

U.S. DEPARTMENT OF COMMERCE
National Technical Information Service

AD-A024 144

AN APPLICATION OF BAYESIAN ANALYSIS IN
DETERMINING APPROPRIATE SAMPLE SIZES FOR
USE IN U. S. ARMY OPERATIONAL TESTS

GEORGIA INSTITUTE OF TECHNOLOGY

PREPARED FOR
ARMY OPERATIONAL TEST AND EVALUATION AGENCY

JUNE 1975

KEEP UP TO DATE

Between the time you ordered this report—which is only one of the hundreds of thousands in the NTIS information collection available to you—and the time you are reading this message, several *new* reports relevant to your interests probably have entered the collection.

Subscribe to the **Weekly Government Abstracts** series that will bring you summaries of new reports as soon as they are received by NTIS from the originators of the research. The WGA's are an NTIS weekly newsletter service covering the most recent research findings in 25 areas of industrial, technological, and sociological interest—invaluable information for executives and professionals who must keep up to date.

The executive and professional information service provided by NTIS in the **Weekly Government Abstracts** newsletters will give you thorough and comprehensive coverage of government-conducted or sponsored re-

search activities. And you'll get this important information within two weeks of the time it's released by originating agencies.

WGA newsletters are computer produced and electronically photocomposed to slash the time gap between the release of a report and its availability. You can learn about technical innovations immediately—and use them in the most meaningful and productive ways possible for your organization. Please request NTIS-PR-205/PCW for more information.

The weekly newsletter series will keep you current. But *learn what you have missed in the past* by ordering a computer **NTISearch** of all the research reports in your area of interest, dating as far back as 1964, if you wish. Please request NTIS-PR-186/PCN for more information.

WRITE: Managing Editor
5285 Port Royal Road
Springfield, VA 22161

Keep Up To Date With SRIM

SRIM (Selected Research in Microfiche) provides you with regular, automatic distribution of the complete texts of NTIS research reports *only* in the subject areas you select. SRIM covers almost all Government research reports by subject area and/or the originating Federal or local government agency. You may subscribe by any category or subcategory of our WGA (**Weekly Government Abstracts**) or **Government Reports Announcements and Index** categories, or to the reports issued by a particular agency such as the Department of Defense, Federal Energy Administration, or Environmental Protection Agency. Other options that will give you greater selectivity are available on request.

The cost of SRIM service is only 45¢ domestic (60¢ foreign) for each complete

microfiched report. Your SRIM service begins as soon as your order is received and processed and you will receive biweekly shipments thereafter. If you wish, your service will be backdated to furnish you microfiche of reports issued earlier.

Because of contractual arrangements with several Special Technology Groups, not all NTIS reports are distributed in the SRIM program. You will receive a notice in your microfiche shipments identifying the exceptionally priced reports not available through SRIM.

A deposit account with NTIS is required before this service can be initiated. If you have specific questions concerning this service, please call (703) 451-1558, or write NTIS, attention SRIM Product Manager.

This information product distributed by

NTIS

U.S. DEPARTMENT OF COMMERCE
National Technical Information Service
5285 Port Royal Road
Springfield, Virginia 22161



AD A024144

AN APPLICATION OF BAYESIAN ANALYSIS IN DETERMINING
APPROPRIATE SAMPLE SIZES FOR USE IN U. S.
ARMY OPERATIONAL TESTS

A THESIS

Presented to

The Faculty of the Division of Graduate Studies

By

Robert L. Cordova

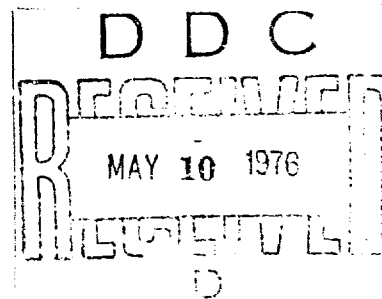
PRICES SUBJECT TO CHANGE

In Partial Fulfillment
of the Requirements for the Degree
Master of Science in Industrial Engineering

REPRODUCED BY
NATIONAL TECHNICAL
INFORMATION SERVICE
U. S. DEPARTMENT OF COMMERCE
SPRINGFIELD, VA. 22161

Georgia Institute of Technology

June, 1975



DISTRIBUTION STATEMENT A

Approved for public release;
Distribution Unlimited

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER	2. GOVT ACCESSION NO.	3. RECIPIENT'S CATALOG NUMBER
4. TITLE (and Subtitle) AN APPLICATION OF BAYESIAN ANALYSIS IN DETER- MINING APPROPRIATE SAMPLE SIZES FOR USE IN U.S. ARMY OPERATIONAL TESTS		5. TYPE OF REPORT & PERIOD COVERED Thesis
7. AUTHOR(s) Robert L. Cordova		6. PERFORMING ORG. REPORT NUMBER
9. PERFORMING ORGANIZATION NAME AND ADDRESS The School of Industrial & Systems Engineering Georgia Institute of Technology Atlanta, Georgia 30332		8. CONTRACT OR GRANT NUMBER(s) DAAG39-75-C-0097
11. CONTROLLING OFFICE NAME AND ADDRESS USA Operational Test & Evaluation Agency ATTN: CSTE-TDD; 5600 Columbia Pike Falls Church, VA 22041		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		12. REPORT DATE June 1975
		13. NUMBER OF PAGES 76
		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		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number)		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number)		

AN APPLICATION OF BAYESIAN ANALYSIS IN DETERMINING
APPROPRIATE SAMPLE SIZES FOR USE IN U. S.
ARMY OPERATIONAL TESTS

A THESIS

Presented to

The Faculty of the Division of Graduate Studies

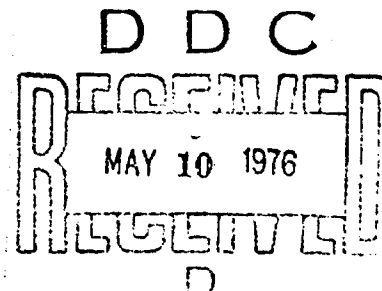
By

Robert L. Cordova

In Partial Fulfillment
of the Requirements for the Degree
Master of Science in Industrial Engineering

Georgia Institute of Technology

June, 1975



DISTRIBUTION STATEMENT A

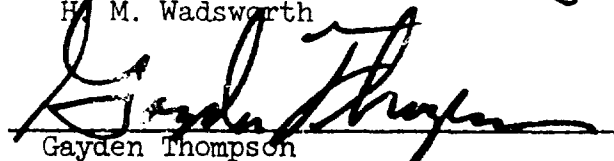
Approved for public release;
Distribution Unlimited

AN APPLICATION OF BAYESIAN ANALYSIS IN DETERMINING
APPROPRIATE SAMPLE SIZES FOR USE IN U. S.
ARMY OPERATIONAL TESTS

Approved:


Russell G. Heikes, Chairman


H. M. Wadsworth


Gayden Thompson

Date approved by Chairman:

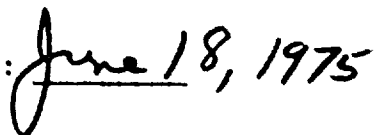


TABLE OF CONTENTS

	Page
LIST OF TABLES	iii
LIST OF FIGURES	iv
SUMMARY	v
Chapter	
I. INTRODUCTION	1
The General Problem	
The Specific Problem	
Background	
Review of the Literature	
II. THE TEST METHODOLOGY	10
The Assumptions of Normality	
The Prior Information	
The Basic Alternatives of Determining Sample Size	
The Classical Method	
A Bayesian Approximation	
III. DEMONSTRATION OF THE METHODOLOGY	34
Programming the Model	
Demonstrating the Model	
IV. CONCLUSIONS AND RECOMMENDATIONS	47
Conclusions	
Recommendations	
Appendices	
I. FORTRAN PROGRAM FOR THE APPROXIMATE BAYESIAN PROCEDURE	50
II. FORTRAN PROGRAM FOR THE CHI-SQUARE TEST OF NORMALITY	61
BIBLIOGRAPHY	64

LIST OF TABLES

Table	Page
1. Minimum Sample Size - Classical Method	17
2. Test of Normal Random Generator	35
3. Data for the Bayesian Approximation Model Based on One Hundred Runs for Each Value of Delta $\sigma^2/\sigma'^2 = 16, m' - \mu = 5$	41
4. Data for the Bayesian Approximation Model Based on One Hundred Runs for Each Value of Delta $\sigma^2/\sigma'^2 = 16, m' - \mu = 10$	42
5. Data for the Bayesian Approximation Model Based on Various Values of $ m' - \mu $ $\sigma^2/\sigma'^2 = 16$	44
6. Data for the Bayesian Approximation Model Based on Various Values of the Ratio of the Variances $ m' - \mu = 5$	46

LIST OF FIGURES

Figure	Page
1. Bayes Theorem for Continuous Random Variables	18
2. Bayes Theorem for Normal Distributions	23
3. Separation of the Posterior and Sample Means	39

SUMMARY

The research is devoted to modifying the Bayesian techniques associated with determining the minimum sample size required to construct interval estimates of the true mean of an experimental or sampling process which is modeled by a normal distribution with unknown parameters. The procedure considers only the case where the prior information can be represented by a normal distribution with known mean and known variance.

Rigorous Bayesian analysis of this situation would result in using a posterior distribution which has a normal-gamma density to construct interval estimates. In order to circumvent the obvious difficulties of working with this rather complex density, a procedure, which is felt to be more compatible to the U. S. Army Operational Testing environment, is offered for approximating the required Bayesian sample size.

If the variance of the sampling or experimental process, σ^2 , were known, the minimum Bayesian sample size required to construct an interval estimate about the true mean, μ , with confidence coefficient, α , and width, k , is:

$$n_b^* = \left[\frac{2 Z_{\alpha/2} \sigma}{k} \right]^2 - \frac{\sigma^2}{\sigma'^2}$$

where $Z_{\alpha/2}$ refers to the percentage points of the standard normal distribution such that $P(Z > Z_{\alpha/2}) = \alpha/2$, and σ'^2 is the variance of the prior distribution or, the prior variance of the true sampling mean, μ .

Substituting the sample variance, S^2 , for the true process variance and the term $t_{\alpha/2, n^*_c - 1}$ for $Z_{\alpha/2}$ in the above expression, and defining the width of the confidence interval as a function of the sample variance, i.e., $k = \delta S$, results in the following expression for the approximate Bayesian sample size:

$$n^*_b = \left[\frac{2 t(\alpha/2, n^*_c - 1)}{\delta} \right]^2 - \frac{S^2}{\sigma^2}$$

where n^*_c is the classical sample size required to construct an interval estimate of width k about the true mean, μ , of the sampling process.

The term $t(\alpha/2, n^*_c - 1)$ refers to the percentage points of the Student's t distribution with $n^*_c - 1$ degrees of freedom such that $P[(t > t(\alpha/2, n^*_c - 1))] = \alpha/2$.

The expression for the approximate Bayesian sample size is solved iteratively, starting with a fraction of the classical sample size required for the interval estimate of the same specified confidence and accuracy, as the first approximation. The iterative procedure is programed for a UNIVAC 1108 computer and applied to a hypothetical example to demonstrate the effectiveness of the methodology.

If accurate prior information is available, the results achieved by the procedure developed to approximate the Bayesian sample size and construct interval estimates of the unknown sampling mean, μ , of specified confidence and accuracy are comparable to the results obtained using classical techniques. These results, however, are achieved using smaller samples sizes than required for the classical case. A pro-

cedure was suggested for examining the accuracy of the prior information and ascertaining whether or not Bayesian analysis was appropriate for a given sampling or experimental situation. However, the expected results of using this procedure were not obtained.

CHAPTER I

INTRODUCTION

The General Problem

This study is an investigation of the problem of determining the minimum sample size of an experiment, that is, the minimum number of replications of the experiment, required to estimate the mean of the experimental variable to within a predetermined accuracy. In general, the study is limited to a certain type of testing situation which has the following characteristics:

- a. The test variable, whose mean is to be estimated, can be modeled by a normal distribution with unknown mean and unknown variance.
- b. Information is available, prior to sampling or experimenting, from which a probability distribution of the mean of the test variable can be constructed.
- c. This prior information can be represented by a normal distribution with known mean, m' , and known variance, σ'^2 .

The Specific Problem

A U. S. Army Operational Test (OT) is an overall evaluation of a system which has been developed for general use within the U. S. Army structure (2). In this context, a "system" may be not only hardware, but also doctrinal concepts, and is usually a mixture of both. The test is conducted in an environment which duplicates or closely simulates those conditions under which the system will be employed if

it is adopted for general use. It is this fact that generally differentiates OT from engineering, developmental, pre-production, and other tests which may be conducted on the same system and which are usually a part of the total development scheme of the system. In fact, OT is required to be an independent evaluation of the system. Thus, OT is a vital part of the process by which new equipment and concepts are incorporated into the U. S. Army structure.

An Operational Test is in essence a systematic plan for evaluating the total system being tested. It is composed of numerous subtests which address specific issues (unknown parameters) which are considered critical or paramount to the total evaluation of the system (3). The specific critical issues to be evaluated by each subtest and the order of these tests govern the overall structure of the Operational Test.

Once the specific structure of the Operational Test has been established, a decision must be made as to the number of replications of each subtest to conduct in order to properly evaluate the critical issue in question. Time and budget constraints place emphasis on conducting the minimum number of replications possible; while the disastrous consequences that could result if a critical issue is not properly evaluated, make it imperative that accuracy is not sacrificed for economy. Thus, the problem reduces down to one of determining the minimum number of replications of each subtest to conduct in order to evaluate the critical issues in question to within a predetermined accuracy.

Current procedural and policy documents governing the conduct of Operational Tests (3, 4, 11) suggest that for the most part, sample

sizes are determined using non-Bayesian or classical statistical methods. These methods do not consider prior information available concerning the variable being tested; therefore, inferences and decisions about the variable are based entirely on the experimental or sampling results. Bayesian techniques, on the other hand, attempt to use both the prior information and the experimental results in making inferences and decisions about the variable. Thus, this investigation is essentially a search for a practical procedure for applying Bayesian techniques to Operational Testing. The principal Operations Research tools used in this study are statistical inference and estimation techniques to develop the methodology, and computer simulation techniques to demonstrate the procedures developed.

Operational Testing is an expensive undertaking which must operate in an environment constrained by budget and time considerations. The author believes that a methodology which effectively reduces the number of replications required to evaluate the critical issues addressed by each subtest and also maintains the accuracy and confidence desired of the test, is a worthwhile pursuit directly applicable to the Operational Testing environment.

Background

During the last decade, there has been an increasing emphasis and drive within the military community to develop and formalize a methodology to adequately identify and evaluate the risks associated with the development and procurement of major weapons systems. The underlying premise which initiated this action was that unanticipated

cost and time over-runs and performance shortcomings, which had become increasingly prevalent, were the result of inadequate assessment of the risks involved with the materiel acquisition process. The methodology which grew out of this effort is known as decision risk analysis. In a report prepared for the Army Materiel Systems Analysis Agency (AMSAA), Atzinger, Brooks, et al., (6) present a brief history and description of the major concepts of the decision risk analysis process. The authors define risk analysis as follows: "Decision risk analysis is a discipline of systems analysis, which in a structured manner, provides a meaningful measure of the risks associated with various alternatives." The purpose of the report is to structure this decision risk analysis process so that the trade-offs inherent in the alternatives are visibly and meaningfully displayed. It cites the following four major areas as the underlying concepts of decision risk analysis:

- a. Subjective Probability
- b. Monte Carlo Methods
- c. Network Analysis
- d. Bayesian Statistics

Bayesian statistics and Bayes Theorem have attracted renewed interest in many fields of applied and theoretical statistics in recent years. This theorem is essentially a mechanism for combining new information with previously available information so that decisions or inferences can be based on all the information available. Over the years, a controversy has developed between the Bayesian and the more orthodox classical statistical concepts. Anscombe (1) provides a brief

but concise history of the development of both philosophies.

During the last few years there has been a revival of interest among statistical theorists in a mode of argument going back to the Reverend Thomas Bayes¹ (1702-61), Presbyterian minister at Tunbridge Wells in England, who wrote an "Essay Towards Solving a Problem in the Doctrine of Chances," which was published in 1763 after his death. Bayes work was incorporated in a great development of probability theory by Laplace and many others, which had general currency right into the early years of the century. Since then there has been an enormous development of theoretical statistics, by R. A. Fisher, J. Neyman, E.S. Pearson, A. Wald and many others, in which the methods and concepts of inference used by Bayes and Laplace have been rejected.

The orthodox statistician, during the last twenty-five years or so, has sought to handle inference problems (problems of deciding what the figures mean and what ought to be done about them) with the utmost objectivity. He explains his favorite concepts, significance level, confidence coefficient, unbiased estimates, etc., in terms of what he calls probability, but his notion of probability bears little resemblance to what the man in the street means (rightly) by probability. He is not concerned with probable truth or plausibility, but he defined probability in terms of frequency of occurrence in repeated trials, as in a game of chance. He views his inference problems as matters of routine, and tries to devise procedures that will work well in the long run. Elements of personal judgment are as far as possible to be excluded from statistical calculations. Admittedly, a statistician has to be able to exercise judgment, but he should be discreet about it and at all costs keep it out of the theory. In fact, orthodox statisticians show a great diversity in their practice, and in the explanations they give for their practice; and so the above remarks, and some of the following ones, are no better than crude generalizations. As such, they are, I believe, defensible. (Perhaps it should be explicitly said that Fisher, who contributed so much to the development of the orthodox school, nevertheless holds an unorthodox position not far removed from the Bayesian; and that some other orthodox statisticians, notably Wald have made much use of formal Bayesian methods, to which no probabilistic significance is attached.)

The revived interest in Bayesian inference starts with another posthumous essay on "Truth and Probability," by F. P. Ramsey² (1903-30), who conceived of a theory of consistent behavior by a person faced with uncertainty. Extensive developments were made by B. de Finette and (from a rather different point of view) by J. Jefferys. For mathematical statisticians the most thorough study of such a theory is that of L. J. Savage^{3,4}. R. Schlaifer⁵ has persuasively illustrated the new approach by reference to a variety of business and industrial problems. Anyone curious to obtain some insight into the Bayesian method, without mathematical hardship, cannot do better than browse in Schlaifer's book.

The Bayesian statistician attempts to show how the evidence of observations should modify previously held beliefs in the formation of rational opinions, and how on the basis of such opinions and of value judgments a rational choice can be made between alternative available actions. For him probability really means probability. He is concerned with judgments in the face of uncertainty, and he tries to make the process of judgment as explicit and orderly as possible.

Atzinger, Brooks, et al., (6) obviously consider Bayesian statistical procedures to have great potential in the decision risk analysis process; they state:

Bayesian statistics enjoys a unique position in risk analysis. There frequently exist situations where the analyst has both data and expert judgment to draw upon in constructing the probability distribution of interest in the consolidation activity. Bayesian statistics provides the analyst with a tool for synthesizing all of this information into one probability distribution which can then be used to directly estimate risks.

Review of the Literature

The statistical literature dealing with sample size determination is quite extensive, particularly in the area of classical techniques.

- ¹ Bayes, T., Essay Towards Solving a Problem in the Doctrine of Chances, reprinted with bibliographical note by G. A. Barnard, *Biometrika*, 45 (1958), 293-315.
- ² Ramsey, F. P., The Foundations of Mathematics, London: Rowtledge and Kegan Paul, 1931.
- ³ Savage, L. J., The Foundations of Statistics, New York, John Wiley, 1954.
- ⁴ Savage, L. J., Subjective Probability and Statistical Practice, to be published in a Mehtuen Monograph.
- ⁵ Schlaifer, R., Probability and Statistics for Business Decisions: An Introduction to Managerial Economics Under Uncertainty, New York, McGraw-Hill, 1959.

Mace (12) provides an excellent and thorough coverage of classical procedures for determining the optimum sample size of a research experiment. This publication is applications oriented and provides procedures, formulas, and tables for determining economical sample sizes for some forty different types of research objectives. Unfortunately, the author considers only one rather limited application of Bayesian techniques to sample size determination. The limitation in this particular example, that the variance of the sampling process must be known, seems to occur quite frequently in the literature of Bayesian techniques for determining minimum sample sizes.

There has been extensive research in the application of Bayesian techniques to reliability engineering and quality control. White (15) presents a promising methodology for periodic reliability assessment using Bayesian techniques to combine analytical predictions with limited test results to obtain greater precision in the reliability estimate. The main limitation of this paper is that it considers only the gamma distribution in the analysis. Gilbreath (8) has devised sampling procedures for use in sequential sampling models which have direct application in quality control and in economic lot size determination. These techniques, however, are more applicable to hypothesis testing than to the estimation problem.

Atzinger and Brooks (5) provide an excellent comparison of Bayesian and classical decision making under uncertainty for a class of problems where the decision variable is the Bernoulli success probability, p . If the outcome of any particular test or experiment is

viewed as a success or failure, the resulting data classification is characteristic of a Bernoulli process. The authors persuasively argue that historically, one of the major objectives in test and evaluation processes has been to estimate this unknown Bernoulli success parameter. Unfortunately, such an analysis does not address the actual parameters of the sampling or experimental process itself.

Winkler (16) provides a rather detailed and complete development and treatment of Bayesian applications to inference and decision theory at the introductory level. Although the concepts developed in this publication are very thoroughly covered, the scope of the material is rather limited. That is, only two specific sampling processes are analyzed in detail: the sampling process modeled by the Bernoulli distribution, and the sampling process represented by the normal distribution with known variance.

Raiffa and Schlaifer (13) provide an extensive mathematical development of Bayesian techniques applied to statistical decision theory. However, once again, extensive analysis of the normal distribution is generally restricted to the case where the variance of the sampling population is known.

Thus, Bayesian applications to the problem of sample-size determination deal only with very specialized situations in the current literature. There appears to be no substantial research into the examination of the general problem. On the other hand, classical statistical techniques commonly apply iterative type algorithms to the to the general problem of sample size determination. The author believes that these techniques can be validly extended to Bayesian analysis and

produce equally valid results. The aim of this investigation, then, is to extend the application of these well known techniques to the general sampling situation using Bayesian analysis.

CHAPTER II

THE TEST METHODOLOGY

The Assumptions of Normality

The normality assumptions stated in the introduction introduction are crucial, albeit restrictive, to this investigation. The assumption that the prior distribution, which represents the distribution of the mean of a random variable, is normally distributed has solid support in the Central Limit Theorem. Hines and Montgomery (9) state the essence of this important theorem as follows:

If $X_1, X_2, \dots, X_n$ is a sequence of n independent random variables with $E(X_i) = \mu_i$ and $V(X_i) = \sigma_i^2$ (both finite) and $Y = X_1 + X_2 + \dots + X_n$, then under some general conditions

$$Z_n = \frac{Y - \sum_{i=1}^n \mu_i}{\sqrt{\sum_{i=1}^n \sigma_i^2}}$$

has an approximate $N(0,1)$ distribution as n approaches infinity.

The "general conditions" mentioned in the theorem are informally summarized as follows: The terms X_i , taken individually, contribute a negligible amount to the variance of the sum, and it is not likely that a single term makes a large contribution to the sum.

The principal implication of this theorem, then, is that in general the sum of n independent random variables is approximately normally distributed for sufficiently large n , regardless of the distribution of the n individual random variables. Unfortunately, the

assumption that the random variable to be tested is normally distributed is much more restrictive. However, in many cases, real-world situations can be satisfactorily approximated by a normal process. Also, statistical inference and estimation procedures, particularly those concerning the mean of random variables, are generally robust (insensitive) to the normality assumption (12).

The Prior Information

At first glance, the requirement that Operational Testing be independent of other testing conducted on the same system may seem an insurmountable obstacle in attempting to obtain adequate prior information. This, however, is usually not the case; other sources of prior information do exist. For example, most new systems undergoing testing have been specifically designed to replace older or outmoded systems which are currently a part of the U. S. Army structure. These older systems represent a vast source of historical data from which prior distributions for nearly any critical issue can be developed. In those rare cases where no historical data exist from which to construct a prior distribution for a specific critical issue, the Delphi technique or other proven methods of developing subjective assessments of uncertainties can be used to develop the prior distribution (6).

In any event, to the Bayesian statistician, the prior information represents the best available estimate about an uncertain quantity, regardless of its source. This fact even suggests that it is reasonable and logical to modify the prior distribution developed from historical data to reflect the improved design characteristics of the

new system. Suppose, for example, that one of the critical issues being evaluated during OT of a new weapons system is the accuracy of the weapon at a specified range. The distribution of the mean-error of similar weapons currently in use can be determined from historical data. If the new system is expected to be significantly more accurate because of new design characteristics, the mean of the prior distribution developed from the historical data can be adjusted to reflect the expected increase in the performance of the new system. In discussing techniques for the assessment of prior distributions and the use of diffuse prior distributions to represent the situation where no prior information is available, Winkler (16) states:

It should be stressed that in general, there is no such thing as a "totally informationless" situation and the use of particular distributions to represent diffuse prior states of information is a convenient approximation that is applicable only when the prior information is "overwhelmed" by the sample information. In most real-world situations, non-negligible prior information (non-negligible relative to the sample information) is available, and the concept of a diffuse prior distribution is not applicable.

The Basic Alternatives of Determining Sample Size

This study considers only two basic approaches to determining the appropriate sample size in an experimental process. One approach is to simply disregard any prior knowledge or information available about the variable of interest, and use classical statistical techniques to solve the problem. The other approach is to combine the prior information with the results of a limited number of replications of the experiment, if possible, and then use these results to solve the problem.

The Classical Method

Classical estimation procedures and techniques are well documented in the literature (9, 10, 12). This method uses only the results of the sampling or experimental process in the estimation procedures and ignores all prior information. Starting from the basic assumption that the sampling process is normally distributed with unknown mean, μ , and unknown variance, σ^2 , the random variable representing the outcome of the sampling process can be represented by:

$$X_i \sim N(\mu, \sigma^2), \quad \text{with } \mu, \sigma^2 \text{ unknown}$$

Let $(X_1, X_2, \dots, X_n)$ represent the results of n replications of the experiment. The sample statistics based on the specific n values obtained from the sampling process can be expressed as:

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i, \quad \text{the sample mean}$$

and

$$s^2 = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n - 1}, \quad \text{the sample variance}$$

The appropriate expression for a $(1 - \alpha)$ percent confidence interval about the unknown mean, μ , for a process which is normally distributed and for which the variance is unknown is constructed using the Student's t distribution, i.e.;

$$P\left(\bar{X} - t(\alpha/2, n-1) \frac{S}{\sqrt{n}} \leq \mu \leq \bar{X} + t(\alpha/2, n-1) \frac{S}{\sqrt{n}}\right) = 1-\alpha \quad (2-1)$$

where the expression $t(\alpha/2, n-1)$ refers to the percentage points of the Student's t distribution with $n-1$ degrees of freedom such that $P(t > t(\alpha/2, n-1)) = \alpha/2$.

Recall from the introduction that the classical interpretation of probability differs considerably from the Bayesian interpretation. Thus, the interpretation of equation (2-1) is based on long-run considerations. That is, the classical statistician would say that if a confidence interval based on a sample size of n is constructed each time, then in the long run, $1-\alpha$ percent of such intervals would contain the true mean of the normally distributed sampling process. The value of α , which is preselected at some low value, can then be thought of as protection against failure of the interval to include the true value of the mean of the sampling process. The value, $\alpha = 0.05$, is often selected for statistical inference and estimation problems because of traditional usage. The second type of error that can occur in interval estimation problems is that the interval constructed based on a set of specific sample results may be too wide, even though the interval does include the true value of the mean of the sampling process. This, then, is a problem of the accuracy associated with the confidence interval. Protection against this type of error is accomplished by controlling the width of the confidence interval constructed. The width of each specific confidence interval is dependent on the

sample size and the value of α specified.

The terms

$$UL = \bar{X} - t(\alpha/2, n-1) \frac{S}{\sqrt{n}}$$

and

$$UU = \bar{X} + t(\alpha/2, n-1) \frac{S}{\sqrt{n}},$$

which are real-valued functions of the sample results, are the lower and upper limits, respectively, of the interval estimate. The Student's t distribution is very similar to the standard normal distribution, and for degrees of freedom, $v = n-1 > 20$, the two distributions are virtually indistinguishable. And in fact, the Student's t distribution is identical to the standard normal distribution for degrees of freedom, $v = \infty$ (10). This fact allows accurate approximations in computing the minimum sample size by approximating the value of $t(\alpha/2, n-1)$ by $t(\alpha/2, \infty) = Z_{\alpha/2}$ for moderate sample sizes. The expression $Z_{\alpha/2}$ refers to the percentage points of the standard normal distribution such that $P(Z > Z_{\alpha/2}) = \alpha/2$.

For the moment, let the preselected width of the confidence interval be simply equal to k . Then from equation (2-1), the half-interval width can be expressed as:

$$t(\alpha/2, n-1) \frac{S}{\sqrt{n}} = \frac{k}{2}$$

Solving this equation for n , results in the following expression for the minimum sample size required for a confidence interval width equal to k .

$$n_c^* = \left[\frac{2t(\alpha/2, n_c^* - 1)S}{k} \right]^2 \quad (2-2)$$

It is more convenient to express the width of the confidence interval in terms of the sample standard deviation in order to simplify equation (2-2). Thus, if $k = \delta S$ is substituted into the equation, the minimum sample size required can then be expressed as:

$$n_c^* = \left[\frac{2t(\alpha/2, n_c^* - 1)}{\delta} \right]^2 \quad (2-3)$$

Equation (2-3) cannot be solved explicitly for n_c^* , since the value of $t(\alpha/2, n_c^* - 1)$ is a function of the sample size n_c^* . But since the value of $t(\alpha/2, n_c^* - 1)$ is approximately equal to $t(\alpha/2, \infty)$, which is equal to $Z_{\alpha/2}$, for moderate sample sizes a good first approximation for the solution of equation (2-3) is obtained by substituting the value of $Z_{\alpha/2}$ for the value $t(\alpha/2, n_c^* - 1)$. This first approximation is known to be too small, although for large sample sizes it is quite close to the actual value of n_c^* . Using this first approximation, call it n_0 , to evaluate $t(\alpha/2, n_0 - 1)$ and to solve equation (2-3) again, to obtain a better second approximation for the value of n_c^* . This iterative procedure can be used to approximate the value of n_c^* to any desired accuracy; however, there is usually no significant improvement in the approximation after the second or third iteration.

Table 1 shows the values of n_c^* obtained for a 95 percent ($\alpha = 0.05$) confidence interval for various values of δ using this iterative procedure. Because of the premium placed on accurate estimates in Operational Testing, values of $\delta > 1.0$ were not considered. The values shown in the table under the heading $P(K)$ are the approximate probabilities of a single observation from the sampling process falling between the lower and upper limits of the confidence interval, i.e., $P(K) = P(\underline{LX} \leq x \leq \overline{UX})$. This value gives a probabilistic measure of the accuracy (width) of the confidence interval. The values of n_c^* in the table have been rounded up to the next highest integer. As illustrated in Table 1, equation (2-3) points out that in order to decrease the width of a confidence interval by one-half, the sample size must be increased approximately by a factor of four.

Table 1. Minimum Sample Size - Classical Method

δ	$P(K)$	n_c^*
1.0	0.383	18
0.9	0.347	22
1.8	0.311	27
0.7	0.274	34
0.6	0.236	46
0.5	0.197	64
0.4	0.159	99
0.3	0.119	174
0.2	0.080	387
0.1	0.040	1537

A Bayesian Approximation

Bayes Theorem for Continuous Random Variables. The essence of Bayes Theorem for continuous random variables is depicted in Figure 1 shown below. The densities $f(\theta)$ and $f(\theta|y)$ represent the prior distribution and the posterior distribution respectively, and $f(y|\theta)$ represents the likelihood or sampling function. It is important to keep in mind always that it is the prior distribution or the statisticians prior state of knowledge that is modified by the sampling results and not the reverse.

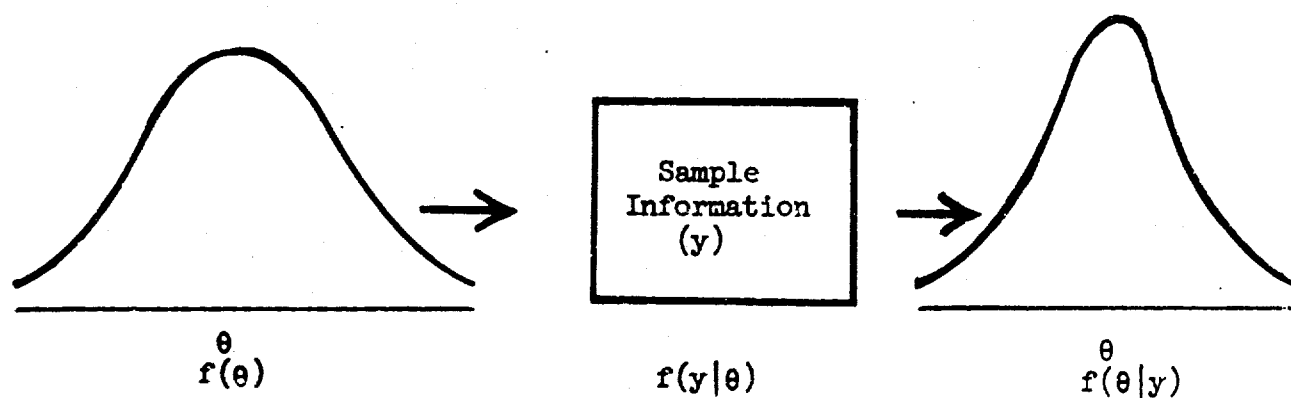


Figure 1. Bayes Theorem for Continuous Random Variables

The prior and posterior distributions must be proper density functions. That is, they must possess the following mathematical properties applicable to the density function of any continuous random variable, $\tilde{x}$, which has range space or domain, R_x :

- (i) $f(x) \geq 0$ for all $x \in R_x$
- (ii) $\int_{R_x} f(x) dx = 1$

The likelihood or sampling function, $f(y|\theta)$, represents the probability of obtaining a given value, y , for the range of possible values of θ . The likelihood function is not a proper density function because the events $f(y|\theta)$ are not mutually exclusive over the range of θ .

As suggested in Figure 1, Bayes Theorem is essentially a process of combining the prior distribution with the sample information to yield the posterior distribution. The resultant posterior density has the following form:

$$f(\theta|y) = \frac{f(\theta) f(y|\theta)}{\int_{-\infty}^{\infty} f(\theta) f(y|\theta) d\theta} \quad (2-4)$$

This result can be expressed in words as:

$$\text{posterior density} = \left[\begin{array}{c} \text{normalizing} \\ \text{constant} \end{array} \right] \left[\begin{array}{c} \text{prior} \\ \text{density} \end{array} \right] \left[\begin{array}{c} \text{likelihood} \\ \text{function} \end{array} \right]$$

where the normalizing constant, $1/\int_{-\infty}^{\infty} f(\theta)f(y|\theta)d\theta$, is needed to make the posterior distribution a proper density function.

Before the advent of the high speed computer which greatly eased the computational burden involved with numerical integration techniques, application of equation (2-4) to revise density functions in the light of sample information often proved extremely difficult because of the integration required to compute the normalizing constant. For this reason, Bayesian statisticians developed the concept of "conjugate" distributions, which are families of distributions that ease the computational burden when they are used as prior distributions (16). Of

course the resultant form of the posterior distribution depends on the likelihood function as well as the prior distribution. Thus, conjugate prior distributions are selected on the basis of the statistical properties of the model chosen to represent the sampling process. When the prior distribution is conjugate to the likelihood or sampling function, the resultant posterior distribution is also a member of the same conjugate family of prior distributions.

Bayes Theorem for Normal Distributions. If it is possible to model the population or process being sampled by a normal distribution, the proper choice for a family of conjugate prior distributions depends on the statistician's knowledge of the parameters of the normal data generating process used. Raiffa and Schlaifer (13) summarize the effects of the statistician's knowledge of the two parameters of the normal distributions on the proper choice of conjugate prior distributions as follows:

Case (i) μ known, σ^2 unknown: The appropriate family of conjugate distributions have a gamma-2 density.

Case (ii) σ^2 known, μ unknown: The appropriate family of conjugate distributions have a normal density.

Case (iii) both μ and σ^2 unknown: The appropriate family of conjugate distributions have a normal-gamma density.

An Approximation Procedure. Since it was assumed that within the context of this study the model representing the sampling process in Operational Testing was normally distributed with unknown mean, μ , and unknown variance, σ^2 , the appropriate family of conjugate distributions to use in this case have a normal-gamma density. In order to overcome the obvious difficulties associated with computing interval

estimates with the normal-gamma density, a procedure is suggested here to modify the Bayesian analysis of this sampling process so that the family of conjugate prior distributions have a normal density function; as is the case when the variance of the population or sampling process is known.

Assume for the moment that the variance of the sampling process is known. Then the conjugate prior distribution has a normal density function of the form:

$$f'(\mu) = \frac{1}{\sqrt{2\pi\sigma'^2}} e^{-(\mu - m')^2/2\sigma'^2}$$

where the prime (') is used to signify a parameter or constant which is associated with the prior distribution. Thus, σ'^2 is the variance of the prior distribution or, the prior variance of the unknown parameter, μ ; and m' is the mean of the prior distribution of this parameter.

If n replications of the experiment were now conducted and a sample mean,

$$m = \frac{1}{n} \sum_{i=1}^n X_i ,$$

and a sample variance,

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - m)^2 ,$$

were observed, the resultant posterior distribution would also have a

normal density function of the form:

$$f''(\mu|y) = \frac{1}{\sqrt{2\pi\sigma''^2}} e^{- (\mu - m'')^2 / 2\sigma''^2}$$

where y represents the sample results, and the double prime (") is used to indicate a parameter or constant which is associated with the posterior distribution. Thus, σ''^2 is the posterior variance of μ , and m'' is the mean of the posterior distribution of μ . These posterior parameters can be computed from the following formulas:

$$\frac{1}{\sigma''^2} = \frac{1}{\sigma'^2} + \frac{n}{\sigma^2} \quad (2-5)$$

and

$$m'' = \frac{(1/\sigma'^2)m' + (n/\sigma^2)m}{(1/\sigma'^2) + (n/\sigma^2)} \quad (2-6)$$

Equations (2-5) and (2-6) indicate that the reciprocal of the posterior variance is equal to the sum of the reciprocal of the prior variance, σ'^2 , and the reciprocal of the variance of the sample mean, σ^2/n . The posterior mean is a weighted average of the prior mean, m' , and the sample mean, m . The weights being the reciprocal of the respective variances.

As depicted in Figure 2., an important feature of the posterior distribution is that the posterior mean, m'' , always lies between the prior mean, m' , and the sample mean, m . The posterior variance, σ''^2 ,

is always smaller than the prior variance, σ'^2 (16). From equation (2-5), if the variance of the prior distribution, σ'^2 , decreases, the amount of prior uncertainty decreases, and the prior information is given more weight in the determination of the posterior distribution. Similarly, as the variance of the sample mean, σ^2/n , decreases, the sampling information is given more weight in the determination of the posterior distribution.

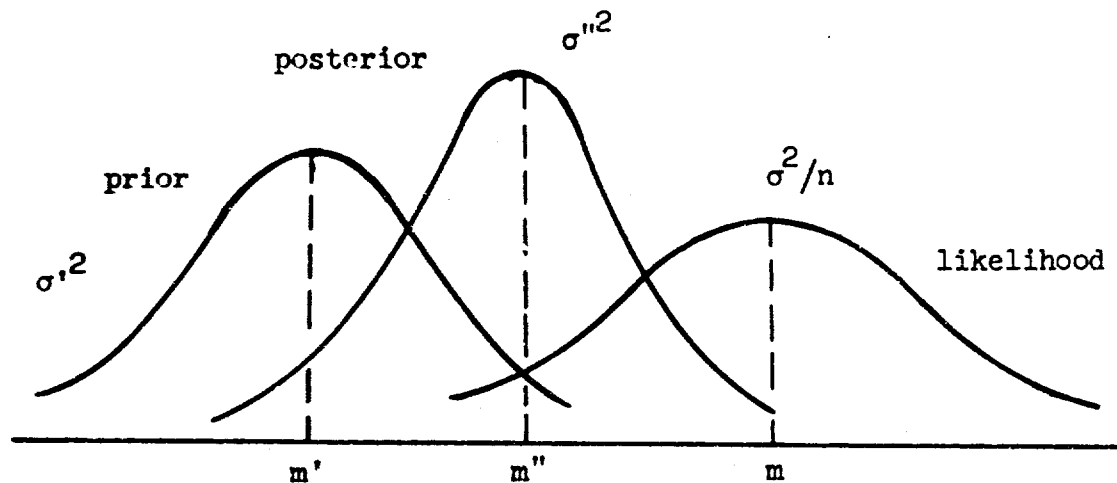


Figure 2. Bayes Theorem for Normal Distributions

A different parameterization of this problem might help clarify the results obtained.

$$\text{let } n' = \frac{\sigma^2}{\sigma'^2}$$

Then the prior variance can be written in terms of n' and the process or sampling variance, thus:

$$\sigma'^2 = \frac{\sigma^2}{n'}$$

Similarly, if

$$n'' = \frac{\sigma^2}{\sigma''^2}$$

then

$$\sigma''^2 = \frac{\sigma^2}{n''}$$

Substituting these results into equations (2-5) and (2-6), the parameters of the posterior distribution are then

$$\frac{n''}{\sigma^2} = \frac{n'}{\sigma^2} + \frac{n}{\sigma^2}$$

or simply,

$$n'' = n' + n$$

and

$$m'' = \frac{(n'/\sigma^2)m' + (n/\sigma^2)m}{(n'/\sigma^2) + (n/\sigma^2)}$$

or simply,

$$m' = \frac{n'm' + nm}{n' + n} \quad (2-10)$$

In his interpretation of the results obtained by using these new parameters, Winkler (16) suggests that the prior distribution can be thought of as roughly equivalent to the information contained in a sample of size n' with a sample mean of m' from a normal sampling process with variance σ^2 . That is, n' appears to be the sample size required to produce a variance of σ'^2 for a sample mean equal to m' , since the variance of the sample mean from a sample size n' is equal to σ^2/n' . Winkler also considers equations (2-7) and (2-3) as formulas for pooling the information from the two samples. Under this interpretation, the posterior or pooled sample size is equal to the sum of the two individual sample sizes, one from the prior distribution and one from the sampling process. The posterior or pooled sample mean is equal to a weighted average of the two individual sample means.

This pooling process suggests that a reasonable estimate of the sample mean, based on all the information available, is the posterior or pooled mean, m'' . Notice that if $n' > n$, then the posterior or pooled mean is closer to the prior mean than to the sample mean. That is, the prior information is given more importance than the sample results in the determination of the posterior parameters. Of course, the posterior mean is closer to the sample mean if $n > n'$; and if $n' = n$, the posterior mean is exactly midway between the prior mean and the

sample mean. Notice also that since the sample mean, m , is as equally likely to fall above as it is to fall below the true population or sampling mean, μ ; it is then equally likely that the sample mean and the mean of the prior distribution, m' , to be on the same or opposite sides of μ . When m' and m fall on the same side of μ , the mean of the posterior distribution, m'' , will be further from μ than the sample mean. That is, the posterior mean will be a less accurate estimate of the true population mean than the sample mean. When m' and m are on opposite sides of μ , then it cannot be determined whether the posterior mean will be closer or further from μ than the sample mean. Each specific case must be examined separately; the results will depend on the sample size, the specific value of the prior mean, and the variances of the prior and sampling distributions.

Since the point estimate of μ based on all information available is the posterior mean which is normally distributed with mean, m'' , and variance, σ''^2 , the statistic

$$Z = \frac{m'' - \mu}{\sigma''}$$

has a standard normal distribution, i.e., $Z \sim N(0, 1)$. Therefore the appropriate expression for a $(1 - \alpha)$ percent interval estimation of μ for this case is constructed using the standard normal distribution, i.e.,

$$P(m'' - Z_{\alpha/2} \sigma'' \leq \mu \leq m'' + Z_{\alpha/2} \sigma'') = 1 - \alpha \quad (2-11)$$

The lower and upper limits of the confidence interval in this case are $UL = m'' - Z_{\alpha/2} \sigma''$ and $UU = m'' + Z_{\alpha/2} \sigma''$, respectively.

If, as was done in the classical case, the width of the confidence interval for the general case is set equal to k , then from equation (2-11) the half-interval width can be expressed as:

$$Z_{\alpha/2} \times \sigma'' = \frac{k}{2}$$

Now substituting the expression for σ''^2 from equation (2-5) into the above expression results in the following:

$$Z_{\alpha/2} \left[\frac{1}{1/\sigma'^2 + n/\sigma^2} \right]^{\frac{1}{2}} = \frac{k}{2}$$

$$Z_{\alpha/2} \left[\frac{\sigma'^2 \sigma^2}{\sigma^2 + n\sigma'^2} \right]^{\frac{1}{2}} = \frac{k}{2}$$

$$\left[\frac{Z_{\alpha/2} \sigma' \sigma}{k/2} \right]^2 = \sigma^2 + n\sigma'^2$$

and finally,

$$n_b^* = \left[\frac{2Z_{\alpha/2} \sigma}{k} \right]^2 - \frac{\sigma^2}{\sigma'^2} \quad (2-12)$$

This then, is the Bayesian solution for the minimum sample size required to establish a confidence interval of width k about the mean of the sampling process under the special condition that the variance

of the sampling process is known. Several characteristics of equation (2-12) deserve mention. First of all, the first term in the equation, $[2 Z_{\alpha/2} \sigma/k]^2$, is in fact the exact expression for the classical solution to the problem of determining the minimum sample size required to establish a confidence interval of width k about the mean of a sampling process with known variance. Second, the last term in the equation, σ^2/σ'^2 , is the expression developed earlier for n' in equation (2-7). Recall Winkler's interpretation of n' as being roughly the equivalent sample size, relative to the sampling process, of the information contained in the prior distribution. The ratio $\sigma^2/\sigma'^2 = 0$ is also used to define a diffuse prior distribution, i.e., an informationless prior state. Assuming that the variance of the sampling process, $\sigma^2 > 0$, the ratio $\sigma^2/\sigma'^2 = 0$ only if the variance of the prior distribution, $\sigma'^2 = \infty$. In this case, the variance of the prior distribution would represent a condition of total uncertainty and since $n' = \sigma^2/\sigma'^2 = 0$, equation (2-12) would yield the same results as in the classical case.

Tying all these facts together, equation (2-12) can be interpreted as follows: the minimum Bayesian sample size required to establish an interval estimation of the mean of any specified width or accuracy is equal to the minimum sample size required to establish the same interval estimation using classical methods, minus the value of the prior information in terms of an equivalent sample size. Or, more clearly:

$$n_b^* = n_c^* - n'$$

Now, consider once again equation (2-12) in order to address the fact that the variance of the sampling process is in fact not known. Substituting the sample variance for the variance of the sampling process and the term $t(\alpha/2, n^*_c - 1)$ for $Z_{\alpha/2}$ in equation (2-12), and once again defining the width of the confidence interval as $k = \delta S$, results in the following expression for the approximate Bayesian sample size:

$$n^*_b = \left[\frac{2 t_{\alpha/2, n^*_c - 1}}{\delta} \right]^2 - \frac{S^2}{\sigma'^2} \quad (2-13)$$

where

$$S^2 = \frac{\sum_{i=1}^n (x_i - m)^2}{n^*_b - 1}$$

and m is equal to the sample mean based on n^*_b observations.

Examination of equation (2-13) reveals that the first term in the equation is identical to equation (2-3), the classical solution to the minimum sample size problem for a normal sampling process with unknown variance. The last term in the equation is an approximation of the equivalent sample size of the information contained in the prior distribution, where the value of $n' = \sigma^2/\sigma'^2$ is approximated by $n' = S^2/\sigma'^2$. Of course equation (2-13) cannot be evaluated explicitly, even though the value of the first term in the equation is exactly known from the results obtained using the classical method, since the value of S^2 depends on the specific observations obtained during the sampling process.

Before suggesting a procedure for approximating a solution to equation (2-13) for the general case, it may be more appropriate at this point to examine the general implications of using the posterior distribution to construct confidence interval estimates about the mean of the sampling process. An interval estimation based on the posterior distribution has as its midpoint the posterior mean, m'' ; while the midpoint of an interval estimation based on the sampling process alone is the sample mean, m . Referring to Figure 2, it is obvious, then, that an interval estimate of width δS which is based on the posterior distribution will not include the sample mean, m , if m'' and m are separated by more than $\frac{1}{2}\delta S$. A large separation between m'' and m is indicative of prior information which is not very compatible to the results obtained from the sampling or experimental results. In other words, the prior information does not predict the behavior of the sampling process very well. This is an important consideration in Operational Testing, since it is important to decide whether or not to use the prior information in estimating the mean of the sampling or experimental process.

It would seem appropriate then, to develop at least a heuristic rule to reject the use of prior information which causes the posterior and sampling means to differ beyond some pre-established limit. The general form of such a rule would be of the form:

$$|m'' - m| \leq q\delta S$$

where the value of q is selected in a manner such that if the inequality

were not satisfied, the application of Bayesian techniques would be aborted and the appropriate sample size for the specific situation would be determined by using classical techniques.

Returning now to the problem of constructing an interval estimate of width δS for the mean of the sampling process using Bayesian techniques, the following procedure is suggested as a reasonable approach to approximating the solution of equation (2-13) for the general case.

a. Determine the minimum sample size, n_c^* , required for the classical method. This value, call it n_0 , is the upper limit of the Bayesian sample size.

b. As a first approximation to the Bayesian sample size, let $n_1 = n_0 / d$. Where the value of d is selected with consideration given to the classical sample size being used. That is, for small values of n_c^* , d should be chosen at some low value (such as 2 or 4) in order that n_1 be large enough to yield suitable sample statistics. For large values of n_c^* , d may be increased since the resulting n_1 samples would still yield suitable statistics. The objective here is to approximate the Bayesian sample size conservatively while insuring that the approximation decided upon is large enough to yield reasonably valid sample statistics.

c. Conduct the n_1 replications of the experiment and from the results compute the sample statistics:

$$m_1 = \frac{\sum_{i=1}^{n_1} X_i}{n_1}$$

and

$$s_1^2 = \frac{\sum_{i=1}^n (x_i - m_1)^2}{n_1 - 1}$$

d. Use these statistics to compute the approximations:

$$n'_1 = \frac{s_1^2}{\sigma'^2}$$

and

$$m''_1 = \frac{n'_1 m' + n_1 m_1}{n'_1 + n_1}$$

e. Determine the second approximation of the Bayesian sample size by using the value obtained for the first approximation and the following relationship:

$$n_2 = n_1 + \Delta(n_0 - n'_1)$$

where Δ is chosen with the same considerations as was the value of $\underline{d}$. The expression for the $n_{j\text{th}}$ approximation of the Bayesian sample size is:

$$n_j = n_{j-1} + \Delta(n_0 - n'_{j-1})$$

f. Determine if sufficient replications of the experiment have

been conducted after each iteration by comparing the computed approximation of the Bayesian sample size to the classical sample size minus the computed value of n' . That is, continue the iterative procedure until $n_j \geq n_0 - n'_j$.

g. After computing the final approximation of the Bayesian sample size, determine if the prior information should be accepted or rejected. That is, if $|m'' - m| \leq q\delta S$, use the n_j replications already conducted to construct the interval estimate of the mean of the experimental process using Bayesian techniques. If $|m'' - m| > q\delta S$, reject the use of the prior information; conduct the remaining $n_0 - n_j$ replications of the experiment and construct the desired interval estimate of the mean of the experimental process using classical techniques.

CHAPTER III

DEMONSTRATION OF THE METHODOLOGY

Programming the Model

The model developed for approximating the minimum Bayesian sample size for the special test situation described in Chapter I is programmed for the UNIVAC 1108 computer using standard Fortran IV language. The program consists of four basic segments designed to perform the following functions: generate the required data and compute the sample statistics; compute the minimum classical sample size required for an interval estimation of specified width; compute the approximate Bayesian sample size required for the same interval width; and construct the confidence intervals desired based on the sampling results.

The Box and Mueller technique (7) is used to generate the normally distributed pseudo random numbers representative of a normal process with specified mean and variance. The random number generator was tested for various sample sizes and values of the model parameters using the chi-square goodness-of-fit test for normality. The results of these tests were quite favorable and are summarized in Table 2.

Equation (2-3) is solved iteratively for the minimum classical sample size by using two standard UNIVAC MATH-STAT library functions (14). The function TINORM is used to compute the value of the inverse of the standard normal distribution given the value of the probability for which the ordinate is to be calculated. The function STUDIN is

Table 2. Test of Normal Random Generator

Specified Mean and Variance	Number of Observations per Trial	Number of Trials	Number of Trials Accepted at the $\alpha = .05$ Significance Level
$\mu = -120.0$	500	30	27
$\sigma^2 = 81.0$	250	30	27
	100	30	29
	25	30	26
$\mu = -5.0$	500	30	26
$\sigma^2 = 4.0$	250	30	27
	100	30	28
	25	30	24
$\mu = 0.0$	500	30	27
$\sigma^2 = 1.0$	250	30	28
	100	30	29
	25	30	23
$\mu = 57.0$	500	30	30
$\sigma^2 = 64.0$	250	30	27
	100	30	28
	25	30	25
$\mu = 200.0$	500	30	28
$\sigma^2 = 400.0$	250	30	29
	100	30	27
	25	30	24
$\mu = 225.0$	500	30	26
$\sigma^2 = 25.0$	250	30	30
	100	30	27
	25	30	23

Table 2. (Continued)

Specified Mean and Variance	Number of Observations per Trail	Number of Trials	Number of Trials Accepted at the $\alpha = .05$ Significance Level
$\mu = 297.72$ $\sigma^2 = 173.18$	337	30	28
$\mu = 199.5$ $\sigma^2 = 46.32$	483	30	29
$\mu = 45.3$ $\sigma^2 = .78$	13	30	24
$\mu = 1542.0$ $\sigma^2 = 5742.6$	486	30	28

used to calculate the inverse of the Student's t distribution for a given confidence coefficient. The results obtained from the subroutine used to calculate the classical sample size for each specified value of δ are shown in Table 1.

Approximations for the Bayesian sample size for a given value of δ are computed using the iterative procedure developed in the preceding chapter. The value of the classical sample size computed for a given value of δ is input to this subroutine which uses this value to calculate the first approximation of the Bayesian sample size.

Confidence intervals are computed by using the STUDIN library function to calculate the value $t(\alpha/2, n^*-1)$, where n^* is the computed classical or Bayesian sample size. The subroutine then computes the lower and upper limits of the confidence interval, i.e.,

$$UL = \bar{m} - t(\alpha/2, n_c^* - 1) \frac{S_c}{\sqrt{n_c^*}}$$

and

$$UU = \bar{m} + t(\alpha/2, n_c^* - 1) \frac{S_c}{\sqrt{n_c^*}}$$

for the classical case, and

$$UL = \bar{m}'' - t(\alpha/2, n_b^* - 1) \frac{S_b}{\sqrt{n''}}$$

and

$$UU = m'' + t(\alpha/2, n_b^* - 1) \frac{S_b}{\sqrt{n''}}$$

for the Bayesian case.

Demonstrating the Model

In order to demonstrate the model developed to approximate the Bayesian sample size in the preceding chapter, various values of the constants, d , Δ , and q used in the iterative procedure were tried in preliminary simulations. The values $d = 4$, $\Delta = 1/4$, and $q = 3/8$ were chosen for the following reasons:

a. Values of $d < 4$ tended to produce first approximations of the Bayesian sample size which were too large when working with small values of the classical sample size, n_0 . That is, $n_1 \geq n_0 - n'_j$ after the first approximation. Larger values of d produced more conservative first approximations of the Bayesian sample size for small values of n_0 , but at the same time resulted in unreliable, i.e., greatly variable, sample statistics.

b. Values of $\Delta < 1/4$ were rejected because for large values of n_0 the number of iterations required to compute the approximate Bayesian sample size was considerably increased. It was felt that this result was undesirable in an Operational Testing mode and, of course, it also meant increased computer times to solve the approximation. A scheme of using a variable value for Δ was tried, i.e., Δ was decreased by one-half after each iteration. This scheme was also rejected because

for larger values of n_0 the iterative procedure quickly evolved into a sequential type of sampling procedure.

c. The value of $q = 3/8$ was selected as a reasonable choice based on the illustration shown in Figure 3. The interval a-d represents an interval estimation based on the posterior distribution. Then from previous definitions, $a-d = \delta S$, and the intervals $a-m'' = m''-d = 1/2 \delta S$. Then if the intervals $a-b = c-d = 1/8 \delta S$, the sample mean, m , is required to be within the interval $b-c = 3/4 \delta S$, i.e., $|m'' - m| \leq 3/8 \delta S$ is the prerequisite for incorporating the prior information into the estimation procedures. It was felt that $1/8 \delta S$ would allow for sufficient variation of the sample mean due to differences in sample results.

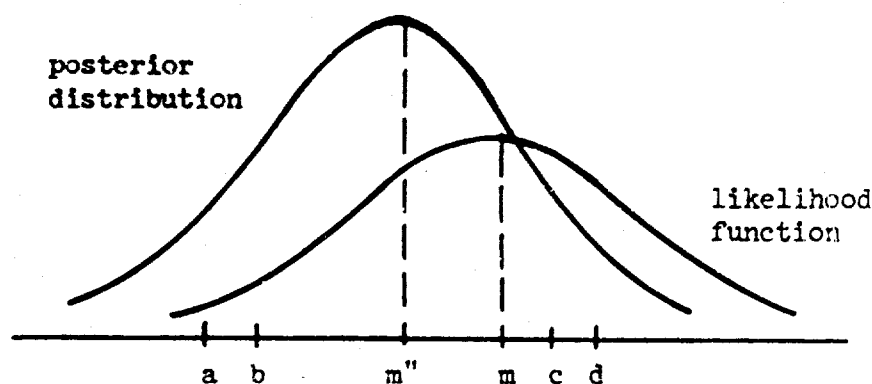


Figure 3. Separation of the Posterior and Sample Means

The procedure to approximate the Bayesian sample size was demonstrated using a hypothetical case having the following characteristics:

- a. the ratio $\sigma^2/\sigma'^2 = 16$.
- b. $|m' - \mu| = 5$ and $|m' - \mu| = 10$.

where μ and σ^2 are the true (but assumed unknown) values of the parameters of the sampling process and m' and σ'^2 are the parameters of the prior distribution.

The first test of the procedure involved a computer simulation of 100 runs for each value of delta from 1.0 to 0.2. The model was not tested for the value of delta equal to 0.1 in this or subsequent tests of the procedure because the large sample sizes involved required an excessive amount of computer time. The results of this first test are summarized in Table 3 for the case where $|m' - \mu| = 5$ and in Table 4 for the case where $|m' - \mu| = 10$. These results appear quite favorable as shown in the percentage of reduction achieved over the classical sample sizes. Note that the computed Bayesian sample size does not depend on the value of $|m' - \mu|$. That is, the Bayesian sample sizes are identical in Tables 3 and 4 for a given value of delta. The confidence and accuracy of the interval estimates produced, i.e., the number of times the true mean of the sampling process is contained within the interval and the width of the interval constructed, is comparable to the results obtained using classical methods for the case where $|m' - \mu| = 5$. For the case where $|m' - \mu| = 10$, the desired confidence is not achieved until the situation involving the two largest sample sizes. The separation between the posterior and sample means

Table 3. Data for the Bayesian Approximation Model Based
on One Hundred Runs for Each Value of Delta

$$\sigma^2/\sigma'^2 = 16, |m' - \mu| = 5$$

Specified Width of Confidence Interval	Computed Classical Sample Size	Number of Times $UL_c \leq m \leq UL_c$	Average Computed Bayesian Sample Size	Number of Times $UL_b \leq m \leq UL_b$	Average Value of $ m'' - m $	Number of Times $\bar{D} \leq q\delta S$	Percent Reduction in Sample Size
δ	n^*_c		$\bar{n}^*_b$		$\bar{D}$		
1.0	18	96	6.7	97	5.589	75	61.1
0.9	22	96	8.9	97	4.814	76	59.1
0.8	27	95	12.6	98	4.355	79	51.9
0.7	34	96	18.7	98	2.541	90	44.1
0.6	46	98	32.4	96	2.175	95	28.3
0.5	64	93	52.8	93	1.185	100	17.2
0.4	99	97	87.7	94	0.756	100	11.1
0.3	174	96	163.1	98	0.470	100	5.7
0.2	387	97	376.4	98	0.217	100	2.6

Table 4. Data for the Bayesian Approximation Model Based
on One Hundred Runs for Each Value of Delta

$$\sigma^2/\sigma'^2 = 16, |m' - \mu| = 10$$

Specified Width of Confidence Interval	Computed Classical Sample Size	Number of Times $UL_c \leq m \leq UU_c$	Average Computed Bayesian Sample	Number of Times $UL_b \leq m \leq UU_b$	Average Value of $ m'' - m $	Number of Times $\bar{m} \leq q \leq S$	Percent Reduction in Sample Size
δ	n^*_c		$\bar{n}_b$		$\bar{D}$		
1.0	18	96	6.7	75	8.300	54	61.1
0.9	22	96	8.9	74	6.008	51	59.1
0.8	27	96	12.6	86	6.993	49	51.9
0.7	34	96	18.7	77	4.438	75	44.1
0.6	46	98	32.4	86	3.874	77	28.3
0.5	64	94	52.8	88	2.359	97	17.2
0.4	99	98	87.7	88	1.545	100	11.1
0.3	174	95	163.1	95	0.930	100	5.7
0.2	387	97	376.4	95	0.416	100	2.6

decreases as the sample size increases and the sampling information is given more weight in the determination of the posterior distribution. For this reason, the test suggested for determining whether or not to use the prior information does not work well at all. For both the case where $|m' - \mu| = 5$ and $|m' - \mu| = 10$, the test rejects the prior information too often for small sample sizes and erroneously allows the use of the prior information in large sample sizes. It appears that a better decision rule as to whether or not to reject the prior information should consider the difference between the prior mean (rather than the posterior mean) and the sample mean. The accuracy of the approximation procedure is quite good; the overall average reduction in the sample size for all values of delta is 12.0 samples, which equates to approximately 75 percent of the true difference between the classical and the Bayesian sample sizes, which is 16 samples for this particular case.

The second test of the procedure involved computing the Bayesian sample size required for each value of delta and for various values of $|m' - \mu|$ ranging from one standard deviation below the true mean of the sampling process to one standard deviation above this value. The specific values chosen for $|m' - \mu|$ and the results of the test are shown in Table 5. The results obtained when the value of $|m' - \mu|$ is within one-half standard deviation on either side of μ are quite favorable, with only three cases out of the total of 63 trials where the Bayesian interval estimate did not include the true value of the mean of the sampling process. Overall, there were a total of 24 cases, out of the 99 total trials, where the Bayesian interval estimate did not include

Table 5. Data for the Bayesian Approximation Model Based on Various Values of $|m' - \mu|$

$$\sigma^2/\sigma'^2 = 16$$

Specified Value of Delta	Classical Sample Size n_c	Computed Bayesian Sample Size for the Specified Value of $ m' - \mu $										
		-20	-15	-10	-5	-2	0	2	5	10	15	20
1.0	18	7*	5*	7	5	5	11	5	5	5	5*	5*
0.9	22	6*	10*	12	19	6	18	6	6	6	13	6*
0.8	27	15	7*	10	14	20	23	20	7	17*	15*	21
0.7	34	23*	22	18*	12	24	21	24	17	9	9*	19*
0.6	46	38*	39*	35	37	36	31	40	32	36	40	36*
0.5	64	51	48	52	49	52	57	58	55	54	52	53*
0.4	99	85*	87*	88	89	87	91	91	90	93	85	87
0.3	174	159	164*	159	163	162	166	164	161	166*	163	164*
0.2	387	378	379	376	375	376	376	377	377	376	376*	378

Figures marked with an astric (*) indicate cases where the confidence interval based on the Bayesian sample size did not contain the true value of the mean of the sampling process.

the true value of the mean of the sampling process.

The final test conducted on the model was to fix the value $|m' - \mu| = 5$ and to compute the Bayesian sample size required for each value of delta and for various ratios of the variances, σ^2/σ'^2 . The specific values chosen for the ratio of the variances and the results of the test are shown in Table 6. The results obtained when the ratio of the sampling and the prior variances was 4 or greater are good, with only one case out of a total of 63 trials where the Bayesian interval estimate did not include the true value of the mean of the sampling process. Overall, there were a total of seven cases out of the 99 trials where the Bayesian interval estimate did not include the true value of the mean of the sampling process.

Table 6. Data for the Bayeslan Approximation Model Based on Various Values of the Ratio of the Variances

$$|m' - \mu| = 5$$

Specified Value of the Ratio of the Variances σ^2/σ'^2	Computed Bayesian Sample Size for the Specified Value of Delta									
	1.0	0.9	0.8	0.7	0.6	0.5	0.4	0.3	0.2	
32	5	6	7	13	22	33	69*	144	368	
16	5	10	7	22	39	48	87	164	379	
8	14	16	16	29	41	58	94	167	382	
7	5	21	23	29	43	58	95	170	382	
6	14	6	16	30	44	60	95	170	384	
5	17	21	26	33	44	61	97	173	385	
4	11	19	25	34	45	64	97	173	385	
3	16	16*	25*	31	45	64*	98	173*	385	
2	17	21	27	33	45	64	99	173*	385	
1	18	22	27	33	46	64	99	174*	387	
Classical Sample Size	18	22	27	34	46	64	99	174	387	

Figures marked with an astric (*) indicate cases where the confidence interval based on the Bayesian sample size did not contain the true value of the mean of the sampling process.

CHAPTER IV

CONCLUSIONS AND RECOMMENDATIONS

Conclusions

The results of this study indicate the following conclusions.

1. The suggested procedure to approximate Bayesian sample sizes and construct interval estimates for the mean of the sampling process should be used for the normal sampling process when accurate prior information is available. That is, when the prior mean is within one-half standard deviation of the true mean of the sampling process.
2. In the worst case, the procedure will yield the same sample sizes as would classical techniques. In this case, the interval estimates should be based on the classical method, since in essence, the prior information has been rejected.
3. The accuracy and confidence levels associated with the interval estimates based on the approximation procedure are comparable to those obtained by using classical techniques if the prior information is accurate.
4. The heuristic rule suggested to determine whether or not to use the prior information did not work well because the value $|m'' - m|$ is a function of the sample size as well as being a function of the value of the prior mean, m' .
5. The results obtained in the demonstration of the procedure for the values of delta selected, indicate that the procedure to approxi-

mate the Bayesian sample size and construct interval estimates is a viable concept which has direct applicability and value in Operational Testing.

Recommendations

As in most cases involving research of a limited scope, perhaps more problems are unearthed than are resolved in this study. The limited results obtained, however, show some merit and applicability to Operational Testing. As a matter of future research in the area covered by this study, the following recommendations are suggested.

1. Further efforts are required to improve the iterative procedure used to approximate the Bayesian sample size. A refined procedure should take into account the need to treat large and small sample sizes as separate problems. Perhaps the increment added to the approximation at any specific iteration should be some function of the number of iterations already conducted. Care must be taken, however, that any procedure developed for this situation be compatible to the Operational Testing environment, where ease of application and similarity are prime objectives.
2. The sample standard deviation, S , in equation (2-13), is the only variable in the equation for a specific sample size. This particular random variable is related to the chi-square distribution. Perhaps further work with this particular element of the expression for the approximate Bayesian sample size would lead to more accurate approximations of the equation.
3. A workable decision rule for determining whether or not to

use the prior information is needed. It is suggested that the relationship between the prior and sample means, i.e., $|m' - m|$, will yield more viable results than the technique used in this study. Obviously, whatever rule is developed, it must treat the differences associated with large and small sample sizes separately.

4. There are obvious limitations in applying this procedure to Operational Testing. Although the procedure holds some potential of reducing costs associated with Operational Testing by reducing the number of replications required of a specific tests, any iterative sampling scheme is inherently difficult and costly to apply because of the problems involved with multiple scheduling and set-up costs. The procedure seems better suited to those testing situations where a large number of samples are required and the cost of sampling is relatively low. For these reasons, a scheme to incorporate the concept of loss functions into this procedure is needed before it can assume the cloak of a true decision making procedure.

APPENDIX I**FORTRAN PROGRAM FOR THE APPROXIMATE
BAYESIAN PROCEDURE**

-RFOR, IS MAIN

```
COMMON/ONE/ X(2,1000), N(3)
COMMON/TWO/ XMEAN(3), XVAR(3)
COMMON/THREE/ XHAT(3), SHAT(3)
COMMON/FIVE/ ALPHA, UL(3), UU(3)
COMMON/SEVEN/ NC(5), DELTA
COMMON/EIGHT/ NB(10), NPRIM(10), DIFF
COMMON/NINE/ WIDTH(2)
COMMON/TEN/ LOOP(2), KEY, DMAX
EXTERNAL UNIF
```

C***** READ IN BASIC PARAMETERS

```
8 FORMAT ( 1 )
  READ(5,8) ALPHA
  READ(5,8) NSU
  READ(5,8) XMEAN(1), YVAR(1)
  READ(5,8) XMEAN(2), XVAR(2)
```

C***** START UP UNIFORM GENERATOR TO RANDOMIZE STARTING POINT

```
DO 10 J=1, NSU
  Q= UNIF(A)
10 CONTINUE

DO 100 KK= 1, 30
  READ(5, 8, END= 999) DELTA
```

C***** DETERMINE THE MINIMUM CLASSICAL SAMPLE SIZE

```
CALL CLASS( N(1) )
CALL RANDN(1)
```

C***** DETERMINE THE MINIMUM BAYESIAN SAMPLE SIZE, IF APPROPRIATE

```
CALL BAYES( N(1), N(2) )
```

C***** COMPUTE CONFIDENCE INTERVALS FOR THE DATA PROCESSES

```
CALL CONFID( N(3), 3 )
CALL ORDER(1)
CALL CONFID( N(1), 1 )
```

C***** PRINT OUTPUT

```
CALL OUTPUT
```

```
100 CONTINUE
```

```
999 CONTINUE
  WRITE(6,70)
  70 FORMAT(1H1)
```

```
STOP
END
```

Copy available to DDC does not
permit fully legible reproduction

-RFOR,IS RANDN

```

C*****THIS SUBROUTINE GENERATES NORMALLY DISTRIBUTED PSEUDO-RANDOM
C      NUMBERS HAVING A SPECIFIED MEAN AND VARIANCE
C
C*****ARGUMENT DEFINITION
C      X IS THE ARRAY OF RANDOM NUMBERS (OUTPUT)
C      N IS THE NUMBER OF RANDOM NUMBERS DESIRED (INPUT)
C      XMEAN IS THE MEAN OF THE RANDOM NUMBERS (INPUT)
C      XVAR IS THE VARIANCE OF THE RANDOM NUMBERS (INPUT)
C
C*****THIS SUBROUTINE USES THE BOX AND MUELLER METHOD FOR
C      GENERATION OF NORMAL PSEUDO-RANDOM NUMBERS

```

SUBROUTINE RANDN(J)

```

COMMON/ONE/ X(2,1000), N(3)
COMMON/TWO/ XMEAN(3), XVAR(3)

```

EXTERNAL UNIF

```

TPI=6.2831852
DO 100 I=1, N(J), 2
  A= UNIF(1)
  B= UNIF(2)
  X(1,I)= XMEAN(2)+ SQRT(-2.0*XVAR(2)*ALOG(A))*COS(TPI*B)
  X(2,I)= X(1,I)
  II= I+ 1
  X(1,II)= XMEAN(2)+ SQRT(-2.0*XVAR(2)*ALOG(A))*SIN(TPI*B)
  X(2,II)= X(1,II)
100 CONTINUE

  RETURN
  FND

```

-RFOR,IS UNIF

```

FUNCTION UNIF(A)
DATA IY/96581/
IY=IY*3125
IF(IY)5,6,6
5 IY=IY+1+34359738367
6 YFL=IY
UNIF= YFL*2.0**(-35)

RETURN
END

```


-RFOR, IS ORDER

```

C***** THIS SUBROUTINE SORTS A GIVEN SET OF DATA FROM THE LOWEST
C        VALUE TO THE HIGHEST, AND COMPUTES THE SAMPLE STATISTICS
C        (MEAN AND STANDARD DEVIATION) OF THE DATA PROCESS
C
C***** ARGUMENT DEFINITION
C        X= THE ARRAY OF DATA VALUES TO BE SORTED (INPUT/OUTPUT)
C        N= THE NUMBER OF DATA POINTS (INPUT)
C        XHAT= THE SAMPLE MEAN OF THE DATA PROCESS (OUTPUT)
C        SHAT= THE SAMPLE STANDARD DEVIATION OF THE DATA
C        PROCESS (OUTPUT)

```

SUBROUTINE ORDER(K)

```

COMMON/ONE/ X(2,1000), N(3)
COMMON/THREE/ XHAT(3), SHAT(3)

```

```

NM1= N(K)- 1
DO 200 I=1, NM1
  IP1=I+1
  DO 100 J= IP1, N(K)
    IF( X(K,I) .LE. X(K,J) ) GO TO 100
    TEMP= X(K,I)
    X(K,I)= X(K,J)
    X(K,J)= TEMP
  100 CONTINUE
200 CONTINUE

```

```

C***** COMPUTE THE SAMPLE STATISTICS FOR THE DATA PROCESS
SUM1=0.0
SUM2=0.0
DO 300 I=1, N(K)
  SUM1= SUM1+ X(K,I)
  SUM2= SUM2+ X(K,I)**2
300 CONTINUE

```

```

YN= N(K)
XHAT(K)= SUM1/YN
RN= YN- 1.0
SUM22= SUM2- (SUM1**2)/YN
SHAT(K)= SORT(SUM22/RN)

```

```

8    RETURN
END

```

-RFOR.15 CLASS

```

C***** THIS SUBROUTINE CALCULATES THE MINIMUM CLASSICAL SAMPLE SIZE
C        REQUIRED TO CONSTRUCT A CONFIDENCE INTERVAL OF GIVEN WIDTH
C        ABOUT THE MEAN OF A NORMAL SAMPLING POPULATION OF UNKNOWN
C        VARIANCE
C***** ARGUMENT DEFINITION
C        ALPHA= THE CONFIDENCE COEFFICIENT (INPUT)
C        DELTA= A FUNCTION OF THE INTERVAL WIDTH (INPUT)
C        NCLASS= THE COMPUTED SAMPLE SIZE (OUTPUT)

```

SUBROUTINE CLASS(NCLASS)

```

COMMON/FIVE/ ALPHA, UL(3), UU(3)
COMMON/SEVEN/ NC(5), DELTA
COMMON/TEN/ LOOP(2), KEY, DMAX

```

```

C***** COMPUTE THE FIRST APPROXIMATION OF THE CLASSICAL SAMPLE
C        SIZE, NC(1), BASED ON THE STANDARD NORMAL DISTRIBUTION, WHICH
C        IS IDENTICAL TO THE T DISTRIBUTION WITH INFINITE DEGREES
C        OF FREEDOM

```

```

ALPHA1= ALPHA/2.0
S= TINORM(ALPHA1, $15)
GO TO 18

```

```

15 WRITE(6,17)
17 FORMAT(//, 10X, 68H ERROR MESSAGE--OVERFLOW ON INVERSE NORMAL DIST
IRIBUTION--FORMAT 15 )
CALL EXIT

```

```

18 CONTINUE
REALC= (2.0*S/DELTA)**2
NC(1)= INT(REALC)
IF( NC(1) .LT. REALC ) NC(1)= NC(1)+ 1

```

```

C***** COMPUTE THE SUCCEEDING APPROXIMATIONS OF THE CLASSICAL SAMPLE
C        SIZE, NC(J), BASED ON THE T DISTRIBUTION WITH DEGREES OF
C        FREEDOM EQUAL TO NC(J-1)- 1. STOP THE ITERATIVE PROCEDURE
C        WHEN NC(J) IS EQUAL TO NC(J-1)

```

```

DO 30 J=2, 10
NDF= NC(J-1)- 1
T= STUDIN(ALPHA, NDF, $21)
GO TO 24

```

```

21 WRITE(6,23)
23 FORMAT(//, 10X, 74H ERROR MESSAGE--OVERFLOW ON STUDENTS T DISTRIBU
TION FUNCTION--FORMAT 21 )
CALL EXIT

```

```

24 CONTINUE
REALC= (2.0*T/DELTA)**2
NC(J)= INT(REALC)
IF( NC(J) .LT. REALC ) NC(J)= NC(J)+ 1
IF ( NC(J) .EQ. NC(J-1) ) GO TO 35

```

8 30 CONTINUE

C***** ASSIGN THE COMPUTED SAMPLE SIZE NCLASS
35 NCLASS= NC(J)
LOOP(1)= J
RETURN
END

```

-RFOR, IS BAYES
C   THIS SUBROUTINE CALCULATES THE MINIMUM BAYESIAN SAMPLE SIZE,
C   IF APPROPRIATE, TO CONSTRUCT A CONFIDENCE INTERVAL OF GIVEN
C   WIDTH ABOUT THE MEAN OF A NORMAL SAMPLING POPULATION
C   WITH UNKNOWN VARIANCE
C
C***** ARGUMENT DEFINITION
C       K= THE MINIMUM CLASSICAL SAMPLE SIZE (INPUT)
C       NBAYES= THE COMPUTED SAMPLE SIZE (OUTPUT)
C
SUBROUTINE BAYES ( K, NBAYES )

COMMON/ONE/ X(2,1000), N(3)
COMMON/TWO/ XMEAN(3), XVAR(3)
COMMON/THREE/ XHAT(3), SHAT(3)
COMMON/SEVEN/ NC(3), DELTA
COMMON/EIGHT/ NB(10), NPRIM(10), DIFF
COMMON/TEN/ LOOP(2), KEY, DMAX

C***** COMPUTE THE FIRST APPROXIMATION, N(1), OF THE BAYESIAN
C   SAMPLE SIZE
REALB= FLOAT(K)/4.0
NB(1)= INT( REALB )
IF( NB(1) .LT. REALB ) NB(1)= NB(1)+ 1
N(2)= NB(1)

C***** TAKE N(1) SAMPLES AND COMPUTE THE SAMPLE STATISTICS FOR THE
C   DATA PROCESS AND THE POSTERIOR PARAMETERS BASED ON THESE
C   N(1) OBSERVATIONS
CALL ORDER(2)
APPN= SHAT(2)**2/XVAR(1)
NPRIM(1)= INT( APPN )
IF( NPRIM(1) .LT. APPN ) NPRIM(1)= NPRIM(1)+ 1
XMEAN(3)= ( NPRIM(1)*XMEAN(1)+ NB(1)*XHAT(2) ) /
1 FLOAT( NPRIM(1)+ NB(1) )
DIFF= ABS( XMEAN(3)- XHAT(2) )
DMAX= DIFF
KEY= 1
IF( N(2) .GE. K- NPRIM(1) ) GO TO 55

C***** COMPUTE THE SUCCEEDING APPROXIMATIONS, N(J), OF THE BAYESIAN
C   SAMPLE SIZE. STOP THE ITERATIVE PROCEDURE WHEN N(J) IS
C   GREATER THAN OR EQUAL TO K- NPRIM(J)
DO 100 J=2, 20
RINC= FLOAT( K- NPRIM(J-1) )/4.0
INC= INT( RINC )
IF( INC .LT. RINC ) INC= INC+ 1
NB(J)= NB(J-1)+ INC
N(2)= NB(J)

C***** TAKE EACH SUCCEEDING N(J) SAMPLES AND COMPUTE THE SAMPLE
C   STATISTICS FOR THE DATA PROCESS AND THE POSTERIOR PARAMETERS
C   BASED ON THESE N(J) OBSERVATIONS
CALL ORDER(2)

```

```

      APPN= SHAT(2)**2/XVAR(1)
      NPRIM(J)= INT( APPN )
      IF( NPRIM(J) .LT. APPN ) NPRIM(J)= NPRIM(J)+ 1
      XMEAN(3)= ( NPRIM(J)*XMEAN(1)+ NB(J)*XHAT(2) )/
1     FLOAT( NPRIM(J)+ NB(J) )
      DIFF= ABS( XMEAN(3)- XHAT(2) )
      IF( DIFF .LE. DMAX ) GO TO 35
      DMAX= DIFF
      KEY= J
35  CONTINUE
      IF( N(2) .GE. K- NPRIM(J) ) GO TO 45

100 CONTINUE

C*****  ASSIGN THE SAMPLE SIZE COMPUTED ABOVE TO NBAYES AND DETERMINE
C        THE POSTERIOR (POOLED) SAMPLE SIZE
55  CONTINUE
      NBAYES= N(2)
      IF( NBAYES .GT. K ) NBAYES= K
      N(3)= NBAYES+ NPRIM(1)
      XHAT(3)= XMEAN(3)
      SHAT(3)= SHAT(2)
      LOOP(2)= 1
      GO TO 999

45  CONTINUE
      NBAYES= NB(J)
      IF( NBAYES .GT. K ) NBAYES= K
      N(3)= NBAYES+ NPRIM(J)
      XHAT(3)= XMEAN(3)
      SHAT(3)= SHAT(2)
      LOOP(2)= J

999 RETURN
      END

```

```

      8
      -RFOR,IS CONFID

C***** THIS SUBROUTINE CALCULATES A CONFIDENCE INTERVAL FOR THE MEAN
C      OF A NORMAL POPULATION WHEN THE VARIANCE IS UNKNOWN
C
C***** ARGUMENT DEFINITION
C      N= THE NUMBER OF DATA POINTS IN THE SAMPLE (INPUT)
C      ALPHA= THE CONFIDENCE COEFFICIENT (INPUT)
C      XHAT= THE SAMPLE MEAN OF THE DATA PROCESS (INPUT)
C      SHAT= THE SAMPLE STANDARD DEVIATION OF THE DATA
C      PROCESS (INPUT)
C      UL= THE LOWER CONFIDENCE LIMIT FOR THE MEAN (OUTPUT)
C      UU= THE UPPER CONFIDENCE LIMIT FOR THE MEAN (OUTPUT)

      SUBROUTINE CONFID(N, J)
      COMMON/THREE/ XHAT(3), SHAT(3)
      COMMON/FIVE/ ALPHA, UL(3), UU(3)

C***** COMPUTE THE DEGREES OF FREEDOM ASSOCIATED WITH THE SAMPLE
      NDF= N-1

C***** DETERMINE THE VALUE OF THE STUDENT(S T DISTRIBUTION AT A
C      SIGNIFICANCE LEVEL = ALPHA
C      NOTE--THIS OPERATION USES A STAT*PACT FUNCTION CALLED STUDIN
C      TO CALCULATE THE INVERSE STUDENTS T VALUE GIVEN THE
C      CONFIDENCE COEFFICIENT ALPHA

      T= STUDIN(ALPHA, NDF, $10)
      GO TO 700

10 WRITE(6,15)
15 FORMAT(//,10X, 74H ERROR MESSAGE--OVERFLOW ON STUDENT(S T DISTRIBUTION
      FUNCTION--FORMAT 700          )

      CALL EXIT

700 CONTINUE

      YN=N

C***** COMPUTE THE LOWER CONFIDENCE LIMIT
      UL(J)= XHAT(J)- T*(SHAT(J)/SQRT(YN))

C***** COMPUTE THE UPPER CONFIDENCE LIMIT
      UU(J)= XHAT(J)+ T*(SHAT(J)/SQRT(YN))

      RETURN
      END

-RMAP
LIB SYSTEMS*MATHSTAT.
-XQT

```

-RFOR.IS OUTPUT

SUBROUTINE OUTPUT

```
COMMON/ONE/ X(2,1000), N(3)
COMMON/TWO/ XMEAN(3), XVAR(3)
COMMON/THREE/ XHAT(3), SHAT(3)
COMMON/FIVE/ ALPHA, UL(3), UU(3)
COMMON/SEVEN/ NC(5), DELTA
COMMON/EIGHT/ NB(10), NPRIM(10), DIFF
COMMON/NINE/ WIDTH(2)
COMMON/TEN/ LOOP(2), KEY, DMAX
```

C***** PRINT HEADINGS FOR PRINTED OUTPUT

```
DO 100 J=1, 2
WRITE(6,15)
15 FORMAT(1H1)
IF( J .EQ. 2 ) GO TO 40
WRITE(6,35)
35 FORMAT(///, 40X, 45H DATA VALUES USED IN THE CLASSICAL ANALYSIS )
GO TO 50

40 WRITE(6,45)
45 FORMAT(///, 40X, 45H DATA VALUES USED IN THE BAYESIAN ANALYSIS )
```

C***** PRINT BASIC PARAMETERS ASSOCIATED WITH EACH DATA PROCESS

```
50 CONTINUE
WRITE(6,52) N(J)
52 FORMAT(///, 10X, 26H NUMBER OF OBSERVATIONS = , I3)
WRITE(6,54) XMEAN(2)
54 FORMAT(10X, 19H LIKELIHOOD MEAN = , F8.3)
8 IF( J .EQ. 2 ) WRITE(6,56) XMEAN(1)
56 FORMAT(1H+, T82, 14H PRIOR MEAN = , F8.3)
WRITE(6,58) XVAR(2)
58 FORMAT(10X, 23H LIKELIHOOD VARIANCE = , F8.3)
IF( J .EQ. 2 ) WRITE(6,60) XVAR(1)
60 FORMAT(1H+, T82, 18H PRIOR VARIANCE = , F8.3)
WRITE(6,62) DELTA
62 FORMAT(10X, 9H DELTA = , F4.2)
```

C***** PRINT THE DATA VALUES GENERATED BY RANDN

```
WRITE(6,65) ( X(J,I), I=1, N(J) )
65 FORMAT(///, 10(3X, F8.3) )
```

C***** PRINT THE SAMPLE STATISTICS OF THE DATA PROCESS

```
IF( J .EQ. 2 ) WRITE(6,72) XHAT(3)
72 FORMAT(///, 10X, 47H THE MEAN OF THE POSTERIOR DISTRIBUTION, M-- =
1 , F10.5)
WRITE(6,75) XHAT(J), SHAT(J)
75 FORMAT(///, 10X, 45H THE SAMPLE MEAN OF THE DATA PROCESS, XHAT = ,
1 F10.5, ///, 10X, 52H THE SAMPLE STANDARD DEVIATION OF THE DATA P
ROCESS, SHAT = , F10.5 )
```

```

C***** PRINT THE (1-ALPHA) CONFIDENCE INTERVAL ASSOCIATED WITH EACH
C PROCESS
      WIDTH(J)= DFLTA*SHAT(J)
      WRITE(6,85) WIDTH(J)
85  FORMAT(//, 10X, 52H THE DESIRED WIDTH OF THE CONFIDENCE INTERVAL I
      IS =      , F6.2 )
      UL(2)= UL(3)
      UU(2)= UU(3)
      WRITE(6,95) ALPHA, UL(J), UU(J)
95  FORMAT(//, 10X, 47H THE (1-ALPHA) CONFIDENCE INTERVAL FOR THE MEAN
      1 . /, 10X, 38H WITH CONFIDENCE COEFFICIENT, ALPHA =      , F4.3,
      2 8H, IS = (      , F8.3, 2H,      , F8.3, 1H) )

      IF( J .EQ. 1) GO TO 97

      WRITE(6,98) DIFF
98  FORMAT(//, 10X, 71H THE ABSOLUTE DIFFERENCE BETWEEN THE POSTERIOR
      1AND SAMPLE MEANS, DIFF =      , F6.3)
      WRITE(6,99) DMAX, KEY
99  FORMAT(//, 10X, 7H DMAX =      , F6.3, 10H AT LOOP =      , 12)

97  CONTINUE
      WRITE(6,96) LOOP(J)
8  96  FORMAT(//, 10X, 9H LOOPS =      , 12)

100 CONTINUE

      RETURN
      END

```


APPENDIX II

FORTRAN PROGRAM FOR THE CHI-SQUARE

TEST OF NORMALITY

-RFOR.IN CHISO

```

C***** THIS SUBROUTINE TAKES A SET OF ORDERED DATA (ARRANGED FROM
C THE LOWEST TO THE HIGHEST VALUE) AND
C (1) ESTABLISHES K EQUAL-PROBABILITY CELLS, WHERE K DEPENDS ON
C THE SAMPLE SIZE, I.E., K= 20 FOR N .GE. 100, K= 10 FOR N .GE.
C 50 .AND. .LT. 100, AND K= 5 FOR N .LT. 50
C (2) PERFORMS A CHI-SQUARE GOODNESS-OF-FIT TEST FOR NORMALITY
C ON THE DATA SAMPLE AND DETERMINES THE SIGNIFICANCE LEVEL
C AT WHICH WE CAN ASSUME THAT THE DATA SAMPLE IS IN FACT
C REPRESENTATIVE OF A NORMAL PROCESS
C***** ARGUMENT DEFINITION
C X= THE ARRAY OF DATA VALUES TO BE TESTED (INPUT)
C N= THE NUMBER OF DATA POINTS (INPUT)
C K= THE NUMBER OF CELLS INTO WHICH THE DATA IS DIVIDED (INPUT)
C XHAT= THE SAMPLE MEAN OF THE DATA PROCESS
C SHAT= THE SAMPLE STANDARD DEVIATION OF THE DATA PROCESS
C CHIS= THE CHI-SQUARE STATISTIC COMPUTED FROM
C THE DATA (OUTPUT)
C SIGL= THE SIGNIFICANCE LEVEL OF THE TEST (OUTPUT)

```

SUBROUTINE CHISO

```

COMMON/ONE/ X(500), N
COMMON/TWO/ XMEAN, XVAR
COMMON/FOUR/ K, KLESS1, CHIS, SIGL
COMMON/SIX/ CBSTRD(19), CBNORM(19), KOUNT(20)
DIMENSION ALPHA(19)

```

```

C***** SET ALL CELL COUNTERS TO ZERO
DO 5 I=1, K
  KOUNT(I)= 0
5 CONTINUE

```

```

C***** COMPUTE THE CELL-BREAK POINTS FOR THE GIVEN DATA
C NOTE--THIS OPERATION USES A STAT-PACK FUNCTION CALLED TINORM
C TO COMPUTE THE VALUE OF THE INVERSE OF THE NORMAL (0,1) DISTR.
IF(K-1)10, 20, 30
10 DO 15 I=1, KLESS1
  ALPHA(I)= 0.2*I
15 CONTINUE
  GO TO 50
20 DO 25 I=1, KLESS1
  ALPHA(I)= 0.1*I
25 CONTINUE
  GO TO 50
30 DO 35 I=1, KLESS1
  ALPHA(I)= 0.05*I
35 CONTINUE
50 DO 100 I=1, KLESS1
  CBSTRD(I)= TINORM(ALPHA(I), 370)
  CBNORM(I)= CBSTRD(I)*SQRT( XVAR )+ XMEAN
  GO TO 100

70 WRITE(6,20)
20 FORMAT(//,10X,68H ERROR MESSAGE--INVERTED IN PROBABILITIES IN ALL
  11RIBUTION--FORMAT 100 1
10 CONTINUE

```

```

C***** COUNT THE NUMBER OF OBSERVATIONS FALLING IN EACH CELL
DO 300 I=1, N
  DO 200 J=1, KLESS1
    IF( X(I) .GT. CBNORM(J) ) GO TO 190
    KOUNT(J) = KOUNT(J) +1
    GO TO 300
  190 IF(J .EQ. KLESS1) KOUNT(K)= KOUNT(K)+1
  200 CONTINUE
  300 CONTINUE

C***** COMPUTE THE CHI-SQUARE STATISTIC, CHIS
  CH1=0.0
  RN= FLOAT( N ) / FLOAT( K )
  DO 500 I=1, K
    CH1= CH1+(KOUNT(I)-RN)**2
  500 CONTINUE
  CHIS= CH1/RN

C***** DETERMINE THE SIGNIFICANCE LEVEL OF THE TEST
C      NOTE--THIS OPERATION USES A STAT-PACK FUNCTION CALLED CHI TO
C      DETERMINE THE CHI-SQUARE DISTRIBUTION GIVEN THE POINT AND
C      THE DEGREES OF FREEDOM
  NDF= K-3
  CUMD= CHI( CHIS, NDF, .6600 )
  SIGL= 1.0- CUMD
  GO TO 600

600 WRITE(5,610) CHIS
610 FORMAT(//,10X, 74H ERROR MESSAGE--OVERFLOW ON CHI-SQUARE DISTRIBUT
      ION FUNCTION--FORMAT 600 .24H CHI-SQUARE STATISTIC = , F5.2)

      CONTINUE

620 RETURN
      END

```

BIBLIOGRAPHY

1. Anscombe, F. J., "Bayesian Statistics," The American Statistician, Vol. 15, 1961, pp. 21-24.
2. Army Regulation No. 10-4, "U. S. Army Operational Test and Evaluation Agency," Headquarters, Department of the Army, Washington, D. C., 1974.
3. Army Regulation No. 71-3, "User Testing," Headquarters, Department of the Army, Washington, D. C., 1974.
4. Army Regulation No. 1000-1. "Basic Policies for Systems Acquisition by the Department of the Army," Headquarters, Department of the Army, Washington, D. C., 1971.
5. Atzinger, E. M., and Brooks, W. J., Comparison of Bayesian and Classical Analysis for a Class of Decision Problems, Technical Report No. 59, April 1972, U. S. Army Materiel Systems Analysis Agency, Aberdeen Proving Ground, Maryland.
6. Atzinger, E. M., Brooks, W. J., et al., Compendium on Risk Analysis Techniques, Special Publication No. 4, July 1972, U. S. Army Systems Analysis Agency, Aberdeen Proving Ground, Maryland.
7. Box, G. E. P., and Muller, M. E., "A Note on the Generation of Random Normal Deviates," Annals of Mathematical Statistics, Vol. 29, 1958, pp. 21-24.
8. Gilbreath, S. G., A Bayesian Procedure for the Design of Sequential Sampling Plans, Unpublished Doctoral Dissertation, Georgia Institute of Technology, Atlanta, Georgia, 1966.
9. Hines, W. W., and Montgomery, D. C., Probability and Statistics in Engineering and Management Science, New York: The Ronald Press Company, 1972.
10. Hoel, P. G., Port, S. C., and Stone, C. J., Introduction to Statistical Theory, Boston: Houghton Mifflin Company, 1971.
11. OTEA, "MICV Firing Port Weapon Operational Test I (MICV FPW OT I) Test Design Plan," U. S. Army Operational Test and Evaluation Agency, Ft. Belvoir, Virginia, September 1973.
12. Mace, A. E., Sample Size Determination, New York: Reinhold Publishing Corp., 1964.

13. Raiffa, H., and Schlaifer, R., Applied Statistical Decision Theory, Boston: The MIT Press, 1961.
14. UNIVAC, "Large Scale Systems STAT-PACK and MATH-PACK Programmer's Reference," New York: Sperry Rand Corp., 1st Revision, 1970.
15. White, L. R., Bayesian Reliability Assessment for Systems Program Decisions, Unpublished Masters Dissertation, Air Force Institute of Technology, Wright-Patterson AFB, Ohio, 1971.
16. Winkler, R. L., Introduction to Bayesian Inference and Decision, New York: Holt, Rinehart and Winston, Inc., 1972.