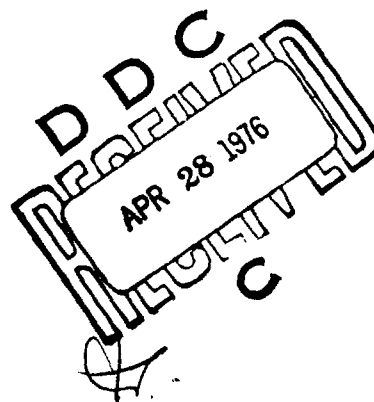


construction
engineering
research
laboratory

8 *12*
SPECIAL REPORT M-209
March 1977
Seismic Design Criteria for Interior Utility and Lifeline Systems

FRAGILITY DATA ANALYSIS AND TESTING GUIDELINES
FOR ESSENTIAL EQUIPMENT USED
IN CRITICAL FACILITIES

by
Paul N. Sonnenburg



Approved for public release; distribution unlimited.

ADA 038768

FILE COPY

The contents of this report are not to be used for advertising, publication, or promotional purposes. Citation of trade names does not constitute an official indorsement or approval of the use of such commercial products. The findings of this report are not to be construed as an official Department of the Army position, unless so designated by other authorized documents.

**DESTROY THIS REPORT WHEN IT IS NO LONGER NEEDED
DO NOT RETURN IT TO THE ORIGINATOR**

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER CERL-SR-M-209	2. GOVT ACCESSION NO.	3. RECIPIENT'S CATALOG NUMBER (9)
4. TITLE (and Subtitle) FRAGILITY DATA ANALYSIS AND TESTING GUIDELINES FOR ESSENTIAL EQUIPMENT USED IN CRITICAL FACILITIES		5. TYPE OF REPORT & PERIOD COVERED FINAL / Rept.
7. AUTHOR(s) (10) Paul N. Sonnenburg		6. PERFORMING ORG. REPORT NUMBER (17)
9. PERFORMING ORGANIZATION NAME AND ADDRESS CONSTRUCTION ENGINEERING RESEARCH LABORATORY P.O. Box 4005 Champaign, Illinois 61820		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS (16) 4A762731AT41 04-002
11. CONTROLLING OFFICE NAME AND ADDRESS		12. REPORT DATE (11) March 1977
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office) (12) 46p.		13. NUMBER OF PAGES 48
		15. SECURITY CLASS. (of this report) UNCLASSIFIED
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) Approved for public release; distribution unlimited.		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
18. SUPPLEMENTARY NOTES Copies are obtainable from National Technical Information Service, Springfield, VA 22151		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) fragility test specifications fragility data analysis critical facilities		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) Fragility data reports from tests of many items of tactical support equipment used at missile sites were reviewed to analyze failure character- istics. This type of equipment is closely related to that used in essential systems of critical facilities. Results were organized to formulate speci- fications for fragility test report requirements and guidelines for planning fragility tests for the essential equipment of utilities and lifelines used in critical facilities. A method of statistically analyzing		

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)

fragility data was developed to facilitate calculating the probability of failure when the cause of failure was not consistent or obvious. The proposed specifications and guidelines were to be incorporated into a more complete set of testing specifications to be finished at a later date.

UNCLASSIFIED

FOREWORD

This research was performed for the Directorate of Military Construction, Office of the Chief of Engineers (OCE) under Project 4A7-62731AT41, "Design, Construction, and Operations and Maintenance Technology for Military Facilities"; Task 04, "Construction Systems Technology"; Work Unit 002, "Seismic Design Criteria for Interior Utility and Lifeline Systems." The QCR number is 1.03.003. Mr. George Matsumura is the OCE Technical Monitor.

The work was performed by the Structural Mechanics Branch (MSS), Materials and Science Division (MS), U.S. Army Construction Engineering Research Laboratory (CERL). Dr. P. N. Sonnenburg was the Principal Investigator for this project. Dr. W. E. Fisher is Chief of MSS, and Dr. G. R. Williamson is Chief of MS.

COL J. E. Hays is Commander and Director of CERL, and Dr. L. R. Shaffer is Technical Director.

ACCESSION #	
NTIS	✓
DDC	✓
UNANNOUNCED	✓
JUSTIFICATION	
BY	
DISTRIBUTION/AVAILABILITY CODES	
DECL.	AVAIL. AND SP. CHAR.
A	

CONTENTS

	DD FORM 1473	1
	FOREWORD	3
	LIST OF FIGURES AND TABLES	5
1	INTRODUCTION	7
	Background	
	Purpose	
	Outline of Report	
2	PROPOSED SPECIFICATIONS.	9
	Definitions	
	Applicability	
	Test Report Requirements	
	Test Summary Format	
	Hardness Assurance or Assessment	
	Guidelines for Selection of Test Levels	
3	OBSERVED FAILURE INFORMATION	20
	Summary of SAFEGUARD Failure Data	
	Types of Failures	
	Typical Failure Modes	
4	SUMMARY AND CONCLUSIONS.	24
	Summary	
	Conclusions	
	APPENDIX A: Overview of Related Destructive Testing Methods	26
	APPENDIX B: Statistical Analysis of Fragility Data (Independent Failures)	29
	REFERENCES	48
	DISTRIBUTION	

LIST OF FIGURES

<u>Number</u>		<u>Page</u>
1	Typical Test Summary Format	13
2	Probability of Failure	15
3	Average Results for Estimator and Accuracy, Independent Failures	18
A1	Fatigue Data Presentation	27
A2	Distribution and Density Functions for Consistent and Independent Failures	28
B1	Parameter and Accuracy for Exponential Distribution	41
B2	Parameter and Accuracy for the One-Parameter Pareto Distribution	42
B3	Parameter and Accuracy for the One-Parameter Limited Distribution	43
B4	α -Parameter and Accuracy for the Two-Parameter First Asymptotic Distribution	44

LIST OF TABLES

1	Failure Summary From General Equipment Tests	21
B1	Analysis of Motor Control Center Data Using Exponential Distribution	46

FRAGILITY DATA ANALYSIS AND TESTING
GUIDELINES FOR ESSENTIAL EQUIPMENT
USED IN CRITICAL FACILITIES

1 INTRODUCTION

Background

The necessity for critical facilities to provide post-earthquake recovery services requires survival not only of the building structures, but also of the utility and lifeline systems which support the functions most needed after an event. Assuming the structure will survive, the major problem is the isolation or hardening of all systems, subsystems, equipment, and components needed to support the essential functions.

Before a decision is made to harden or isolate certain equipment, it is necessary to determine the fragility properties of the equipment. Some components may be so delicate that they must be isolated, while others may withstand the environment with no degradation. It is presently too costly and otherwise impractical to require that all essential equipment of critical facilities be tested for fragility. However, testing of related off-the-shelf equipment has been done for the assessment of hardness of tactical support equipment at missile sites. A considerable amount of experience can be gained from the analysis of the resulting hardness assessment reports. The major results of analyzing these reports are recorded in this work, and form the basis for proposing some preliminary specifications and guidelines for fragility testing for essential equipment of critical facilities.

Purpose

The purpose of this report is to present preliminary guidelines and test report specifications derived from analyzing fragility data test results of the type of equipment used in essential systems of critical Army facilities.

Outline of Report

Chapter 2 presents the primary results of this work in the form of proposed testing and test report specifications. Other guidelines provided will enable the test engineer to categorize failures and identify problems in time to reschedule testing when certain typical test results are encountered.

Chapter 3 analyzes many fragility reports from which the specifications in Chapter 2 were formed. Data were collected from pertinent off-the-shelf equipment similar to the type of equipment found in essential systems of utilities and lifelines of critical facilities. Appendix A describes the fundamental theoretical principles of fragility testing in perspective with other forms of destructive testing more familiar to engineers. Appendix B provides a method for analyzing fragility test data to determine probability of failure. This method was applied to hypothetical failure data to establish the guidelines.

2 PROPOSED SPECIFICATIONS

The specifications proposed in this section are directed toward improving fragility testing and test report procedures. They have been assembled from an analysis of failure data contained in the Army Corps of Engineers Huntsville Division's technical reports on SAFEGUARD systems testing; however, they are not yet complete. Completion would entail the generation of specifications for controlling the accuracy of shock environments and instrumentation measurements, and for providing guidelines to selecting appropriate shock environments for various classes of equipment. Ultimately, specifications will be completed for testing, analysis, design, and procurement of essential systems and components.

Definitions

Test Unit: Lists of essential equipment are usually displayed according to systems, subsystems, and components. However, for a particular test setup or analysis, it is more convenient to refer to a complete unit of equipment. A unit is defined as any system, subsystem, component, or combination of these which may be treated independently, either as an assembly or as a detailed part with a specific function. For example, individual valves in a piping system may be sufficiently rugged to avoid testing each one. Since the weakest points may be the joints of the valves and the piping, it may be desirable to test or analyze the entire piping system as a unit. The size of a test unit is limited to its capability of being subjected to a single defined shock environment; i.e., a piping system of a building would be too large to test the entire unit.

Hardness: A unit's hardness is defined as its probability of failure under expected environmental loading conditions.

Fragility: A unit's fragility is defined by stating the value of a dependent variable, such as acceleration, at which it will fail.

Fragility Envelope: A fragility envelope is determined by expressing the dependent variable descriptive of failure as a function of frequency.

Test Capability: Test capability is the capability to duplicate a prescribed shock environment for a test unit and to provide instrumentation for monitoring the shock environment and response of the unit in accordance with MIL-C-45662, "Calibration System Requirements" for instrumentation standards.

Test Analysis Capability: Test analysis capability is the capability to process and reduce failure data and to calculate the probability of failure of a test unit under prescribed loading conditions. This operation usually requires a computer facility and personnel knowledgeable in statistical analysis.

Hardness Assessment: The assessment of a unit's hardness is achieved by calculating the unit's probability of failure under (possibly numerous) specified shock loading conditions.

Hardness Assurance: The assurance of hardness is achieved by reducing the probability of a unit's failure below an acceptable value. This can be accomplished either by redesign to eliminate failures under a prescribed shock environment, or by isolating the unit from the prescribed shock environment.

Test Level: The shock environment for fragility testing is usually defined as a shaped shock spectrum, referred to as the 100 percent level. Tests may be conducted with the same spectrum shape, but with a uniform amplitude change in parameters. The level of the test refers to the amplitude of a single parameter such as displacement, velocity, or acceleration, which can be used for comparison with the 100 percent level. The same terminology is used for application to any type of shock environment.

Flaw: A flaw is a fault in a unit which is sensitive to one or more variables such as acceleration, velocity, or displacement. The flaw will cause the unit to fail if one of these variables exceeds a certain limit.

Flaw Level: The lowest limit of the variable which will cause the unit to fail is the flaw level; thus, if the test level exceeds the flaw level, the test unit will fail.

Consistent Failure: When a failure occurs repeatedly and can be predicted at a well-defined test level value, the flaw or failure is said to be consistent. A unit will exhibit 100 percent probability of failure at test levels at or above this well-defined level, and the failure will be caused by the same (consistent) flaw. Numerous consistent failures may occur simultaneously. A fragility envelope can be formed from flaw levels if the flaws are consistent.

Independent Failure: When a failure occurs erratically, at different levels of the variables, the flaw is said to be independent. A well-defined fragility envelope cannot be formed from flaw levels if the flaws are independent. When this situation occurs, it is necessary to consider the probability of failure of a unit at any test level.

Qualifying Failure: A qualifying failure is one which can be corrected almost immediately, or which has a degrading influence not directly affecting the unit's function or that of any other interfacing unit.

Lingering Failure: A lingering failure is one which requires an intolerable time delay to correct. The failure may occur directly

within the test unit, or a faulty output of the unit may cause an interfacing system to malfunction.

Applicability

These specifications should apply whenever testing is authorized for components of essential systems and functions of critical facilities. Equipment manufacturers who do not have the specified testing capability should not be authorized to conduct such tests. Likewise, a statement about the probability of a component's failure should not be requested unless the test analysis capability exists.

Test Report Requirements

Care must be taken to insure that test data can be interpreted with a minimum of subjectivity. If test capability exists and authority has been granted to conduct fragility tests, a complete test report should be required. The cost of such a report is expected to be relatively small compared to the costs of equipment, time, and labor. A typical report should include the following:

1. Purpose. Provide a detailed description of why the unit must be tested and what results are considered important.
2. Authorization. Provide the authorizing agency, the funding, and the time schedule restrictions.
3. Unit Description. Describe the unit to be tested in terms of where or how it fits into a critical system or subsystem, and what components are to be tested. Describe critical interfaces with other systems, and define what structural or functional aspects constitute failure.
4. Ambient Conditions. Describe the operating conditions under which the unit must be tested, including such information as pressures, temperatures, flow rates, etc.
5. Loading Requirements. Loading requirements should be specified by contract before testing is authorized. The required 100 percent test level should be recorded as specified, usually in terms of shock spectrum. Variations of this loading for special conditions should be stated in detail.
6. Test Machine Description. Describe the dynamic capabilities of the testing machine in terms of maximum displacement, velocity, acceleration, frequency range, table weight, test mass weight, or any other pertinent specifications.

7. Method of Procedures. Describe all testing methods and all detailed procedures for conducting the tests. The procedures should be recorded as testing progresses. Each action should be listed on a tabulated form so that events can be read in chronological order; in particular, all anomalies, failures, and corrective actions should be described.

8. Test Summary Tabulation. A standard tabulation format should be used to summarize test results. This format, as described in the next section, should include sufficient information to assess the hardness of the unit without further reference to detailed methods and procedures.

9. Appendix. All raw test data should be shown in a tabular format.

Test Summary Format

The purpose of the test summary format is to provide a standard presentation from which enough significant information can be extracted to facilitate the unit's hardness assessment. Figure 1 illustrates a suitable format and typical entries. The following minimum information should be provided.

1. Heading information should include a title, description of the unit, the test machine used, and the date or span of dates over which testing was conducted.

2. The axis or axes of testing must be identified.

3. The tests should be identified by number in chronological order.

4. The loading information should be coded for reference to a time history, shock spectrum, or other authorized loading requirement shown elsewhere in the report; the percent of full-scale level should also be shown.

5. A brief description of every failure should be entered. Corrective action need not be listed, since it should be provided elsewhere in the report.

6. The test engineer should enter an opinion about whether the unit's failure is qualifying (Q) or lingering (F). When there is sufficient doubt, the engineer should consult an expert.

7. The test engineer should enter an opinion about whether the failure is consistent (C) or independent (I). Again consultation with an expert may be necessary. This information is not complete until the

UNIT TITLE: <u>Air Compressor Control Panel</u>			DATES: Jan 15-18, 1974	
TEST MACHINE: <u>C.O.E. B, Uniaxial</u>				
TEST	AXIS	LOADING	FAILURE DESCRIPTION	CLASSIFICATION
				DELAY CONSISTENCY
5	X	50%	None	
::				
9	Y	50%	Hi-after cooler temp. fault trip	Q I
::				
::				
14	Y	100%	Fault on disable vibration switch	F C

Figure 1. Typical test summary format.

scheduled testing is finished and the types of failures are reviewed. Further testing may be recommended if independent failures are recognized.

Hardness Assurance or Assessment

If testing will be conducted by the unit's manufacturer, achieving hardness assurance may be possible, since in this case, all consistent failures should be identified and eliminated by redesign, and the unit retested for verification of hardness. The report should include the results of the original design's test if the unit is already operational at any critical facility. Authorization for the mass production and purchasing of hardened units must not be automatically assumed by a manufacturer, since further contract negotiations will be necessary before hardening all production units.

If independent failures are identified, hardness assurance may be difficult or impossible to achieve with available time and funds. In this case the testing facility should request authority for more extensive testing than originally planned. The goal will then be to collect enough test data to provide a reasonably accurate hardness assessment, since hardness assurance may be beyond consideration. If test analysis capability for hardness assessment does not exist and independent failures occur, the hardness assessment and calculation of the probability of failure may be accomplished later if a complete and accurate report is provided.

Guidelines for Selection of Test Levels

It is never known before testing whether consistent or independent failures (or both) will occur in a hardness test program. When a consistent failure occurs at or below the expected environment level, the probability of failure is 100 percent. Hence, the decision must be made either to temporarily condone the existence of the flaw, to harden the unit by eliminating the flaw, or to isolate the unit to prevent damage. When independent failures occur, attempts to harden may be fruitless, and isolation of the unit impractical. In this case, the probability of failure will be less than or equal to 100 percent, and it is usually advisable to collect sufficient test data to enable a reasonably accurate calculation of the probability of failure at appropriate test levels.

In concept, the probability of failure of a unit can be plotted as a function of test level, as shown in Figure 2. It is convenient (and recommended) to select the exponential distribution as the estimated shape of this curve for many units. In this case, the mathematical relation is written as

$$F(x) = 1 - e^{-\alpha x} \quad [\text{Eq 1}]$$

where $F(x)$ = probability of failure

x = test level

α = open parameter, to be determined from failure data

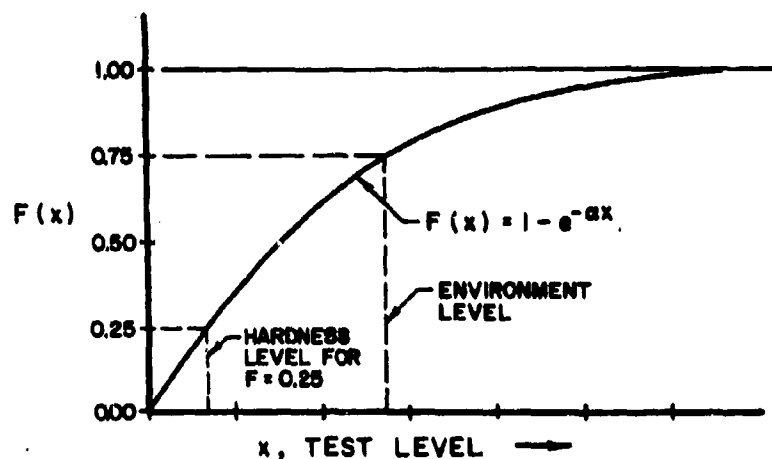


Figure 2. Probability of failure.

The expected environment level may be shown on this plot as a vertical line at the appropriate value of x . By plotting a horizontal line at the intersection of the environment level and the curve for $F(x)$, the probability of failure can be stated. In Figure 2, a 75 percent probability of failure is indicated at the environment level. If an established criterion dictates that the probability of failure shall not exceed 25 percent, this will be represented by a second horizontal line at $F(x) = 0.25$, indicating that the unit has a hardness level somewhat less than the expected environment level for this criterion.

Selection of Test Levels

It should be obvious from Figure 2 that an improved value of α can be estimated if test levels are scheduled so that somewhat uniform values of $F(x)$ are obtained throughout the range $0 < F(x) < 1$. If Eq 1 is solved for test level in terms of F , the result is

$$x = -\frac{1}{\alpha} \ln (1 - F). \quad [\text{Eq 2}]$$

Hence, if α were known, Eq 2 could be used (for the exponential distribution) to solve for optimum test levels by selecting equal increments of F between zero and unity. Of course, α must be taken from the data and is never known in advance. However, it can be seen from Eq 2 that a logarithmic increase in test levels would lead to a more

efficient and accurate estimate of α than other methods. It is common, for example, to increase test levels in equal increments, which could be much less efficient. If an investigator knows beforehand what distribution function he* would like to apply, an equation like Eq 2 could be derived to show what variation in test levels he should use for best efficiency. Such prior knowledge is rare, however, so the simple exponential distribution and the logarithmic increase in test levels is recommended.

A major concern is establishment of an upper test level bound. It is never known beforehand whether the full-scale input shock environment will exceed or fall below the unit's required probability of failure level. The most accurate calculation of open parameters (α , for the exponential distribution) is obtained in the undesirable case when there is a high probability of failure at test levels less than or equal to the full-scale environmental level. Because of time and funding limitations, testing above the full environmental level may not be authorized. Therefore, the least accurate calculation of parameters results when there is a low probability of failure below the full required level.

Numerous hypothetical studies indicate that failure cannot be predicted accurately, for independent failures, unless a unit is tested throughout the full range for $0 < F < 1$. If it is necessary to test above the full environmental level to experience more failures for this purpose, the authority must be granted individually. Thus, the upper test level bound will be a function of the required environmental level, the number of failures experienced in testing up to this level, and the time and funds available to test above this level.

It has also been observed from many hypothetical tests that accuracy in estimating the parameters is increased somewhat if testing is not conducted repeatedly at the same level. That is, for a given number of tests, accuracy is improved by testing at *different* levels, preferably in the range where failures can occur at least 50 percent of the time.

Number of Test Levels

The accuracy with which the distribution parameters can be estimated will also increase as the number of tests increases. Studying the results of hypothetical failure tests through the entire probability range provides a good indication of the minimum number of tests needed to obtain a given degree of accuracy. To accomplish this, three candidate one-parameter distributions were studied in which the number of tests in the range $0 < F < 1$ was varied from 7 to 100. In each case, 100 trials were made; i.e., 100 trials were

*"He" is used throughout this report to represent both masculine and feminine genders.

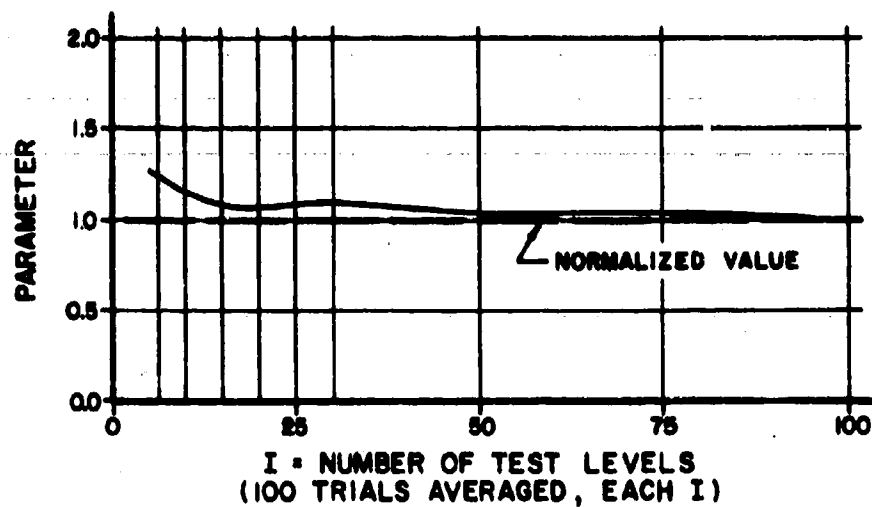
made for seven test levels optimally spaced between $0 < F < 1$, and likewise for 10, 15, 20, 30, 50, 75, and 100 test levels.

The distributions used were the exponential, the Pareto, and the limited types. These yield nonzero probability of failure only for positive values of the test level; hence, the magnitude of the test level is the independent variable. The corresponding density functions are all skewed in the direction of increasing test level. The exponential and Pareto types are unlimited in that the probability of failure approaches (but never quite reaches) 100 percent as the test level is increased. The limited distribution could be used if it were known that 100 percent failure was certain above a given level. The exponential and limited types have all statistical moments (such as mean, variance, skewness, kurtosis, and higher moments) defined, while those for a Pareto distribution would not exist above a certain order. In this respect, a Pareto distribution is a special case of a Cauchy distribution being nonzero only for positive values of test levels. These three distributions were used in this work because they are representative of numerous one-parameter distributions a statistician might select from for application to fragility data.

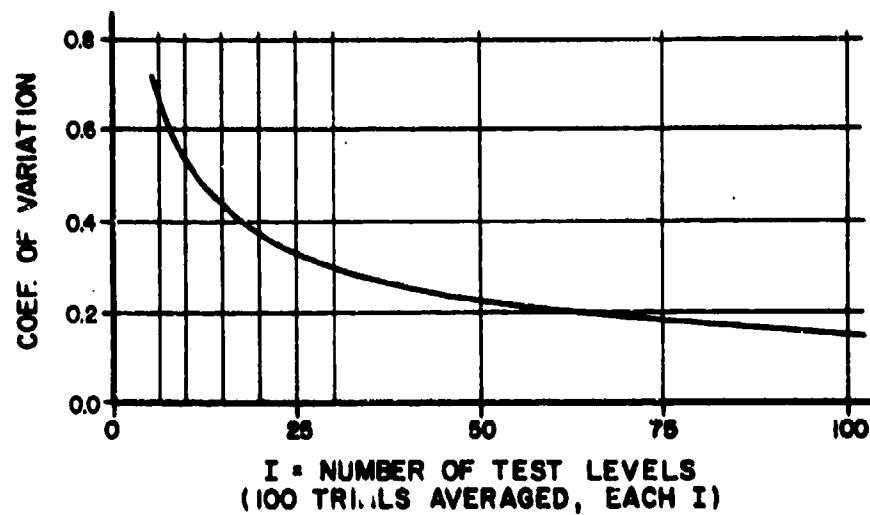
The value of the single parameter α to be estimated in each case was normalized to unity. For simplicity, the measure of accuracy of the estimate was taken as the *coefficient of variation*, which is the ratio of the standard deviation of the estimate to the value of the estimate. The standard deviation is a measure of absolute error of the estimate; hence, this ratio is descriptive of the relative error of the estimate. The results plotted in Figure 3 represent averages of results from the three distributions.

In Figure 3a, α is predicted with some positive bias when the number of test levels is less than approximately 20. The biased estimation of parameters often occurs in statistical problems, and can be accounted for by correcting the parameter by the indicated error. When the number of test levels exceeds 20, however, the average estimation of α approaches unity, as it should.

Figure 3b is a plot of the coefficient of variation of α , and shows the manner in which the estimate of α improves as the number of tests in the full probability range is increased. The interpretation of this accuracy is demonstrated best by an example. For 20 test levels, $c = .392$. This implies that there is 68.4 percent chance that the actual value of α will be within 39.2 percent of the estimated value. It can be seen that this accuracy is rather poor. Again, for 100 test levels, there is 68.4 percent chance that the actual value will be within 16.0 percent of the estimated value. Hence, the improvement of accuracy is demonstrated and shown as a function of the number of tests, optimally spaced between $0 < F < 1$.



a. Estimator prediction.



b. Accuracy prediction.

Figure 3. Average results for estimator and accuracy, independent failures.

The method of *maximum likelihood* was used to obtain the above results. Use of another estimation technique might possibly improve accuracy; however, it is highly unlikely that the improvement would be significant. The important conclusion is that, for independent failures, a large number of tests must be conducted on a unit before failure can be predicted with reasonable accuracy. In practice, it will be difficult to select test levels spaced in an optimum manner, and time and funds may prohibit testing when the probability of failure is high. Hence, the results shown in Figure 3 tend to be optimistic, and additional tests will generally be required to achieve the same accuracy in actual cases.

3 OBSERVED FAILURE INFORMATION

Summary of SAFEGUARD Failure Data

Two different procedures for hardness assurance testing were identified in the SAFEGUARD program.¹ In this program many off-the-shelf items of support equipment were tested which are similar to those used in essential systems of critical facilities. Even though the test envelopes for SAFEGUARD were not what would be required for critical facilities, the experience and qualitative results should be directly applicable.

The method used for most units involved testing at 25 percent, 50 percent, 75 percent, and 100 percent of the expected shock environment. The second method was used exclusively on five motor control centers, each of which was submitted to more than four test levels at and below the 100 percent level. The types of failures recorded in each group were significantly different.

General Equipment Tests

Fifty-eight independent units were tested under the first method; of these, 26 units passed all levels (45 percent), while two units failed all levels (3 percent).

A total of 968 tests were held on all units, often in orthogonal directions. Tests were held primarily at the four levels mentioned above. A few tests were held below the 25 percent level, and others were held as high as the 140 percent level. A total of 430 failures were recorded. Of the failures, 84 percent could be repaired immediately, or had a degrading effect which did not seriously impair the unit's function. Such failures have been defined previously as qualifying--the unit's function might be interrupted momentarily, but normal function can continue almost immediately. The remaining 16 percent of the failures produced lingering effects and required a significant amount of time to correct.

Table 1 lists the percentage of failures at each of the four primary test levels, together with the percentage of qualifying failures (Q), and the percentage of lingering failures (F). At the 100 percent level, for example, there were 411 tests yielding 225 failures (54.7 percent). Of the failures, 33.3 percent were qualifying, while the remaining 21.4 percent produced lingering effects.

¹*Subsystems Hardness Assurance Analysis*, HNDSP-73-161-ED-R (U.S. Army Corps of Engineers, Huntsville Division, December 1974).

Table 1
Failure Summary From General Equipment Tests

Full Test Level	No. of Tests	Combined Failures	Qualifying Failures (Q)	Lingering Failures (F)
25%	70	5.7%	5.7%	0.0%
50%	163	20.2%	14.1%	6.1%
75%	197	37.1%	31.5%	5.6%
100%	411	54.7%	33.3%	21.4%

Motor Control Center Tests

Five independent motor control centers were tested, with failures recorded for all units.

One hundred and seventy-two tests were held at numerous test levels in three orthogonal axes; an average of 34 tests (11 in each axis) was held for each unit. A total of 92 failures (53 percent) was recorded, 22 (13 percent) of which were qualifying, and 70 (40 percent) of which produced lingering effects. The further breakdown of percentages for qualifying and lingering failures does not appear to be meaningful, since correlation with test levels is necessary. In this case, the numerous and inconsistent variety of test levels rendered such a tabulation impractical.

Types of Failures

Failures have already been classified as *qualifying* or *lingering* according to ease of repair or time delay. It is desirable to report *all* failures for future reference. However, in current procedures for recording failures, the required time for repair is usually not indicated. Therefore, it is often difficult to review test reports for the purpose of identifying trivial or significant failures. This experience led to the proposed specification that the test engineer record his opinion about the amount of repair time required. The opinion can be reviewed and corrected by more qualified personnel, if necessary. Even knowledge that the time delay was unknown would be helpful.

Failures observed in the SAFEGUARD data could be classified further according to consistency or independence. Often, more than one consistent failure was observed to occur at a single test level. It

is roughly estimated that more than 90 percent of the failures which occurred during the general equipment tests were of this type. The estimate is rough because the recorded results were not oriented so that this type of failure could be clearly identified. When such a failure occurs, there is no doubt that it must be eliminated by hardening or isolating the unit to withstand the environment.

In contrast, the failures recorded from the motor control center tests were very inconsistent; i.e., the same failure might or might not occur more than once at the same level or at different levels. Usually, repairing the failure after one test would have no significant influence on whether or not the same failure would occur for any other test. In general, the higher the test level, the greater the probability of having one or more failures of this type. Statistically, such failures are independent. Estimating the probability of failure is more difficult in this case, requiring the use of conventional methods of probability and statistics. When independent failures occur, it is especially important to conduct a sufficient number of tests at preferred test levels to more accurately predict the probability of failure.

Appendix B provides a statistical method for estimating the probability of failure for independent failures and for calculating the accuracy of the estimation. The results of this analysis provide criteria for planning the number of tests and test levels when independent failures occur (see summary in Chapter 2).

Typical Failure Modes

Eventually, it will be necessary to generate specifications for designing, mounting, and procuring critical equipment; however, it is not feasible to consider providing specifications for all such equipment in the near future. A more reasonable approach would be to determine what failures have occurred most often during testing or from direct exposure to the hazardous environment. Priorities may then be established for attacking the various modes of equipment failure in order of importance.

The most complete listing of equipment failure modes encountered until now appears to be from SAFEGUARD test data. To complement this information, a survey of earthquake damage reports is scheduled in the near future to identify typical failure modes on the system level for systems considered to be essential for critical facilities. On the component level, failure modes observed in the SAFEGUARD data are:

1. Piping failures:
 - Joint leakage
 - Joint shear
 - Joint separation
 - Pipe burst

Braces bent
Brace bolt sheared
Valve failure (check)
Valve chatter

2. Indicator failures:

Pressure
Temperature
Liquid level
Flow rate

3. Sensing device failures:

Transducer shear-off
Wires cut
Inadvertent switch actuation

4. Machinery failures:

Pump cavitation
Pump leakage
Motor-pump coupling failures
Motor-generator coupling failures
Pump flow setting change
Pump seizure
Motor belt drive separation

5. Mounting failures:

Tank mounting failure
Pump mounting bolt shear
Motor mounts broken
Legs, brackets broken
Displacement interference

6. Electrical failures:

Switch contact chatter
Relay chatter
Relay trip
Circuit breaker trip
Lights broken.

4 SUMMARY AND CONCLUSIONS

Summary

Chapter 2 proposed overall testing and test report specifications that will contribute to a more complete listing of future test specifications. The definitions provided will facilitate the interpretation of failure data.

Two major classifications of failure data have been identified from the review of available fragility test results. First, a failure should be classified according to its significance in preventing a critical function. If a failure causes only a minor interruption of performance or a degrading effect either within the unit or with an interfacing unit, it should be classified as *qualifying*. If it causes an extended interruption, or causes other interfacing units to fail, it should be classified as *lingering*. Attempts to harden equipment should be directed primarily to failures exhibiting lingering effects. Second, a failure should be classified according to its predictability of occurrence. If the same failure occurs repeatedly at or above a given test level, it should be identified as *consistent*. If the failure occurs erratically such that it may or may not occur again at the same or a different test level, it should be identified as *independent*.

Guidelines provided in Chapter 2 will help a test engineer plan for an adequate number of tests at appropriate test levels. The results of analyzing independent failure data from hypothetical tests showed that a relatively large number of tests must be conducted to calculate the probability of failure reasonably accurately. Figure 3 showed how the number of tests was related to the accuracy of the prediction of failure.

Chapter 3 provided evidence of the existence of both consistent and independent failures. When failures occurred consistently, it was clear that the unit had to be hardened or isolated from the environment. On the other hand, the occurrence of independent failures appeared to lead to confusion, since such data were recorded for further subjective assessment of the probability of failure. Also, the importance of identifying a failure as qualifying or lingering was emphasized, since this distinction was not always clear in existing reports.

Typical failure modes of off-the-shelf equipment were listed to identify the most serious failure problems. This tabulation is probably sufficient to set priorities for establishing future design, mounting, and procurement specifications.

Since no basic theoretical reference could be found on fragility testing, Appendix A was written to compare fatigue, strength, and fragility test data in terms of fundamental meaning and treatment. It was considered important for this work to show that fatigue and strength data are formulated as probability density functions, while fragility

data are cumulative and must be formulated as a distribution function. It was significant that a fragility failure could not be identified uniquely with the test level, since the failure might have occurred at any other level below the test level.

Independent failures were common in certain units tested under the SAFEGUARD hardness assurance program. To interpret independent failure data, statistical methods must be used, as demonstrated in Appendix B. The analysis of Appendix B was provided for two purposes: (1) it shows the technical level required to interpret independent fragility failure data, and (2) it provides guidelines for a test engineer to anticipate testing requirements when independent failures are identified. The significant results were reflected in the proposed specifications of Chapter 2.

Conclusions

The hardness assurance of critical equipment involves a series of expensive operations beginning with fragility testing. Therefore, when testing is authorized, the results should be reported in considerable detail. The cost of such reports should be small compared to the cost of testing. Adequate hardness for a unit may never be achieved within the framework of available time and funds, since many manufacturers have only superficial capability of testing their products or interpreting the failure data.

Failure data should be clarified originally by the test engineer and reported, with help from a consultant expert if necessary. Failure data can be used effectively for overall hardness assurance only if the units to be tested are controlled and monitored, and the test data are interpreted in an unbiased manner. Detailed and uniform report information should be required as described herein, since it may be necessary for personnel unfamiliar with the test unit to interpret the failure data at a later time. When independent failures are encountered, the test engineer should recommend continued testing if there is adequate time and funds.

Review of available test data revealed that consistent failures are likely to occur in most mechanical equipment used for internal utility and lifeline systems in critical facilities. Independent failures are likely to occur in electrical control equipment. If a unit exhibits consistent failures, a fragility envelope can be determined. For independent failures, a precise definition of a fragility envelope is not possible, and the prediction of failure becomes a statistical problem.

Failure data analysis results have been recast in the form of proposed specifications. These specifications should be combined with future testing specifications from research that is not directly associated with the objectives of this particular task.

APPENDIX A:

OVERVIEW OF RELATED DESTRUCTIVE TESTING METHODS

No fundamental and readily available texts or other references could be found that address the statistical treatment of fragility data. Since most civil and mechanical engineers are at least superficially familiar with the treatment of fatigue and strength test data, a comparison with these methods is necessary, since there are important differences that may not be immediately obvious to practicing engineers.

All statistical references denote the result of a test as an *outcome*. Input parameters may be fixed or otherwise deterministic, or they may be random. However, the outcome is usually random, since it cannot be predicted deterministically.

In a fatigue test, a unit is loaded in a cyclic (or random cyclic) manner. The loading pattern and time history is controlled by the investigator and is part of the input data. The outcome is the number of cycles (or time) to failure. Note that failure itself is not the outcome, and the test will continue until failure occurs. The purpose of a series of such experiments is to determine the statistical properties of the outcome data, so that time to failure for nominally identical units can be predicted by probability within certain confidence limits. The outcome data can be arranged in the form of a probability density function, such as the bell-shaped curve of the normal density function. The condition is shown in Figure A1, where the peak of the bell occurs at the average time to failure for nominally identical specimens for a specific dynamic loading condition. If the loading condition is varied, the bell-shaped pattern of times to failure will adjust to a new position. If enough loading conditions are tried, a continuous line can be drawn connecting the peaks of the density curves, describing the average time to failure under any of the considered loading conditions.

A similar condition exists with strength testing in which the load on the specimen is increased until failure occurs. The outcome is not failure (which is expected to occur), but the stress or force level at which failure occurred. For nominally identical specimens, the outcome data can again be described by a probability density function in the form of a bell-shaped curve. The peak of the bell occurs at the average force at which failure occurred.

Fragility data is markedly different than either of the above types of data. In this case the test level is controlled and selected by the investigator, and forms part of the input data. There are two possible outcomes: survival and failure. A statistical test of this

type is called a Bernoulli experiment.² The test level is *not* an outcome, and a direct comparison of levels between a fragility test and a strength (or fatigue) test is not appropriate. The major difference is that when a failure occurs at a given test level in a fragility test, it may have occurred at any other test level less than or equal to that level. Hence the test level does not uniquely define the true flaw level at which failure would have occurred. This is a cumulative probability phenomenon and is depicted by a probability *distribution* function, rather than by a *density* function as for fatigue and strength testing data. (A distribution function is simply the integral of a density function.)

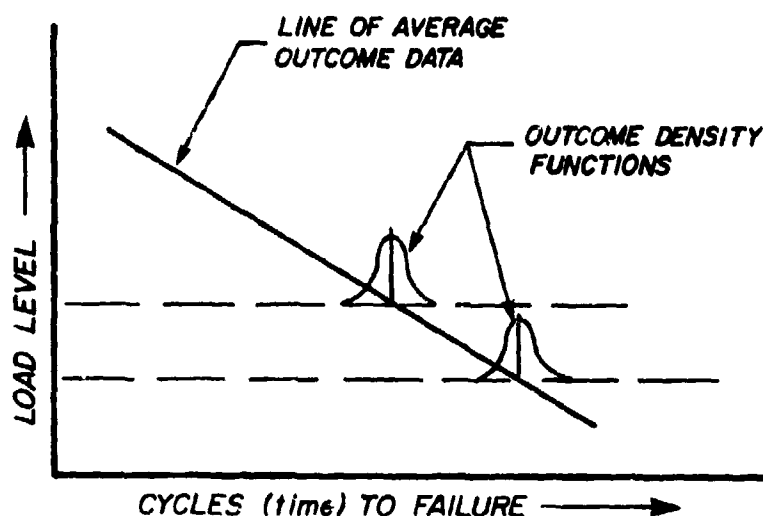
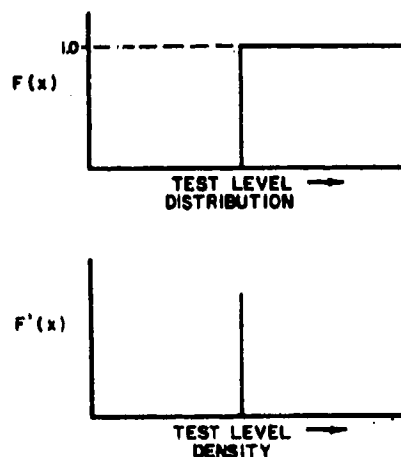


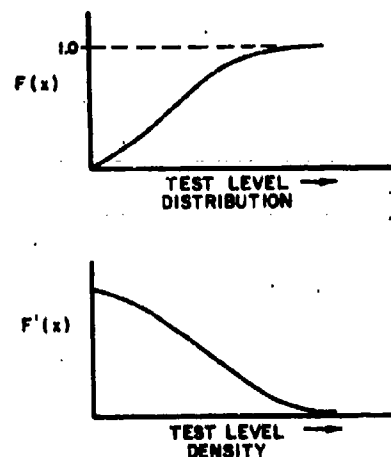
Figure A1. Fatigue data presentation.

When a consistent flaw exists in a fragility specimen, its probability of failure is given by a step-like distribution function, as shown in Figure A2a. That is, below a fairly precise test level, there will be zero probability of failure, and above this level, there will be 100 percent probability of failure. The density function for such a distribution may be depicted theoretically by a mathematical "delta" function, and is formed as the derivative of the step function. This delta function may be compared with the limit of the bell-shaped density function obtained by fatigue or strength testing; the span of failure levels in the delta function is very small.

²B. W. Lindgren, *Statistical Theory* (The MacMillan Company, 1962).



a. Consistent failures.



b. Independent failures.

Figure A2. Distribution and density functions for consistent and independent failures.

When independent flaws occur in a fragility specimen, there is no precise level at which failure will occur, as shown in Figure A2b. The higher the test level, the greater the probability that failure will occur. The derivative of such a distribution function usually appears as a skewed density function, as shown in the lower curve of the figure. If negative test levels are considered (i.e., test shocks in opposite directions), the density function will appear symmetric about the zero level. Hence, the density function derived from fragility data (when independent flaws exist) may appear as a skewed bell-shaped curve, and its peak value will usually be located close to the zero level.

It is expected that most units will exhibit both consistent and independent flaws to some extent. The distinction is often difficult to make unless the test engineer recognizes the type of flaw at the time of the test and records his observation for future reference.

APPENDIX B:

STATISTICAL ANALYSIS OF FRAGILITY DATA (INDEPENDENT FAILURES)

The Fundamental Fragility Problem

This appendix considers only problems occurring with independent failures. All references to an environmental test, flaw, or fragility level implies a shaped spectrum which can be defined by a single parameter, called the "level."

For the simplest fragility test, consider a unit which will either survive (S) or fail (F) if it is subjected to one specified test level.³ Assume that eight tests are held, with the following outcomes:

S, F, S, S, S, F, S, S.

If p denotes the probability of failure, then the probability of survival is $(1 - p)$. Note that failure has occurred twice at this level out of eight trials; hence,

$$p = \frac{2}{8} = 0.25 \quad [\text{Eq B1}]$$

Also note that the test level does not precisely define the level at which failure occurs, since the unit will often survive at this level with $(1 - p) = 0.75$. The only conclusion is that when failure did occur, the actual flaw level could have been any level less than or equal to the actual test level.

For more complicated problems, this basic problem is cast as a problem of *maximum likelihood*. A likelihood function, $L(p)$, is formed by the product of the p and $(1 - p)$ terms as the survivals and failures occurred in the above experiment:

$$\begin{aligned} L(p) &= (1 - p) p (1 - p) \dots (1 - p) \\ &= p^2 (1 - p)^6 \end{aligned} \quad [\text{Eq B2}]$$

Now p may be calculated by maximizing $L(p)$ with respect to p :

$$L'(p) = 0 = p(1 - p)^5 (2 - 8p) \quad [\text{Eq B3}]$$

where $L'(p)$ denotes partial differentiation of L with respect to p .

Note that $L'(p) = 0$ at $p = 0$ and at $p = 1$, and is positive for $0 < p < 1$. Hence, a maximum exists, and p can be calculated. The third term in parentheses on the right of Eq B3 yields

³B. W. Lindgren, *Statistical Theory* (The MacMillan Company, 1962).

$$p = \frac{2}{8} = 0.25 \quad [\text{Eq B4}]$$

as found previously. When p is calculated in this manner, it is called a maximum likelihood (ML) estimator of the true value of p , which is the desired quantity to be estimated from the experiment results. Likewise, $L(p)$ is referred to as the likelihood function (LF).

The fundamental fragility test can be generalized for one test level, where there are n tests and k failures:

$$L(p) = p^k(1 - p)^{n-k} \quad [\text{Eq B5}]$$

$$L'(p) = p^{k-1}(1 - p)^{n-k}(np - k) \quad [\text{Eq B6}]$$

$$p = \frac{k}{n} \quad [\text{Eq B7}]$$

The above relations follow immediately from Eqs B2, B3, and B4.

In the fragility tests performed on the motor control centers discussed in Chapter 3, numerous cases occurred in which Eqs B5, B6, and B7 could be applied at a single test level. However, more information was always available because each unit was subjected to more than one test level. The problem of calculating the probability of failure thereby becomes more complicated, but accuracy improves because more information can be used. The use of the ML estimation method is warranted when various test levels are used.

Extension and Generalization of the Fundamental Fragility Problem

When more than one test level appears in numerous trials of a fragility test, the probability of failure will be a function of test level. Generally more information about the governing probability distribution function for failure must be known or assumed.

The use of the LF to include various test levels is possible if p can be written as a continuous function of the test level, as

$$p = p(x) \quad [\text{Eq 8}]$$

where x is the test level.

Also, the LF given by Eq B5 as

$$L(p) = p^k(1 - p)^{n-k} \quad [\text{Eq B9}]$$

must be an admissible probability density function. The required properties for the LF to be an admissible density function are

$$L(p) \geq 0 \quad [\text{Eq B10}]$$

$$c \int_{-\infty}^{\infty} L(p) dp = 1, c > 0 \quad [\text{Eq B11}]$$

where c is a positive constant.

Inequality B10 was demonstrated to be true, for practical purposes, in the discussion following Eq B3. To prove that Eq B11 is also true, note that p is defined on $[0,1]$ and is zero elsewhere. Then, putting Eq B9 into Eq B11 gives

$$c \int_0^1 p^k (1-p)^{n-k} dp = \frac{\Gamma(k+1)\Gamma(n-k+1)}{\Gamma(n+2)} \quad [\text{Eq B12}]$$

where Γ is the conventional gamma function.

Hence, Eq B11 is satisfied if

$$c = \frac{\Gamma(n+2)}{\Gamma(k+1)\Gamma(n-k+1)} \quad [\text{Eq B13}]$$

It will be shown that c cancels out of subsequent calculations and can be any positive bounded constant. The value of c given by Eq B13 is therefore acceptable. With this value of c , the integral shown in Eq B11 or B12 depicts the conventional beta function, and the LF for a single-level fragility test has the shape of the beta density function.⁴ Again, this condition implies that $L(p)$ is zero at $p = 0$ and $p = 1$, and is positive for $0 < p < 1$. For a single test level, it has a single maximum value from which p can be calculated by the ML method.

The required extension of the LF for more test levels can be obtained as the product of the beta density functions (the LF's for single test levels).⁵

$$L(p) = \prod_{i=1}^I c_i L(p_i) = \prod_{i=1}^I c_i \prod_{j=1}^I L(p_j)$$

or

$$L(p) = c \prod_{i=1}^I p_i^{k_i} (1-p_i)^{n_i-k_i} \left\{ \begin{array}{l} 1 \leq n_i < \infty \\ 0 \leq k_i \leq n_i \end{array} \right\} \quad [\text{Eq B14}]$$

⁴G. H. Hahn and S. S. Shapiro, *Statistical Models in Engineering* (John Wiley and Sons, Inc., 1967).

⁵B. W. Lindgren, *Statistical Theory* (The MacMillan Company, 1962).

where C is the product of the c_i , and I is the total number of test levels.

In Eq B14, the p_i appear as discrete probabilities, but they are not independent, since p is a function of x . It is recalled that a failure will occur if a flaw level in the specimen is less than or equal to the test level. Since more than one flaw level may exist below a given test level, the probability of failure is

$$p_i = P(\text{at least one flaw level} \leq \text{test level } x_i) \quad [\text{Eq B15}]$$

This is a cumulative probability phenomenon, where x_i denotes the test level, and P is any admissible probability *distribution* function. Letting y denote a flaw level, x a test level, and $F(x)$ the governing distribution function for failure, Eq B15 may be written as

$$p_i = P(y \leq x_i) = F(x_i) = F_i \quad [\text{Eq B16}]$$

The probabilities p_i are still discrete but are now expressed as continuous functions of the test levels x_i . So far, the only restriction placed on $F(x)$ is that it satisfy the requirements for an admissible probability distribution function.

The problem of maximizing Eq B14 with respect to each p_i is greatly simplified by replacing p_i with $F(x_i)$, so that $F(x_i)$ is a function of one or more (possibly up to three) parameters α_j , $j = 1, 2, \dots, J$. The maximum number of parameters, J , is usually much less than the total number of tests, I . These parameters are now the unknowns, rather than the p_i , and the function $F(x_i)$ is assumed to be a known function $F(x_i; \alpha_j)$ of the test level and the parameters. The task is now to maximize the LF, Eq B14, with respect to the α_j instead of the p , which is much easier to accomplish.

Putting Eq B16 into Eq B14 gives

$$L(x; \bar{\alpha}) = C \prod_{i=1}^I F_i^{k_i}(x_i; \bar{\alpha}) \{1 - F_i(x_i; \bar{\alpha})\}^{n_i - k_i}$$

or in simplified form,

$$L(x; \bar{\alpha}) = C \prod_{i=1}^I F_i^{k_i} (1 - F_i)^{n_i - k_i} \quad [\text{Eq B17}]$$

where $\bar{\alpha}$ is a vector of the α_j , $j = 1, 2, \dots, J$.

In the usual treatment of a likelihood function, the logarithm of Eq B17 is taken to obtain

$$\ln L(x; \bar{\alpha}) = \ln C + \sum_{i=1}^I [k_i \ln F_i + (n_i - k_i) \ln(1 - F_i)] \quad [\text{Eq B18}]$$

It is proven⁶ that maximizing the logarithm of a likelihood function is equivalent to maximizing the function itself. Thus, the problem is further simplified by dealing with sums instead of products. Eq B18 is maximized with respect to each x_j as

$$\frac{\partial}{\partial \alpha_j} [\ln L(x; \bar{\alpha})] = 0 = \sum_{i=1}^I \left(\frac{k_i}{F_i} - \frac{n_i - k_i}{1 - F_i} \right) \frac{\partial F_i}{\partial \alpha_j}, \quad j=1, 2, \dots, J \quad [\text{Eq B19}]$$

Note that the constant C has been eliminated, as previously stated, and simplification of Eq B19 leads to J simultaneous equations of the form

$$0 = \sum_{i=1}^I \left[\frac{k_i - n_i F_i}{F_i(1 - F_i)} \right] \frac{\partial F_i}{\partial \alpha_j} \quad [\text{Eq B20}]$$

In practice, there appears to be no need for J to exceed unity, although it may be as high as three if hundreds of test results are available for a single unit. Independent failure data are not now readily available, and applying sophisticated probability functions (i.e., with as many as three parameters) does not appear to be warranted.

The solution of Eq B20 for $\bar{\alpha}$ may yield numerous roots,⁷ in which case the correct result is the set of $\bar{\alpha}$ which yields to the largest maximum. When this condition exists, there will also be minimum roots, since one minimum will exist between every pair of maximum points. This condition has not yet been observed.

It is noted in Eq B20 that $F(x; \bar{\alpha})$ must be differentiable with respect to each α_j . Also, although it is not necessary to differentiate $L(x; \bar{\alpha})$, a maximum will exist, provided that

$$\frac{\partial^2}{\partial \alpha_j^2} (\ln L) < 0 \quad [\text{Eq B21}]$$

For further discussion, the roots of Eq B20 will be denoted by $\hat{\alpha}$, where it is implied that $\hat{\alpha}$ is a vector if more than one parameter must be estimated.

⁶E. J. Gumbel, *Statistics of Extremes* (Columbia University Press, 1958).

⁷M. J. Kendall and A. Stuart, *The Advanced Theory of Statistics* (Hafner Publishing Company, 1969).

Accuracy of the Estimator

The solution of Eq B20 yields a value $\hat{\alpha}$ which *estimates* the true value of α . Intuitively, it is expected that $\hat{\alpha}$ will approach α as the number of trials increases. Also, it should be expected that accuracy will improve if the trial test levels are spaced such that increments of F are uniformly distributed in the range $0 < F < 1$.

The dependence of accuracy on the number of trials and test levels is quantified by the calculation of the variance of $\hat{\alpha}$, denoted as $\text{var}(\hat{\alpha})$. Both $\hat{\alpha}$ and $\text{var}(\hat{\alpha})$ can be calculated from a series of numerous tests at various levels.

The general fragility problem for independent failures has been shown to fall quite naturally into a problem of maximum likelihood. A minimum variance bound (MVB) can be calculated if the problem is cast in this manner,⁸ but there may be many other techniques in addition to the ML method to estimate α . However, the highest attainable accuracy is reflected by the MVB, which is determined by the ML method. This does not imply that the MVB is the actual $\text{var}(\hat{\alpha})$ found, even if the ML method is used. Instead, the MVB is a variance which is less than or equal to the actual variance calculated by *any* estimation technique, including the ML method. An estimator which attains such a lower bound for $\text{var}(\hat{\alpha})$ is called an MVB estimator. If the condition for consistency and sufficiency exists (see *The Advanced Theory of Statistics*⁹) and if there is an unbiased MVB estimator, the estimator is given by the ML method.

The MVB, $\bar{v}$, is generally found by taking the negative reciprocal of the expected value of the second derivative of the log-likelihood function, written as

$$\bar{v}(\hat{\alpha}) = 1/R^2(\hat{\alpha}) \quad [\text{Eq B22}]$$

where

$$R^2(\hat{\alpha}) = -E\left(\frac{\partial^2}{\partial \alpha^2} \ln L\right). \quad [\text{Eq B23}]$$

The determination of the expected value, shown in Eq B23, is a standard statistical operation, and can be found in any fundamental statistical reference. In the particular case when L is composed of products of a density function which admits a single sufficient statistic for the parameter α , Eq B23 simplifies to

$$R^2(\hat{\alpha}) = -\left(\frac{\partial^2}{\partial \alpha^2} \ln L\right)|_{\hat{\alpha}} \quad [\text{Eq B24}]$$

⁸M. G. Kendall and A. Stuart, *The Advanced Theory of Statistics* (Hafner Publishing Company, 1969).

⁹M. G. Kendall and A. Stuart.

The density function in this case is from the beta distribution, as given by Eq B12 or Eq B14. Evidence that sufficiency is satisfied is provided in *The Advanced Theory of Statistics*¹⁰ for the binominal distribution, which can be written in the form of a beta distribution.¹¹ Although no further proof of sufficiency is provided here, Eq B24 has been used to calculate $\bar{v}$ in Eq B22. If the analysis of independent failures from fragility testing becomes a common operation in the future, a more rigorous proof of sufficiency should be provided, and Eq B23 used instead of Eq B24 if and when necessary. Also, it would be informative to calculate the actual variance, $\text{var}(\hat{\alpha})$, in addition to the MVB, $\bar{v}(\hat{\alpha})$.

The first derivative of the logarithm of the LF is shown in Eq B19 or Eq B21 as

$$\frac{\partial}{\partial \alpha} \ln L = \sum_{i=1}^I \left\{ \left[\frac{k - nF}{F(1 - F)} \right] \frac{\partial F}{\partial \alpha} \right\}_i \quad [\text{Eq B25}]$$

where each term except α in the outside bracket expression is understood to be subscripted with i . The second derivative is obtained for one parameter and simplified as

$$\frac{\partial^2}{\partial \alpha^2} \ln L = \sum_{i=1}^I \left\{ k \left[\frac{(\frac{\partial F}{\partial \alpha})^2 (1 - 2F)}{F^2 (1 - F)^2} - \frac{(\partial^2 F)/(\partial \alpha^2)}{F(1 - F)} \right] + \frac{nF}{(1 - F)^2} \right\}_i \quad [\text{Eq B26}]$$

However, Eq B26 is valid only if there is one parameter, α , to estimate. For more than one parameter in the expression for F , a matrix of equations like Eq B26 is formed, such as

$$\frac{\partial^2}{\partial \alpha_{\ell} \partial \alpha_m} [\ln L(x; \alpha)] = \sum_{i=1}^I \left\{ k \left[\frac{F_{\ell} F_m (1 - 2F)}{F^2 (1 - F)^2} - \frac{F_{\ell, m}}{(1 - F)} \right] + \frac{nF}{(1 - F)^2} \right\}_i \quad [\text{Eq B27}]$$

where $\ell, m = 1, 2, \dots, J$,

$$F_{\ell} = \frac{\partial F}{\partial \alpha_{\ell}}$$

$$F_m = \frac{\partial F}{\partial \alpha_m}$$

$$F_{\ell, m} = \frac{\partial^2 F}{\partial \alpha_{\ell} \partial \alpha_m}$$

¹⁰M. G. Kendall and A. Stuart, *The Advanced Theory of Statistics* (Hafner Publishing Company, 1969).

¹¹G. J. Hahn and S. S. Shapiro, *Statistical Models in Engineering* (John Wiley and Sons, Inc., 1967).

In this case, a matrix of $\bar{v}$ terms is formed from the negative inverse of the matrix Eq B27:

$$[\bar{v}] = - \left[\frac{\partial^2}{\partial \alpha_k \partial \alpha_m} \ln L \right]^{-1} \quad [\text{Eq B28}]$$

where the brackets denote $J \times J$ square matrices.

The terms on the main diagonal of $[\bar{v}]$ are the MVB's for the associated $\hat{\alpha}_j$. The square roots of these terms will be the minimum standard deviation bound (MSDB) for each $\hat{\alpha}_j$. Letting $\bar{\sigma}_j$, $j = 1, J$, denote the MSDB's, the *coefficient of variance*, may be formed for each $\hat{\alpha}_j$ as¹²

$$c_j = \frac{\bar{\sigma}_j}{\hat{\alpha}_j}, \quad j = 1, 2, \dots, J \quad [\text{Eq B29}]$$

This coefficient provides a relative measure of accuracy. For similar tests of identical test levels, it implies that 58.4 percent of the time the estimator $\hat{\alpha}_j$ will be within c_j (converted to percent) of the calculated value of $\hat{\alpha}_j$. Also 95.4 percent of the time $\hat{\alpha}_j$ will be within $2c_j$ of the calculated value, and 99.7 percent of the time it will be within $3c_j$ of that value.

Selection of a Distribution Function

For a general fragility problem, an assumption must be made about the form of the distribution function for failure, $F(x; \alpha)$. Because of the lack of large quantities of independent failure data, the restrictions on F should be minimized. First, F must be an admissible probability distribution function, satisfying standard requirements in accordance with fundamental statistical theory.

It might be expected that the density function for F (its derivative with respect to x) should be limited at $x = 0$ and possibly skewed in the direction of increasing magnitude of x . This implies that $F(0; \alpha) = 0$ and that F will monotonically increase to unity as test level increases. It does not presently appear that the density function must decrease monotonically with x , although this assumption may be convenient, simple, and warranted for lack of sufficient failure data.

The use of distribution functions which must be written in integral form should be avoided, at least temporarily. For example, the normal, gamma, and beta distributions are written in integral form. (The beta

¹²W. S. Peters and G. W. Summers, *Statistical Analysis for Business Decisions* (Prentice-Hall, Inc., 1968).

function for F should not be confused with the beta function for the distribution of the LF. It is conceivable that F could also be represented as a beta function.) This restriction should be imposed because it may take thousands of times longer to obtain a computer solution for the parameter because of the requirement to integrate many times to obtain F in the iterative solution process.

Typical admissible distribution functions are those presently being used (by coincidence) in extremal analysis.¹³ The simplest and most appropriate form for early studies is recommended as the exponential distribution:

$$F(x;\alpha) = 1 - e^{-\alpha x} \quad [\text{Eq B30}]$$

This is a one-parameter function which has a monotonically decreasing density function for $x \geq 0$ and $\alpha > 0$. It is bounded at $x = 0$, but is unlimited for large values of x . All statistical moments exist. The Pareto distribution is given as

$$F(x;\alpha) = 1 - (x - \beta)^{-\alpha} \quad [\text{Eq B31}]$$

where x may be multiplied by a positive constant if desired.

Here, $\beta \geq 0$, and the lower bound for x is shown by $x \geq 1 + \beta$, while $\alpha \geq 1$. The distribution is unlimited for large values of x , but statistical moments greater than or equal to α do not exist. An investigator may have reason to believe that such a function may be applicable to a particular fragility problem. In comparison to the exponential density function, the derivative of Eq B31 has a peak value (mode) at some value $x > 0$, and therefore does not decay monotonically with increasing x . If x is multiplied by a constant, Eq B31 may be regarded as a three-parameter distribution, but it may be reduced to a two- or one-parameter function by letting the constant equal unity and by setting $\beta = 0$. For cases in which the investigator feels that bounds on x are necessary in both directions, the limited distribution may be used:

$$F(x;\alpha) = 1 - (\beta - x)^{\alpha} \quad [\text{Eq B32}]$$

Again, x may be multiplied by a positive constant to form a three-parameter distribution if desired. As Eq B32 is written, the bounds on x are $\beta - 1 < x \leq \beta$, with $\beta > 0$ and $\alpha > 0$. If the bounds are known (i.e., if β is known), and the constant multiplier for x is assumed as unity, Eq B32 reduces to a one-parameter distribution. Again, the mode of the density function occurs at $x > 0$.

There are more general asymptotic distributions associated with each of the above distributions. For example, the first asymptotic

¹³E. J. Gumbel, *Statistics of Extremes* (Columbia University Press, 1958).

distribution is associated with the exponential type. It is a two-parameter function given as

$$F(x;\bar{\alpha}) = \exp [-e^{-\alpha(x - B)}] \quad [\text{Eq B33}]$$

Restrictions or choices for α and β are discussed in detail in *Statistics of Extremes*.¹⁴ The advantage of using such an asymptotic function is that it will generally reflect the behavior of any distribution function of a class defined as the exponential type for sufficiently large values of x . Hence, the investigator need only assume the *type* of distribution function if he is willing to test at relatively high levels where failure is likely to occur. Two other related asymptotic distributions are available for more general Pareto and limited distributions.

For this work, only those functions presented above have been programmed and analyzed.

Solution

The solution of the fragility problem for the independent failures consists of estimating the parameters $\bar{\alpha}$ for the assumed distribution function $F(x;\bar{\alpha})$. If the parameters are known, the probability of failure can be calculated easily for any test level, x . An adequate computer program may be almost completely established from generally accepted and published subroutines.^{15,16} The following steps are necessary to obtain a solution:

1. Assume a suitable distribution function, such as one suggested in Eqs B29 - B32. The exponential distribution is recommended until more data become available to warrant using other distributions.
2. If J parameters must be estimated, set up J simultaneous equations, as given by Eq B20. There will be n tests and k failures at each test level x_i obtained from the experimental test data. Here, $0 \leq k \leq n$ and $n \geq 1$ at each test level.
3. Use a conventional iteration method, such as the Newton-Rapson technique, to calculate the estimators $\hat{\alpha}_j$, $j = 1, 2, \dots, J$.
4. Calculate the $J \times J$ matrix $[\bar{V}]$ (the MVB's) given by Eq B28. This matrix should be reasonably well-conditioned, implying that the

¹⁴E. J. Gumbel, *Statistics of Extremes* (Columbia University Press, 1958).

¹⁵System/360 Scientific Subroutine Package (360A-7M, 03X) Version III, Programmer's Manual (IBM, 1968).

¹⁶B. Carnahan, H. A. Luther, and J. O. Wilkes, *Applied Numerical Methods* (John Wiley and Sons, Inc., 1969).

terms on the main diagonal should be generally larger than the off-diagonal terms. If this situation is not observed, it is evident that the selection of $F(x;\bar{\alpha})$ was poor.

5. Take the square roots of the diagonal terms of $[V]$ from Eq B28 to obtain the MSDB's $\bar{\sigma}_j$, $j = 1, J$. Then calculate the coefficients of variation c_j given by Eq B29.

As suggested previously, it may be desirable in the future to calculate the actual covariance matrix as an additional step. Presently, however, completion of the above five steps constitutes the solution of the independent failure fragility problem.

Hypothetical Failure Studies

Many computer runs were made to analyze hypothetical failure data by the procedure outlined above to gain information about the relation between the number of tests and the accuracy of the estimators.

The hypothetical data were generated as follows. The total number of trials, I , for a test series was selected. Then $I + 1$ equal increments were taken between 0.0 and 1.0 to establish ideal increments in the function $F(x;\bar{\alpha})$. Hence, F_i , $i = j, 2, \dots, I$, was established and distributed with equal increments over $0.0 < F < 1.0$. Next, the desired algebraic form for $F(x;\bar{\alpha})$ was chosen from any of Eqs B30-B33, and values of α_j were assumed. For each value of F_i , the corresponding value of the test level x_i was calculated as shown in Eq B2 for the exponential distribution. This provided I ideal test levels to assure the best accuracy in calculating the parameters. Only one test was run for the fictitious unit at each level: a random number, selected from numbers uniformly distributed between 0.0 and 1.0, was selected by computer; if this number was less than or equal to F_i , the test outcome was considered a failure; if the number was greater than F_i , the outcome was survival.

The above technique provided I outcomes for each hypothetical test series for a fictitious unit. As the computer programming was established, the value of I and the algebraic expression for $F(x;\bar{\alpha})$ could be varied at will. The input data to the solution program consisted of the test levels x_i , the function $F(x;\bar{\alpha})$, and $k_i = 0$ or 1 , depending on whether survival or failure had occurred. The value of $n_i = 1$ was assumed in these hypothetical studies, since only one test was held at each level. However, this was not a general restriction, since real data could be used in the solution program for all acceptable values of k_i and n_i at each x_i .

For consistency in analyzing results, the number of trials for each of the four distributions was taken as $I = 7, 10, 15, 20, 30, 50, 75$, and 100 . For each of these values of I , 100 distinct hypothetical

fragility problems were solved. Solution of each problem consisted of calculating the $\bar{\alpha}$ vector to estimate the known values of α_j which had been assumed to generate the failure data. The MVB matrix, $[V]$, was also calculated. After 100 problems had been run for each value of I , the values of each $\hat{\alpha}_j$ were averaged. It is not an accepted statistical practice to average the coefficients of variance directly, since they are calculated by using the square roots of the diagonal terms of $[V]$. Hence, each diagonal term of $[V]$ was averaged first; then the square roots of these averages were taken to obtain the average $\bar{\sigma}_j$ values. Finally, the average $\bar{\sigma}_j$'s were divided by the average $\bar{\alpha}_j$'s to obtain the average coefficients of variance, c_j , in accordance with Eq B29.

These results are shown in Figures B1-B4 as plots of average $\hat{\alpha}_j$ vs. I , and c_j vs. I for each parameter in each of the four distributions. All four distributions in Figures B1-B3 are shown as reduced to one parameter ($J = 1$) distributions. Convergence of solutions was more difficult for low values of I , but was always obtained for the one-parameter distributions. Usually, convergence for two-parameter distributions was very difficult for $I < 10$ and for $I < 30$. Convergence for three-parameter distributions has never been obtained for values of $I < 100$; hence, no results for the three-parameter distributions were obtained. Figure B4 shows the results for a two-parameter distribution (the first asymptotic). Note that all assumed values of α_j are shown normalized to unity for simplicity of presentation. (The results for the one-parameter distributions were averaged and are presented in Chapter 4.)

It should be noted that some positive bias usually results for the estimators when $I < 20$. However, the $\hat{\alpha}_j$ are usually close to unity, as required, for $I > 20$. The bias can be corrected for low values of I directly from the plots. The consistent positive error implies that the ML method does not yield precisely unbiased estimators in this application. This conclusion would be substantiated only if the programming used was found to be faultless in this respect.

Motor Control Center Results

The ML solution technique was applied to all motor control center data discussed in Chapter 3. (It will be recalled that no consistent failures were identified in the five units tested.) Table B1 provides the results of this technique where the probability-of-failure law was assumed to follow the exponential distribution. Other distributions were tried, but no significantly different results were obtained to warrant additional tabulation. Only those test series have been tabulated for which at least one failure occurred, since the parameter could not be calculated if no failure occurred.

The poorest accuracy (i.e., the poorest coefficients of variation) occurred in Table B1 when only one failure was recorded. Although

FUNCTION: $F(x) = 1 - e^{-\alpha x}$

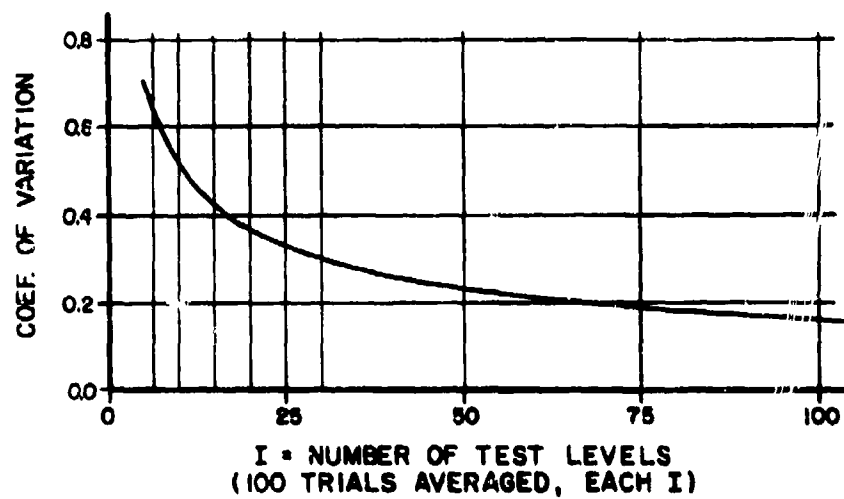
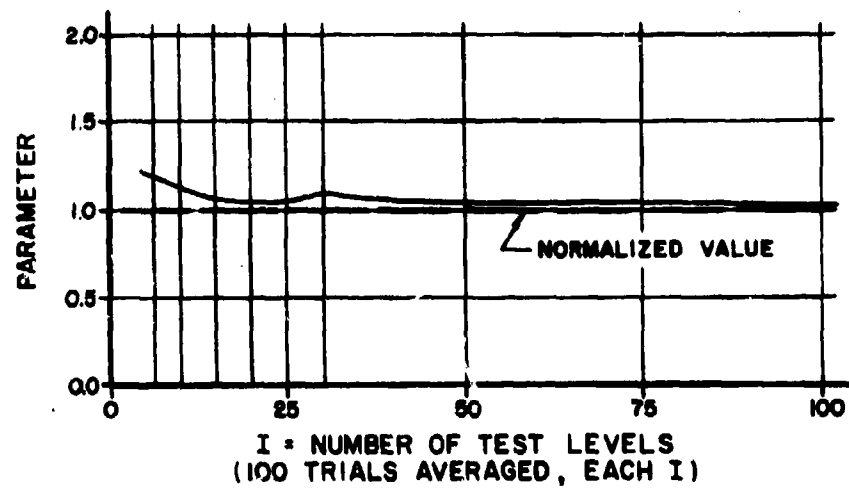


Figure B1. Parameter and accuracy for exponential distribution.

$$\text{FUNCTION: } F(x) = 1 - (x-1)^{-\alpha}$$

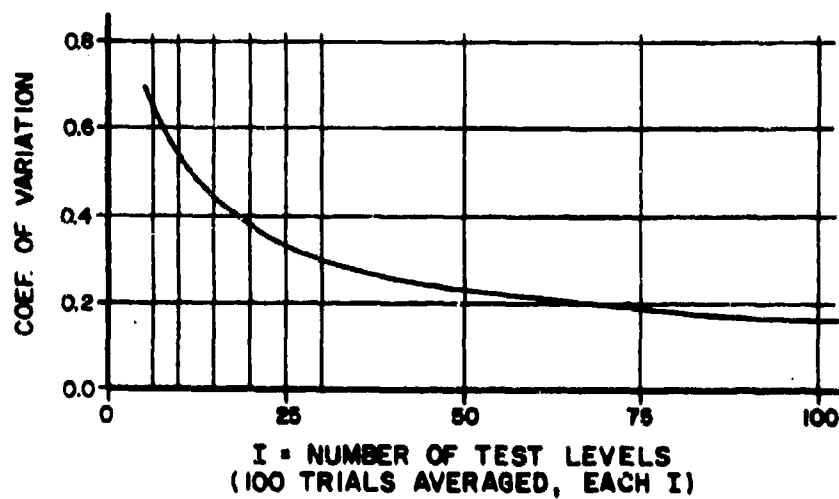
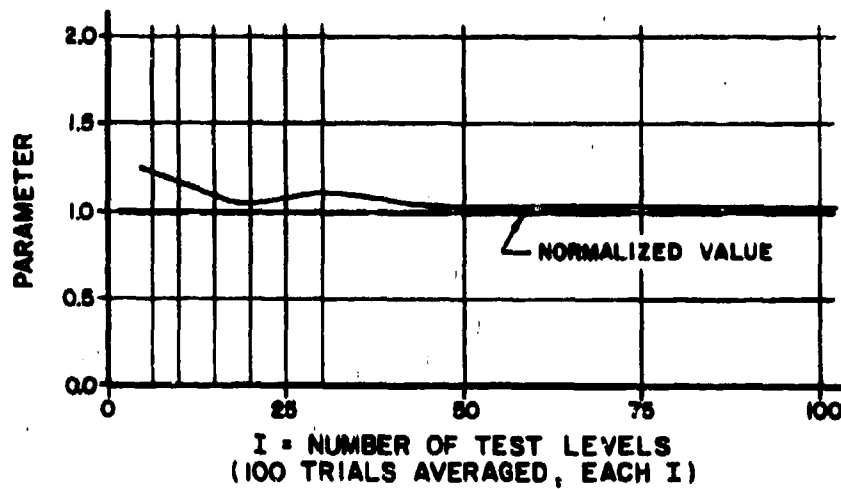


Figure B2. Parameter and accuracy for the one-parameter Pareto distribution.

FUNCTION: $F(x) = 1 - (1-x)^{\alpha}$

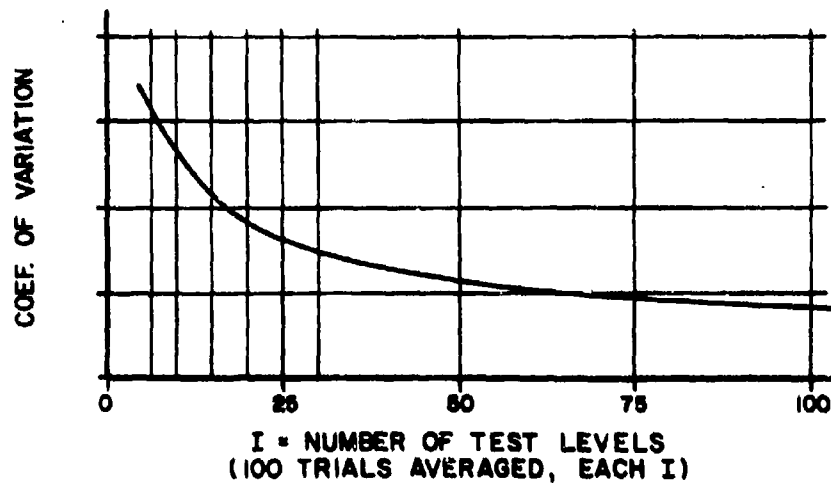
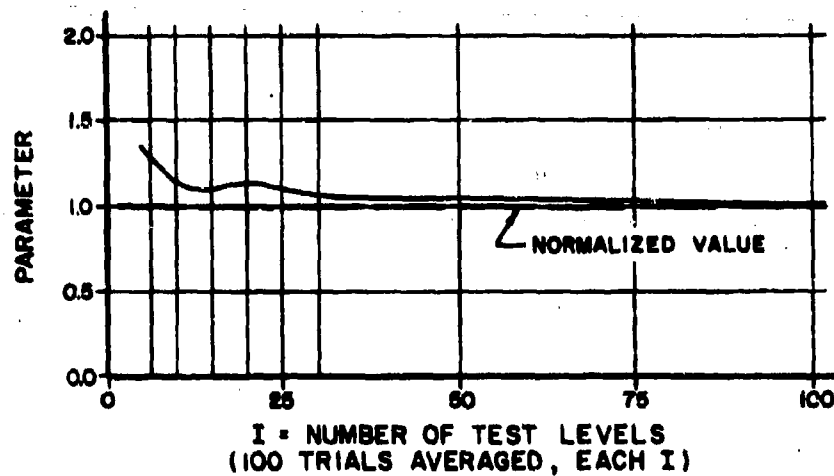


Figure B3. Parameter and accuracy for the one-parameter limited distribution.

$$\text{FUNCTION: } F(x) = \exp[-e^{-\alpha(x-\beta)}]$$

α - PARAMETER

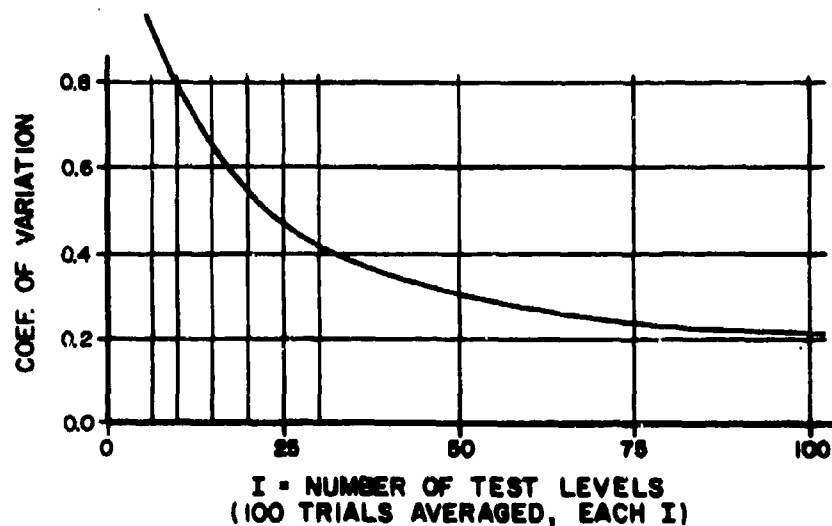
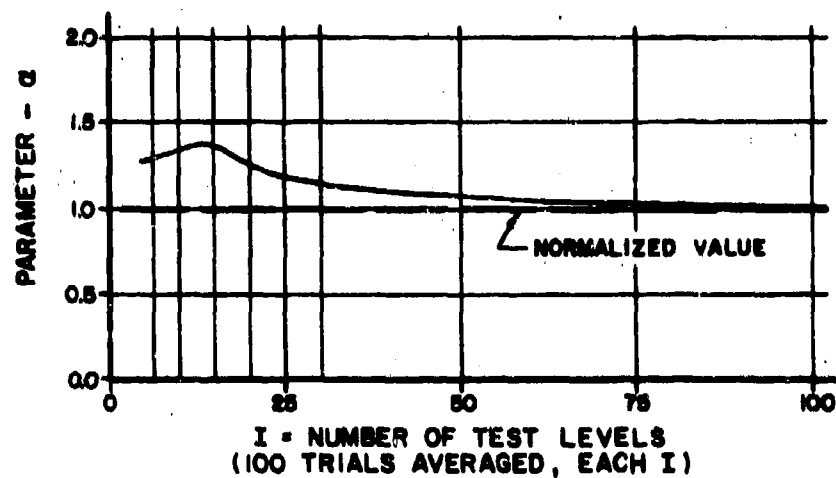


Figure B4. α -Parameter and accuracy for the two-parameter first asymptotic distribution.

$$\text{FUNCTION: } F(x) = \exp \left[-e^{-\alpha(x-\beta)} \right]$$

β - PARAMETER

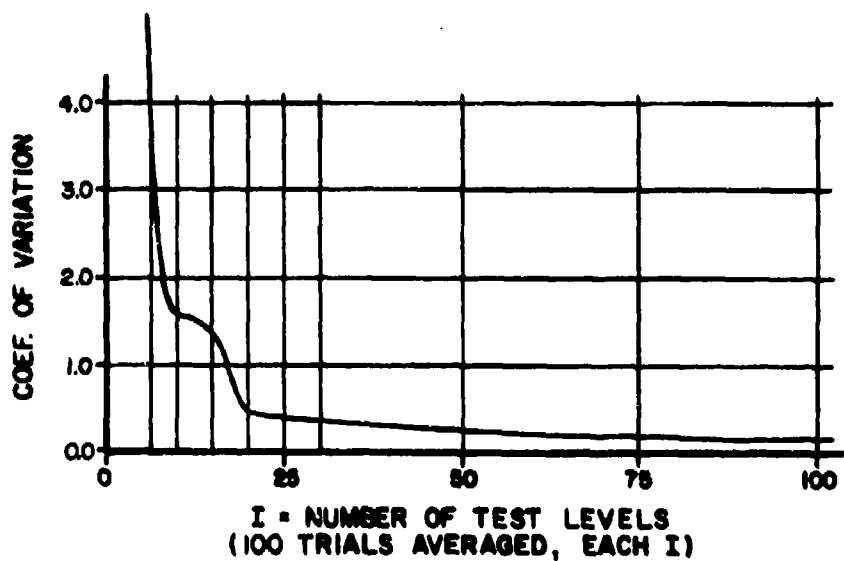
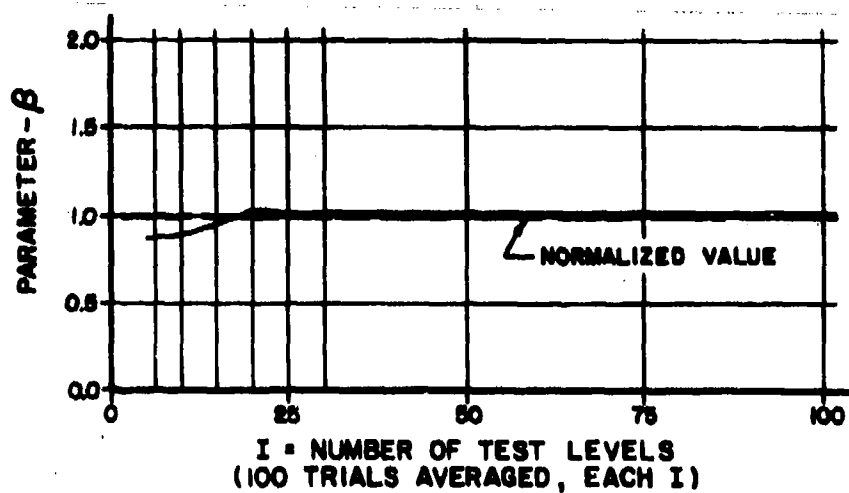


Figure B4 (cont'd).

Table B1

**Analysis of Motor Control Center Data Using
Exponential Distribution**

Probability Function: $F(x) = 1 - e^{-\alpha x}$

Unit	Axis	No. of Levels	Total Tests	Total Failures	α	c
E52MC	Z	14	27	8	0.077	0.371
	X	16	26	10	0.161	0.359
	Y	14	26	10	0.104	0.328
E089MC	X	8	13	1	0.028	1.004
	Y	10	15	1	0.020	1.002
E87MC	Z	15	21	4	0.064	0.511
	X	12	18	7	0.101	0.388
	Y	16	18	5	0.064	0.451
E06MC1	Z	13	15	2	0.052	0.711
E06MC2	Z	8	11	1	0.032	1.000
	X	8	11	1	0.032	1.000
	Y	8	14	3	0.078	0.657

there is no direct relationship, accuracy tends to improve as the percentage of failures of total tests increases. The lack of a more obvious direct relationship is probably the result of a clustered and generally poor selection of test levels.

REFERENCES

- Carnahan, B., H. A. Luther, and J. O. Wilkes, *Applied Numerical Methods* (John Wiley and Sons, Inc., 1969).
- Gumbel, E. J., *Statistics of Extremes* (Columbia University Press, 1958).
- Hahn, G. J. and S. S. Shapiro, *Statistical Models in Engineering* (John Wiley and Sons, Inc., 1967).
- Kendall, M. G. and A. Stuart, *The Advanced Theory of Statistics* (Hafner Publishing Company, 1969).
- Lindgreen, B. W., *Statistical Theory* (The MacMillan Company, 1962).
- Peters, W. S. and G. W. Summers, *Statistical Analysis for Business Decisions* (Prentice-Hall, Inc., 1968).
- Subsystems Hardness Assurance Analysis*, HNDSP-73-161-ED-R (U.S. Army Corps of Engineers, Huntsville Division, December 1974).
- System/360 Scientific Subroutine Package (360-CM 03X) Version III, Programmer's Manual* (IBM, 1968).

CERL DISTRIBUTION

Chief of Engineers
ATTN: DAEN-ASI-L (2)
ATTN: DAEN-MCZ-S (2)
ATTN: DAEN-RD
ATTN: DAEN-ZCI
Dept of the Army
WASH DC 20314

Defense Documentation Center
ATTN: TCA (12)
Cameron Station
Alexandria, VA 22314