
WRIGHT-PATTERSON AFB OH SCHOOL OF ENGI... J D YODER
16 DEC 83 AFIT/GOR/ENC/83D-7

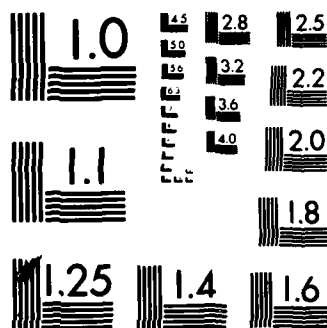
2/2

UNCLASSIFIED

F/G 12/1

NL





MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS-1963-A

```

*
  A2=0.
  A2=(-N)-(SUM/N)
*
*****  CALCULSTE W(SQ) STATISTIC
*
  SUM=0.
  DO 40 I=1,N
    SUM=SUM+(FX(I)-((2.*I)-1.)/(2.*N))**2.
40  CONTINUE
*
  W2=0.
  W2=(1./(12.*N))+SUM
*
  RETURN
  END

```

```

*****
*****
*
*   SUBROUTINE ENDPT: PURPOSE; TO DETERMINE THE LAST POINT
*   IN A GIVEN SERIES OF STATISTICS
*
*   REP=      # OF POINTS IN THE SERIES
*   M=        SLOPE
*   B=        INTERCEPT
*
*****
*
*   SUBROUTINE ENDPT(REP,XX,YY,POINT)
*   REAL XX(REP),YY(REP),POINT,M,B
*   INTEGER REP
*
*   M=0.
*   M=(YY(REP)-YY(REP-1))/(XX(REP)-XX(REP-1))
*   B=0.
*   B=YY(REP-1)-(M*XX(REP-1))
*   POINT=(1.-B)/M
*
*   RETURN
*   END
*****
*****
*
*   SUBROUTINE CV: PURPOSE; TO DETERMINE THE 80,85,90,95,99TH
*   PERCENTILES OF THE DISTRIBUTION OF
*   THE GIVBEN STATISTIC
*
*   REP=      # OF POINTS IN SERIES
*   M=        SLOPE
*   B=        INTERCEPT
*   PER=      PERCENTILE
*
*****
*
*   SUBROUTINE CV(REP,XX,YY)
*   REAL PER,YY(REP+1),XX(REP+1),M,B,P80,P85,P90,P95,P99
*   INTEGER REP
*
*   P80=0.
*   P85=0.
*   P90=0.
*   P95=0.
*   DO 10 J=80,95,5
*     DO 20 K=2,REP
*       IF(YY(K).GE.(J/100.))THEN
*         M=0.
*         B=0.
*         PER=0.
*         M=(YY(K)-YY(K-1))/(XX(K)-XX(K-1))
*         B=YY(K-1)-(M*XX(K-1))
*         PER=((J/100.)-B)/M

```

```

                GOTO 666
            ENDIF
20      CONTINUE
666     IF(J.EQ.80)P80=PER
        IF(J.EQ.85)P85=PER
        IF(J.EQ.90)P90=PER
        IF(J.EQ.95)P95=PER
10      CONTINUE
*
        P99=0.
        DO 30 K=2,REP
            IF(YY(K).GE.0.99)THEN
                M=0.
                B=0.
                M=(YY(K)-YY(K-1))/(XX(K)-XX(K-1))
                B=YY(K-1)-(M*XX(K-1))
                PER=(0.99-B)/M
                GOTO 777
            ENDIF
30      CONTINUE
*
777     P99=PER
*
        WRITE(3,800)
800     FORMAT(///2X,' 80,85,90,95,99 PERCENTILES  ')
        WRITE(3,802)P80,P85,P90,P95,P99
802     FORMAT(2X,5(F10.7,2X))
*
        RETURN
        END

```

```

PROGRAM POWER
REAL KS,A2,W2,AA,WW,BB,SSQ,X(30),Z(30),FX(30),Y(30)
REAL M2,SD2,MM,SS,SUM1,SUM2,PIE,CONST,DIFF(30)
REAL KS95(6),A295(6),W295(6)
INTEGER N,IER,LOOP,REP
DOUBLE PRECISION SEED
EXTERNAL GGUBFS,VSRTA,GGEXN,GGNQF
COMMON N,CONST,PIE,LOOP

*
***** INITIALIZATION *****
*
      OPEN(3,FILE='NORMOUT',STATUS='NEW',FORM='UNFORMATTED')
      WRITE(3,101)
101  FORMAT(2X,' CRITICAL VALUES FOR NORMAL ')
*
      SEED=453689621.D0
      CONST=3.**0.5
      PIE=3.14159265358
      N=0
9995  N=N+5
      IF(N.GT.30)GOTO 9999
      WRITE(3,511)N
511  FORMAT(////2X,' SAMPLE SIZE= ',I5)
      DO 5 I=1,N
          X(I)=0.
          Y(I)=0.
          Z(I)=0.
5      CONTINUE
      DO 6 I=1,6
          KS95(I)=0.
          A295(I)=0.
          W295(I)=0.
6      CONTINUE
*
      REP=1000
*
      DO 10 LOOP=1,REP
*
2000      DO 20 J=1,N
          Y(I)=GGNQF(SEED)
          X(I)=(Y(I)*2.)+1.
20      CONTINUE
*
***** INITIAL ESTIMATES USING SAMPLE MEAN & STD DEV
*
      SUM1=0.
      SUM2=0.
      DO 30 J=1,N
          SUM1=SUM1+X(J)
30      CONTINUE
*
      MM=0.
      MM=SUM1/N
*
      BB=0.

```

```

DO 40 J=1,N
  BB=(X(J)-MM)**2.
  SUM2=SUM2+BB
40 CONTINUE
*
  SSQ=0.
  SS=0.
  SSQ=SUM2/(N-1)
  SS=SSQ**0.5
*
***** CALL MLE SUBROUTINE
*
  CALL MLE(X,MM,SS,M2,SD2,IER)
  IF(IER.EQ.1)GOTO 2000
  IF(IER.EQ.3)GOTO 2000
*
***** CALCULATE F(X) AND S(X)
*
  CALL VSRTA(X,N)
*
DO 85 J=1,N
  Z(J)=(PIE*(X(J)-M2))/(SD2*CONST)
  FX(J)=(1./(1.+EXP(-Z(J))))
85 CONTINUE
*
***** CALCULATE KS, A(SQ) & W(SQ)
*
  CALL STATS(DIFF,FX,A2,W2,KS)
  IF(N.EQ.5.AND.KS.GT.0.309)KS95(1)=KS95(1)+1
  IF(N.EQ.5.AND.A2.GT.0.620)A295(1)=A295(1)+1
  IF(N.EQ.5.AND.W2.GT.0.096)W295(1)=W295(1)+1
  IF(N.EQ.10.AND.KS.GT.0.230)KS95(2)=KS95(2)+1
  IF(N.EQ.10.AND.A2.GT.0.640)A295(2)=A295(2)+1
  IF(N.EQ.10.AND.W2.GT.0.095)W295(2)=W295(2)+1
  IF(N.EQ.15.AND.KS.GT.0.191)KS95(3)=KS95(3)+1
  IF(N.EQ.15.AND.A2.GT.0.665)A295(3)=A295(3)+1
  IF(N.EQ.15.AND.W2.GT.0.098)W295(3)=W295(3)+1
  IF(N.EQ.20.AND.KS.GT.0.170)KS95(4)=KS95(4)+1
  IF(N.EQ.20.AND.A2.GT.0.663)A295(4)=A295(4)+1
  IF(N.EQ.20.AND.W2.GT.0.097)W295(4)=W295(4)+1
  IF(N.EQ.25.AND.KS.GT.0.655)KS95(5)=KS95(5)+1
  IF(N.EQ.25.AND.A2.GT.0.663)A295(5)=A295(5)+1
  IF(N.EQ.25.AND.W2.GT.0.098)W295(5)=W295(5)+1
  IF(N.EQ.30.AND.KS.GT.0.141)KS95(6)=KS95(6)+1
  IF(N.EQ.30.AND.A2.GT.0.649)A295(6)=A295(6)+1
  IF(N.EQ.30.AND.W2.GT.0.098)W295(6)=W295(6)+1
*
*
***** END OF MAJOR LOOP
*
10 CONTINUE
*
DO 89 K=1,6
  A=0.

```

```

      B=0.
      C=0.
      A=KS95(K)/1000.
      B=A295(K)/1000.
      C=S295(K)/1000.
      W=K*5
      WRITE(3,511)W
      WRITE(3,888)A,B,C
888   FORMAT(//2X,' KS%= ',F6.3,' A2%= ',F6.3,' W2%= ',F6.3)
89    CONTINUE
      GOTO 9995
9999  ENDFILE(3)
      CLOSE(3)
      END

```

```

*****
      SUBROUTINE MLE(X,MM,SS,M2,SD2,IER)
      REAL MM,SS,M1,M2,M3,SD1,SD2,SD3,PIE,X(N),TOL,CONST
      REAL CC,DD,EE,GG,HH,Q(4),SUM(6),AE1,AE2
      INTEGER N,IER
      COMMON N,CONST,PIE,LOOP
*
***** INITIAL VALUES
*
      TOL=0.000005
      M1=0.
      M2=MM
      M3=0.
      SD3=0.
      SD2=SS
      SD1=SD2+0.5
*
***** START SECANT MTHD
*
      DO 5 K=1,50
*
      SUM(1)=0.
      SUM(2)=0.
      DD=0.
      EE=0.
*
      DO 10 I=1,N
          DD=((X(I)-M1)*PIE)/(SD2*CONST)
          EE=((X(I)-M2)*PIE)/(SD2*CONST)
          SUM(1)=SUM(1)+(EXP(-DD)/(1+EXP(-DD)))
          SUM(2)=SUM(2)+(EXP(-EE)/(1+EXP(-EE)))
10      CONTINUE
*
***** 1ST DERIVATIVE FUNCTIONAL VALUES FOR MEAN
*
      CC=0.
      CC=PIE/(SD2*CONST)
*
      Q(1)=0.
      Q(2)=0.
      Q(1)=CC*(N-(2.0*SUM(1)))
      Q(2)=CC*(N-(2.0*SUM(2)))
*
      M3=M2-(Q(2)*((M2-M1)/(Q(2)-Q(1))))
      AE1=ABS(M3-M2)
*
      M1=M2
      M2=M3
*
      SUM(3)=0.
      SUM(4)=0.
      SUM(5)=0.
      SUM(6)=0.
      GG=0.

```

```

      HH=0.
*
      DO 20 J=1,N
        GG=((X(J)-M2)*PIE)/(SD1*CONST)
        HH=((X(J)-M2)*PIE)/(SD2*CONST)
        SUM(3)=SUM(3)+GG
        SUM(4)=SUM(4)+HH
        SUM(5)=SUM(5)+((GG*EXP(-GG))/(1.+EXP(-GG)))
        SUM(6)=SUM(6)+((HH*EXP(-HH))/(1.+EXP(-HH)))
20    CONTINUE
*
***** 1ST DERIVATIVE FUNCTIONAL VALUES FOR STD DEV
*
      Q(3)=0.
      Q(4)=0.
      Q(3)=(1./SD1)*(SUM(3)-(2.0*SUM(5))-N)
      Q(4)=(1./SD2)*(SUM(4)-(2.0*SUM(6))-N)
*
      SD3=SD2-(Q(4)*((SD2-SD1)/(Q(4)-Q(3))))
      AE2=ABS(SD3-SD2)
*
*
      IF(SD3.LE.0.1)THEN
        IER=3
        GOTO 1200
      ENDIF
*
      SD1=SD2
      SD2=SD3
*
      IF(AE1.LT.TOL.AND.AE2.LT.TOL)GOTO 1000
5    CONTINUE
*
      IF(K.GE.50.AND.AE1.GT.TOL.OR.AE2.GT.TOL)THEN
        IER=1
        GOTO 1200
      ENDIF
1000 IER=0
*
1200 RETURN
*
      END

```

```

*****
*****
*
      SUBROUTINE STATS(DIFF,FX,A2,W2,KS)
      REAL A2,KS,W2,FX(N),DIFF(N),ONE,TWO
      INTEGER N
      COMMON N,CONST,PIE,LOOP
*
***** CALCULATE KS STATISTIC
*
      DO 2 I=1,N
        DIFF(I)=0.
2      CONTINUE
*
      DO 10 I=1,N
        ONE=0.
        TWO=0.
        Q=0.
        QQ=0.
        QQQ=0.
        Q=I*1
        QQ=(Q-1)/N
        QQQ=Q/N
        ONE=ABS(FX(I)-QQ)
        TWO=ABS(FX(I)-QQQ)
        IF(ONE.GT.TWO)THEN
          DIFF(I)=ONE
        ELSE
          DIFF(I)=TWO
        ENDIF
10     CONTINUE
*
      KS=0.
      DO 20 I=1,N
        IF(DIFF(I).GT.KS)KS=DIFF(I)
20     CONTINUE
*
***** CALCULATE A(SQ) STATISTIC
*
      SUM=0.
      DO 30 I=1,N
        SUM=SUM+((2.*I)-1.)*(LOG(FX(I)))+(LOG(1.-FX(N+1-I)))
30     CONTINUE
*
      A2=0.
      A2=(-N)-(SUM/N)
*
***** CALCULSTE W(SQ) STATISTIC
*
      SUM=0.
      DO 40 I=1,N
        SUM=SUM+(FX(I)-((2.*I)-1.)/(2.*N))**2.
40     CONTINUE
*
      W2=0.

```

W2=(1./(12.*N))+SUM

*

RETURN

END

Vita

John D. Yoder was born on 1 February 1950, in Portland, Oregon. After graduating from high school, he enlisted in the U.S. Air Force. He spent eight years in the communications service and eventually became a Base Communications Center Supervisor. He separated from the Air Force in 1976 in order to pursue his goal of a college degree. He obtained a Bachelor of Science degree in the Physical Sciences from Portland State University in 1979. He received a commission as a 2nd Lieutenant upon graduation and entered the Air Force again in January of 1980. His assignment prior to entering the Air Force Institute of Technology in June of 1982 was to the 3246th Test Wing, Eglin AFB where he was placed in charge of the Wings' Management Information System.

Permanent Address: 1355 Evergreen N.E.
Salem, Oregon 97301

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE

REPORT DOCUMENTATION PAGE

1a. REPORT SECURITY CLASSIFICATION UNCLASSIFIED			1b. RESTRICTIVE MARKINGS	
2a. SECURITY CLASSIFICATION AUTHORITY			3. DISTRIBUTION/AVAILABILITY OF REPORT Approved for public release; distribution unlimited	
2b. DECLASSIFICATION/DOWNGRADING SCHEDULE				
4. PERFORMING ORGANIZATION REPORT NUMBER(S) AFIT/ENR/ENC/83D-7			5. MONITORING ORGANIZATION REPORT NUMBER(S)	
6a. NAME OF PERFORMING ORGANIZATION School of Engineering Air Force Institute of Technology		6b. OFFICE SYMBOL (If applicable) AFIT/EN	7a. NAME OF MONITORING ORGANIZATION	
6c. ADDRESS (City, State and ZIP Code) Wright-Patterson AFB, Ohio 45433			7b. ADDRESS (City, State and ZIP Code)	
8a. NAME OF FUNDING/SPONSORING ORGANIZATION		8b. OFFICE SYMBOL (If applicable)	9. PROCUREMENT INSTRUMENT IDENTIFICATION NUMBER	
8c. ADDRESS (City, State and ZIP Code)			10. SOURCE OF FUNDING NOS.	
			PROGRAM ELEMENT NO.	PROJECT NO.
11. TITLE (Include security classification) See Doc 1j				
12. PERSONAL AUTHOR(S) John D. Roder, 1Lt., USAF				
13a. TYPE OF REPORT Tech. Thesis		13b. TIME COVERED FROM TO		14. DATE OF REPORT (Yr., Mo., Day) 83/12/16
15. PAGE COUNT 107				
16. SUPPLEMENTARY NOTATION By W. L. R. 10017 Rel. 84 Air Force Research and Development Report AFTR-83-10017				
17. COSATI CODES			18. SUBJECT TERMS (Continue on reverse if necessary and identify by block number)	
FIELD	GROUP	SUB. GR.	Statistical Functions, Probability Distribution Functions, Statistical Analysis Theory	
12	01			
19. ABSTRACT (Continue on reverse if necessary and identify by block number) MODIFIED KOLMOGOROV-SMIRNOV, ANDERSON-DARLING AND CRAMER-VON MISES TESTS FOR THE LOGISTIC DISTRIBUTION WITH UNKNOWN LOCATION AND SCALE PARAMETERS Capt. Brian Woodruff				
20. DISTRIBUTION/AVAILABILITY OF ABSTRACT UNCLASSIFIED/UNLIMITED ☒ SAME AS RPT. ☐ DTIC USERS ☐			21. ABSTRACT SECURITY CLASSIFICATION UNCLASSIFIED	
22a. NAME OF RESPONSIBLE INDIVIDUAL Brian Woodruff Assistant Professor			22b. TELEPHONE NUMBER (Include Area Code) 5-5533	22c. OFFICE SYMBOL AFIT/ENC

The method of maximum likelihood is used to determine invariant estimates of the unknown location and scale parameters of a sample from the Logistic distribution. The partial derivatives of the likelihood function can not be solved explicitly, therefore the Secant method is used to iteratively determine the roots of the partial derivatives. Using these estimates modified Kolmogorov-Smirnov, Anderson-Darling and Cramer-von Mises statistics are calculated for a given sample. This procedure is repeated 5000 times for sample sizes of $n = 5(5)30$. The 10th, 85th, 90th, 95th and 99th percentiles of the distribution of each statistic, for each sample size, is then calculated. These values are then used to generate tables of critical values for the Logistic distribution with unknown location and scale parameters. A power comparison between the three tests is performed using samples from various distributions.

The Secant method requires "good" initial estimates of the parameters in order to converge. This thesis uses the sample mean and standard deviation as initial estimates. In four of the total 30,000 samples used, these initial estimates did not allow convergence. While discarding these samples biases the theoretical results, it was determined that discarding these samples would not bias the numerical results. This does however place a constraint on using the Secant method with respect to obtaining maximum likelihood estimates of the parameters. The power of these tests for non-symmetrically convex distributions is very good. However, for symmetrically convex distributions, the power ranges from moderate to only slightly more than the significance level.

END

FILMED

384

DTIC