

Author(s)	Molsberry, Dale M.; Connolly, James P.; Barry, John A.
Title	An algorithm for generating random time delays in manual war games
Publisher	Monterey, California: U.S. Naval Postgraduate School
Issue Date	1963
URL	http://hdl.handle.net/10945/11573

This document was downloaded on May 12, 2015 at 03:21:58



<http://www.nps.edu/library>

Calhoun is a project of the Dudley Knox Library at NPS, furthering the precepts and goals of open government and government transparency. All information contained herein has been approved for release by the NPS Public Affairs Officer.

**Dudley Knox Library / Naval Postgraduate School
411 Dyer Road / 1 University Circle
Monterey, California USA 93943**



<http://www.nps.edu/>

NPS ARCHIVE
1963
MOLSBERRY, D.

AN ALGORITHM FOR GENERATING RANDOM
TIME DELAYS IN MANUAL WAR GAMES

DALE M. MOLSBERRY, JAMES P. CONNOLLY, II,
AND JOHN A. BARRY.

LIBRARY
U.S. NAVAL POSTGRADUATE SCHOOL
MONTEREY, CALIFORNIA

DUDLEY KNOX LIBRARY
NAVAL POSTGRADUATE SCHOOL
MONTEREY CA 93943-5101

AN ALGORITHM FOR
GENERATING RANDOM TIME DELAYS
IN MANUAL WAR GAMES

* * * * *

Dale M. Molsberry
James P. Connolly II
John A. Barry

AN ALGORITHM FOR
GENERATING RANDOM TIME DELAYS
IN MANUAL WAR GAMES

by

Dale M. Molsberry

Major, United States Marine Corps

and

James P. Connolly II

Major, United States Marine Corps

and

John A. Barry

Captain, United States Marine Corps

Submitted in partial fulfillment of
the requirements for the degree of

MASTER OF SCIENCE

United States Naval Postgraduate School
Monterey, California

1 9 6 3

AN ALGORITHM FOR
GENERATING RANDOM TIME DELAYS
IN MANUAL WAR GAMES

by

Dale M. Molsberry

and

James P. Connolly II

and

John A. Barry

This work is accepted as fulfilling
the thesis requirements for the degree of

MASTER OF SCIENCE

from the

United States Naval Postgraduate School

ABSTRACT

A method of deriving a probability distribution function to approximate the distribution of times required for the conduct of military operations is developed. A procedure for collecting the necessary data is suggested, and a detailed description of an algorithm for manual generation of random times is given. Thus the capability of introducing random times to simulate time delays in manual war games is provided.

PREFACE

War games are currently used to a wide extent to simulate the interactions of opposing forces in the conflict situation. The validity of conclusions based upon the results of such simulation depends directly upon the quality of the inputs to the games. One important input is the time required by forces to accomplish certain missions or conduct specific operations. Many games presently utilize a method of fixed time delays which remain constant for specific operations throughout the play of the game. Depending upon the objectives of the war gaming effort, this may well be adequate. Under certain conditions, however, it is desirable, if not mandatory, to utilize some method of random time delays.

Where an attempt is made to follow the latter course, the current trend is to introduce times generated from either a Uniform Probability Distribution or a Triangular Probability Distribution. The use of these distributions is quite arbitrary, based upon ease of computation rather than degree of approximation to reality. In effect, they are used only to inject an element of randomness rather than to simulate the random manner in which real times vary.

In introducing random time delays to a manual war game, there are two basic tasks to be accomplished. One must first

derive some probability distribution which varies in a manner similar to the manner in which the real times being simulated vary. Having done so, one must then devise some method to generate random numbers from this distribution, which will be used to simulate random times.

This thesis attempts to present a method of accomplishing these two tasks which offers an advantage over the methods presently in use. A procedure is suggested, where in the absence of recorded or experimental results, professional estimates of the required times may be used to determine an approximating probability distribution. A method of generating random times from this distribution is presented which represents a compromise between ease of computation and degree of simulation.

Chapter I considers the overall problem of determining the specific probability distribution which best describes a real life event, and describes the difficulties inherent in solving this problem.

Chapter II describes a method of collecting certain minimal data about the distribution of times required for the conduct of military operations. It further demonstrates how this data may be used to develop a distribution function which approximates the distribution of the real times.

Chapter III presents an algorithm for generating random time delays based upon the approximation resulting from Chapter II.

The authors wish to express their appreciation for guidance in the preparation of this paper to Harold J. Larson, Assistant Professor of Mathematics, and Jack R. Borsting, Associate Professor of Mathematics, Department of Mathematics and Mechanics, U. S. Naval Postgraduate School; for encouragement, to Lieutenant Colonel Jack W. Philips, USMC, LFWGG Marine Corps Schools, Quantico, Virginia; and for clerical assistance to Mrs. Susan Feuerman. In addition, the authors wish to acknowledge the cooperation of the members of the staffs of the Stanford Research Institute and The Combat Development and Experimentation Center at Fort Ord, California, for generously allowing continued access to their files and technical library.

TABLE OF CONTENTS

CHAPTER		PAGE
	INTRODUCTION	1
	Background - Fixed Times - Random Times - Arbitrary Times - Requirement for Randomness	
I	GENERATION OF RANDOM TIMES	11
	Generation of Time Delays - Random Variables - Density Function - Distribution Function - Curve Fitting - Pearsonian System - Maximum Likelihood - Chi-square Test, Kolmogorov-Smirnov Test	
II	AN ALTERNATIVE USING MINIMAL DATA	26
	Minimal Data - Use of Subjective Estimates - Seven Density Functions Expressed in Terms of Minimal Data - Selection of an Approximating Distribution Function	
III	AN ALGORITHM TO GENERATE RANDOM TIMES	70
	General Description - Instructions for Use - Alternate Method of Use - Construction of C.P.D.F. Graphs - Extensions of the Algorithm	
	BIBLIOGRAPHY	88
	APPENDIX 1	90
	The Air Battle Analyzer	
	APPENDIX 2	97
	Sample Questionnaire	

INTRODUCTION

SECTION 1

War gaming is an attempt to build a workable and meaningful model of the interactions between opposing forces in a conflict situation, and in this paper should be distinguished from war exercises in which troops are actually deployed in the field. Whether on a map, computer, or playing board, war gaming is an abstraction or idealization of reality and is in no sense an attempt to duplicate reality, which attempt would require the philosopher's "perfect" knowledge of the universe. The level of this abstraction or idealization has varied throughout the history of war gaming from the highly idealized games of war chess (circa 1644) through Prussian rigid Kriegsspiel (circa 1865) with its complicated and elaborate rules and tables and the various forms of free Kriegsspiel (circa 1875) with their heavy reliance on control and director groups, free play and flexibility.(1)

Historically, rigid Kriegsspiel was an attempt to move the war game out of the parlor and closer to real operations by using maps, actual rates of fire, actual movement times and a set of rules defining how assessments were to be made. In its later development, however, it suffered from this same elaborate structure of rule books, tables and

formulas necessary for the assessment of play, so that play lost a great deal of flexibility and consumed considerable time. As a reaction against this complication, free Kriegsspiel was developed in which the assessments of play and guidance for the game were made by a director or control group staff, materially reducing the necessity for elaborate formal rules but placing heavy burdens on the operation and method of selection of the control group.

Modern war gaming preserves this distinction between "rigid" and "free" war games quite well. Until recently, most war gaming was of the "free" play type mainly because of its relative ease of play, its flexibility and the ability of the control group to determine the depth of detail to be investigated. With the development of high-speed digital computers to handle the extensive tables and rules, however, "rigid" type games have again come into their own. So the modern manual war game (say, the CPX) is a direct descendant from rigid Kriegsspiel.

As there is no such thing as the "right" type of war game but only the "most useful" for a particular purpose, these two extreme forms of the war game are often combined in varying degrees and as a group are called composite war games. The games which are more formal and rigid are particularly well suited for investigations of operations in which a great deal of detailed information is available and

the interrelationships of the various components of the game are well understood. These games can often provide useful information for current planning, research and evaluation. The freer type games are more useful for training and investigation of the gross effects of changing battle conditions on an operation.

SECTION 2

War gaming in any form may be characterized as a time varying process in which decisions must be made in the course of a conflict of some type by the opposite players. The variables which describe the game (the state variables) change in a framework of time called game time, which may be in discrete increments (every half hour, day or week, etc.) or may be continuously varying. The decisions which are made within the framework of the rules of a particular game have an effect on the forces of each side which would be instantaneous, if some time delay were not provided for in the rules of the game, and if the model is to have any foundation in a reality which is determined by time and space, so must the model reflect action in time.

In a game, these time delays for the occurrence of an event may be assessed against a side in three general ways. First, for actions of a particular type in a particular situation a fixed time delay, always the same, may be assessed. This time may be multiplied by a degradation factor for weather, light, climate, etc., but the base time will remain the same. Second, for actions of a particular type in a particular situation, a random time delay may be assessed. This time delay will be determined from a known or assumed probability density function which ideally approximates the probability density function of the real event. Thus we avoid the pitfalls of having an always perfectly functioning military machine

operating in the game and, while more difficult to handle than fixed delays, the extra effort may be justified by the type investigation that the game is intended to make. Third, arbitrary time delays may be assessed which arise out of factors outside the game itself. With no known or assumed underlying probability density function, they are products of expedience (to speed up closure of forces, for instance).

Fixed time delays are constant delays in game time assessed against a particular event. These delays are in general an aggregate of specific time delays set down in a rule book and held constant throughout the game.

Fixed time delays have some distinct advantages. In a mechanical game they are easy to program and require a minimum amount of computer core memory in the game iterations.

In manual war gaming the use of fixed time delays greatly simplifies the already complicated assessment process, and materially aides in reducing the time required for a play of the game.

Fixed time delays can frequently be obtained from existing manuals which give a time for various events. In cases where there is no published estimate for the time to complete an event the assessor group or the programmer must establish these times prior to the play of the game.

There are obvious disadvantages to the use of fixed time delays. In war gaming, fixed time delays allow the thinking antagonist an advantage in coordination which he

would seldom be afforded in a real life situation. As we will subsequently demonstrate, there are instances in which the use of fixed time delays can well raise questions as to the validity of conclusions arrived at through simulations using this type of delay.

Arbitrary time delays are adjustments in game time to maintain the flow of the game. There are obviously situations in which rigorous application of fixed time delays in game time may produce impractical delays in the real time progress of a specific game while yielding no contribution to the basic objectives of the game. In these situations the control group or programmer may assess arbitrary time delays; or no time delay at all, for some event, or events, in order to maintain the flow of the simulation at a practical rate. The advantage and necessity of this procedure is undeniable; however, it must be recognized that such action is a reorientation of the game and as such must be considered as a subsequent phase of the original game.

Random time delays are random variables obeying an assumed probability law used to better approximate real time in a simulated environment. The use of random time delays will in many cases lend validity to a war game in that it denies the undue advantage to the planner of exact coordination of rendezvous times, fuel and ammunition requirements as functions of time etc.

The use of random time delays overcomes many of the objections to the use of fixed time delays in simulation, however, compared to the fixed or arbitrary assessment of time delays the method of random delays is slow and cumbersome.

Since there are few instantaneous reactions possible on the part of a military commander, time delays to reflect normal reaction times are very important in war gaming simulations. These time delays are generated in response to specific game actions and vary in length from minutes to hours depending on the nature of the event being simulated. One can easily imagine real life events in which a few minutes difference in the time required to perform a specific task can produce a difference of magnitude in the outcome of an event. In a like manner time delays assessed against planned actions in a war game can produce marked differences in results.

As an example, let us consider a simple war game, the "Air Battle Analyzer." (Appendix (1))

The "Air Battle Analyzer" is a scenario type, free-play war game in which each opponent is given a specific capability and plan of action. It attempts to provide, and display, a readily accessible means for recording and displaying chronologically the principle movements and operations of the different surface and airborne units involved in a battle. This display points up the interactions between

different units, in particular their ordering in time.

Let us look at the description of the First Intercept Range nomograph, enclosure (1) to Appendix (1). It will be noted that a delay of 3 minutes was arbitrarily selected as the "delay for decision," etc. A glance at the Intercept Range Versus Detection Range nomograph shows that it is scaled for delays from zero to four minutes.

Using this nomograph, let us look at the intercept ranges using delays of $1/2$ minute and 4 minutes rather than the delay of three minutes assumed in the example. We note that an assumed delay of $1/2$ minute, a detection range of 150 miles and a target speed of Mach 1.2 yields a detection range of 95 miles while an assumed delay of 4 minutes holding the detection range and target speed parameters fixed at 150 miles and Mach 1.2 respectively yields a range of first intercept of 65 miles.

Using the Fire Power Analyzer with first intercept ranges of 95 miles and 65 miles we obtain the expected number of intercepts for 95 miles as 6.9, and for 65 miles, 5.7. Entering the target speed scaling nomograph with these expected values and assuming a kill probability of .7 for our missile while again using a target speed of Mach 1.2, we ultimately obtain 6.1 as our expected number of kills for a 95 mile intercept range and 3.9 as the expected number of kills for a 65 mile range of first intercept.

Here one should note that the Analyzer missile system is defined as a one launcher, one director system. One can

hardly assume that the expected number of kills for a single missile system can be multiplied by the number of launchers in the entire defense to estimate the grand total of expected kills; this would assume independence which does not exist in reality. As an illustration, let us assume that we are simulating 10 missile launchers defending a perimeter against a large wave of attacking aircraft. Further, let us assume that at least $1/2$ of the missiles launched have a target which can be engaged by no other missile. This is at least the equivalent of 5 independent launchers. Then, 5 (2.2 expected kills each) implies a difference of at least 10 attackers destroyed depending only on the extremes in time delays that could be assessed. Were random time delays used, discrete times (say fifteen seconds apart) could be randomly selected according to a probability law approximating the true distribution of delays and allowing a positive probability of selection of each discrete time including the end points.

If the assumed time of $1/2$ minute or of 4 minutes was known to be correct in every case, there would be no justification for the use of random times; however, it is intuitively obvious that such a "correct time all of the time" does not exist. In such cases, presumably a random time from a distribution approximating the true distribution of times for such events would be more realistic, on the average.

Thus in the above example, the magnitude of the difference in the expected number of kills, and the knowledge that the delay for a decision frequently varies with the decision maker and the situation, would seem to indicate that the use of random time delays was in order.

CHAPTER I

GENERATION OF RANDOM TIMES

SECTION 1

Random time delays for a chosen event in a particular play of a war game are commonly produced by making a probability transformation of a random number from a uniformly distributed population to a corresponding number from a population which is assumed to follow the true probability law governing the event.

By appropriate scaling, which will be covered in a later chapter, this random number is converted to a delay time, which is a number within the range of those possible for the event. It is this number which is then assessed against one of the players as the time necessary for the accomplishment of that event.

Delay time is itself a random variable, that is, a function whose possible values lie in some continuous interval of time. The interval is defined as the range of possible values the delay time may take on.

That the observation of a specific delay time for an event is a random phenomenon, i.e., one obeying probabilistic rather than deterministic laws, is obvious from any observation of the real world. The function which assigns a real number to each outcome of the event or experiment is called a random variable and is denoted by X , Y , Z , etc., and a

particular outcome of the experiment is denoted by x, y, z , etc.

The probability density function, $f(x)$, of a continuous random variable X , is a function such that

$$f(x) \geq 0 \quad \text{where } x \in R \text{ and } R = (a, b)$$

$$f(x) = 0 \quad \text{elsewhere}$$

and such that

$$\int_{-\infty}^{\infty} f(x) dx = 1$$

A probability density function for a continuous random variable, gives through integration, a description of the relative frequencies of the occurrence of particular intervals on the real line.

Associated with these relative frequencies or probabilities of numbers in a continuum is another function, $F(x)$, such that if x is the number associated with some outcome of the experiment, then

$$F(x) = \text{Probability } (X \leq x) = P(X \leq x)$$

and, $F(x)$, is related to the probability density function, $f(x)$, by

$$f(x) = \frac{d}{dx} F(x)$$

that is

$$F(x) = \int_{-\infty}^x f(x) dx \quad -\infty < x < \infty$$

and is called the cumulative probability or distribution function. (14)

If Y is a uniformly distributed random variable on the

interval $[0, 1]$, then

$$\begin{aligned} f(y) &= 1 & 0 < y < 1 \\ &= 0 & \text{elsewhere} \end{aligned}$$

Letting $Y = G(x)$, where $G(x)$ is the cumulative distribution function of the random variable, X , which we may identify as the true distribution of the random variable, X , the delay time; then by letting

$$x = G^{-1}(y)$$

we have a correspondence between the random variables X and Y and, through a knowledge of the distribution function of X , namely $G(x)$, are able to pass from a knowledge of a particular x to a particular y and from y to x . If $G(x)$ is the true distribution of the random delay times of an event, we may, by taking a random number from a uniformly distributed population, produce a number from the true distribution by the transformation above, called the probability integral transformation. (15)

This generation of the random time delay for the event presupposes the knowledge of the true distribution of the times for the accomplishment of the event, that is, knowledge of $G(x)$. Indeed with such knowledge, the generation of random time delays would be relatively easy, albeit, laborious.

In practice, the true distribution of times for the accomplishment of the event will generally be unknown and,

without this knowledge, we can proceed no further with the generation of random time delays and must perforce resort to arbitrary or fixed time delays.

If we would persist in the desire for the introduction of randomness in the war game we still, however, have two courses of action open to us.

These are: First, assume some distribution of delay times arbitrarily; for instance, the uniform distribution might be selected or, for that matter, any other distribution.

Second, seek some knowledge of the distribution of the particular event of interest in the real world.

The remainder of this paper will investigate the second course of action.

The true distribution of the times for the accomplishment of the event can be approximated with some knowledge of the times that the event actually takes when occurring in the real world. These times (the observations), if they are a random sample from the true distribution, are independent. Furthermore, if they are dependent in a known way, we may still apply many standard statistical techniques.

Having such data, we are faced with a standard curve fitting problem. First, to determine a curve which best approximates the true distribution and then to determine some quantitative measure of its goodness of fit. This measure reflects the "confidence" with which we will state

that the curve approximates the true distribution. Classical curve fitting consists in estimating the parameters of some distribution, assumed to be the true distribution on intuitive grounds, by using moments calculated from the sample. A further refinement, called the method of moments, developed by Karl Pearson in a series of papers (22) from 1895-1916, sets up a system of objective criteria by which a particular curve would be selected to approximate the true distribution. The selection of a particular curve depends on the values of certain quantities which are functions of the first four moments of the sample and the system presents a method of calculation of the curve.

The Pearsonian system is essentially the solution of the differential equation

$$\frac{dy}{dx} = \frac{y(m - x)}{(a + bx + cx^2)}$$

where the values of the parameters, a , b , c and m , determine the shape of the curve for a particular treatment. For example, the normal curve is generated if b and c are zero, m is the mean and a is the variance. Under this system, the sample mean m , and the sample moments about the mean, m_2 , m_3 , and m_4 , (i.e., the first three central moments) are calculated and are used to produce certain statistics dependent on the sample size, N , (the number of observations).

These statistics are:

$$k_1 = m$$

$$k_2 = \frac{N^2}{(N-1)(N-2)} m_2$$

$$k_3 = \frac{N^2}{(N-1)(N-2)} m_3$$

$$k_4 = \frac{N^2 [(N+1)m_4 - 3(N-1)s^4]}{(N-1)(N-2)(N-3)}$$

where s is the sample standard deviation

These statistics are used to provide estimates

$$\hat{\delta}_1 = \frac{k_3}{k_2^2}$$

$$\hat{\delta}_2 = \frac{k_4}{k_2^2}$$

of δ_1 and δ_2 which are Pearsonian measures of skewness and kurtosis. These quantities $\hat{\delta}_1$ and $\hat{\delta}_2$ are used to calculate $\beta_1 = \sqrt{\hat{\delta}_1}$ and $\beta_2 = \hat{\delta}_2 + 3$, which are tabled to indicate the proper form of the curve to be used.

Further, solution of the relations

$$b = -m = \frac{\hat{\delta}_1}{2(1+\delta)}$$

$$c = \frac{\delta}{2(1+\delta)}$$

$$a = 1 - 3c$$

$$\delta = \frac{(2\hat{\delta}_2 - 3\hat{\delta}_1^2)}{(\hat{\delta}_2 + 6)}$$

will yield the values of the parameters of the differential equation, which need only be integrated to give the density function.

Besides the labor of calculating the moments of the sample, there are some problems with the system in that except for the Poisson, normal and binomial distributions (18) the moments of the sample used to estimate the corresponding moments of the assumed theoretical distribution do not possess the minimum sampling variance; that is, they are not most efficient

estimates in the statistical sense.

This difficulty may be overcome by the use of maximum likelihood estimators (19), which, however, require that the form of the density function (Normal, Beta, Gamma, etc.) be known or at least assumed. This is not entirely unreasonable.

Basically, the principle of maximum likelihood is to consider every possible value that the true parameter of the distribution, which we have assumed, might have and for each of these values compute the probability that the particular set of observations at hand (the sample) would have occurred if that value were the true value of the parameter of our assumed distribution.

Now of all these possible values of the parameter, we would choose the one for which the probability of our sample occurring would be the greatest. Instead of working with the probability density function of the assumed distribution, we may work with its logarithm L (purely as a matter of convenience). By elementary calculus if

$$\frac{dL}{d\theta} = 0 \quad \text{and} \quad \frac{d^2L}{d\theta^2} < 0$$

where θ is a parameter of the assumed density function; we indeed have a maximum and solution of the equation,

$$\frac{dL}{d\theta} = 0$$

for θ , will yield a maximum likelihood estimator for θ , called $\hat{\theta}$. In the case of the maximum likelihood estimator

of θ , if it exists, it can be proved that the quantity

$\sqrt{N}(\hat{\theta} - \theta)$ is asymptotically normal; and that $\hat{\theta}$ has the smallest limiting variance of any estimator. This implies that, as N becomes very large, $\hat{\theta}$, the maximum likelihood estimator, becomes most efficient. (20)

Unfortunately, for this sophistication of the estimators, we have sacrificed the crude, laborious but practical Pearsonian System, which has the advantage of requiring no prior knowledge of the form of the distribution function for its application. But often, there will exist some body of past experience which will lead us to choose some particular shape or type of curve as being most probable (as, for example, the general concensus that distributions of times for simple actions performed by human beings will generally be unimodal).

This knowledge or the willingness in a particular situation to make an intuitive leap will lead us to assume some particular type of distribution to be the best approximation to the true one. In this case, the most fruitful method will be to estimate the parameters of the assumed distribution by the method of maximum likelihood, assuming such estimates exist. These methods or derivative ones such as the Gram-Chalier Series, Edgeworth Series, etc. (18) all provide the equations of curves which will fit the data to a greater or lesser extent.

To recapitulate; having been given a set of data, which can reasonably be assumed to be a random sample for a particular event, we would apply the Pearsonian procedure to determine the shape of the curve to be used as the approximation to the true distribution in the case of no prior knowledge about the nature of the true distribution. If some knowledge of the form of the true distribution is available, we would estimate the parameters of the assumed distribution by the method of maximum likelihood. The principal methods of establishing a quantitative measure as to what extent the particular distributions that have been derived or assumed to be true by the above approaches actually approximate the true distribution are: the Chi-square Test and the Kolmogorov Test.

The Chi-square Test (χ^2) is a measure of the degree to which a series of observed frequencies deviate from assumed theoretical frequencies over the whole range of values. The relative magnitude of this discrepancy is defined by the quantity

$$\chi^2 = \sum_{t=1}^n \left\{ \frac{(f_{ot} - f_t)^2}{f_t} \right\}$$

where f_{ot} is the relative frequency of the observations in the interval t , f_t is the relative frequency of the assumed distribution in the interval t , n , is the number of intervals. This value of χ^2 is then compared with the tabulated values of the χ^2 distribution which give, for various degrees of

freedom, the probability with which one would expect such a value of χ^2 . We may therefore set up a test of the hypothesis that this sample indeed comes from some assumed distribution by calculating the value of χ^2 and, on comparison with the tabulated values, accept or reject the hypothesis depending on whether the observed value of χ^2 is greater or less than the theoretical value at any probability level we desire. Now it will sometimes occur that several assumed distributions will fit the data and, in this case, there is some justification for choosing the one with the minimum value of χ^2 as the curve of best fit. But extreme caution must be exercised in the application of the Chi-square Test. The reason for this is that no account is taken in the test of the distribution of the discrepancies between f_{ot} and f_t at particular intervals over the range. It is conceivable, because of the nature of the statistic, that, having calculated a value of Chi-square by fitting a curve very closely over part of the range and relatively far away in the other part of the range, this value would be lower than that which would be produced by calculating a value of Chi-square for a curve which displayed a more constant discrepancy over the particular intervals. Yet this distribution of discrepancies may materially influence our judgment as to the goodness of fit.

For small sample sizes, where the Chi-square Test breaks down, we have still available to us the Kolmogorov-Smirnov statistic, which, in essence, measures the maximum

vertical difference between the sample cumulative distribution function (a step function) and the assumed distribution. Then comparison of the value of the statistic is made with the tabulated value as in the Chi-square Test and acceptance or rejection made on the basis of the magnitude of the computed statistic.

The statistic (21) for the two sided case is;

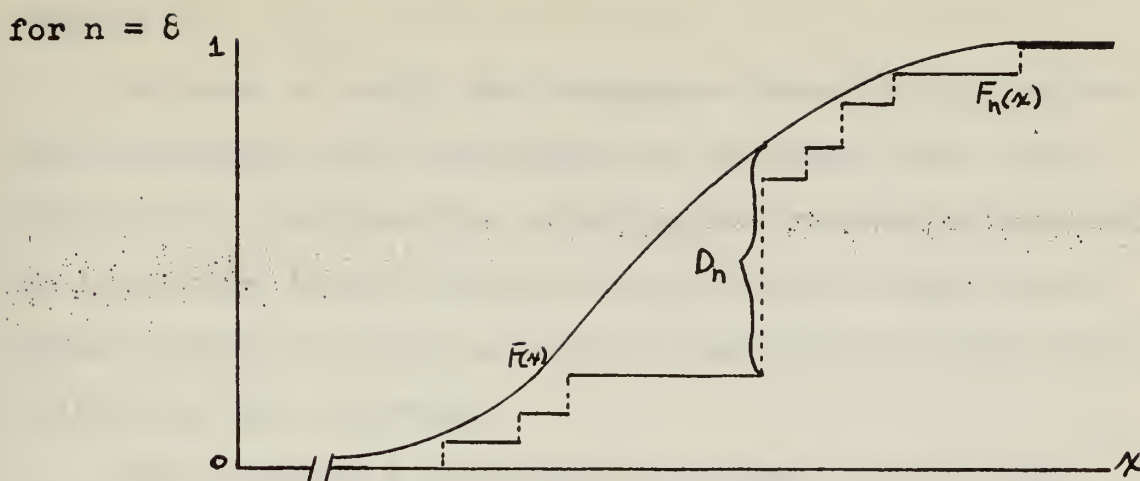
$$D_n = \sup_{\text{all } x} |F_n(x) - F(x)|$$

where $F(x)$ is the assumed distribution function and

$$\begin{aligned} F(x) &= 0, \quad x < X_{(1)} \\ &= \frac{j}{n}, \quad X_{(j)} \leq x < X_{(j+1)}, \quad j = 1, \dots, n-1 \\ &= 1, \quad x \geq X_{(n)} \end{aligned}$$

is the sample cumulative distribution function of the ordered observations.

The statistic is itself the least upper bound (sup) of the difference between the cumulative distribution functions taken over all the values of the argument. It should be noted that whereas the assumed distribution $F(x)$, is continuous and non-decreasing, the sample distribution $F_n(x)$, will be discrete and so constant over each interval, for example,



$D_n(F)$ can be considered as having come from a random sample on a uniform distribution in the interval $[0,1]$ (19) and, hence is independent of the choice of F , so the Kolmogorov-Smirnov Statistic is distribution free. This statistic has an advantage over the Chi-square Test in that it deals directly with the value of the observations and avoids the problems inherent in the summing process of the Chi-square Test, which, as has been mentioned, may tend to obscure lopsided fits. The choice of one or both of these tests of goodness of fit will, of course, depend on the particular situation and the degree of approximation desired.

This brief summary has been intended as the general outline of the procedure to be followed, if sufficient data is available for approximating the distribution of delay times. It should be borne in mind that the importance of the knowledge of an approximation to the true distribution function must justify the expenditure of time and effort necessary to obtain it.

SECTION 2

In order to apply the techniques discussed in Section 1, it is necessary that the analyst be provided data. Presumably the procedure for obtaining the requisite data would be to analyze times measured for particular events under actual combat situations. We can immediately foresee difficulties in this approach.

First, there is little data recorded in sufficient accuracy and detail to provide adequate estimators for such parameters as the analyst might deem necessary.

Second, were the data for some particular event available, one could hardly assume that this data would yield the average times associated with similar events unless the event had been run, and the times recorded, a large number of times.

Last, the time required and the expense of the necessary data search is prohibitive.

One might next logically consider conducting experiments in real time situations and simply recording the desired times in the context of the experiment. Quite a bit of work has been done in this area.

In the course of this thesis, a careful search was conducted of existing published data. Contact was established with the U.S.A. Combat Development Experimentation Center, Fort Ord, California, and a thorough search of their

published and unpublished experiments yielded little of value in determining the underlying distribution of the event being simulated. This is not surprising. The experiments conducted by CDEC, like those conducted by similar organizations, are extremely complicated. In each case, the experiment is designed to consider only some specific portion or element of the entire event. The use of scenario type situations and arbitrary time delays, as defined in the introduction to this thesis, preclude the results of the experiment being used to describe the distribution of the entire event.

In general, the conduct of experiments in real time situations to estimate the parameters of an event in an actual combat situation is prohibitive for one of the three following reasons:

First, the actual combat situations desired are almost impossible to duplicate.

Second, the near simulation of actual combat situations is, not only dangerous to the participants, but prohibitively expensive.

Third, the number and diversity of the events it is necessary to describe is beyond effective description.

In light of the above, we are forced to conclude that in general, the procurement of the necessary data, from past actions, or by experimentation, is not practical for the majority of military events.

The following chapter will suggest a third source of information by which it is felt that a form of data sufficient for the description of the times associated with military events can be obtained.

CHAPTER II

AN ALTERNATIVE USING MINIMAL DATA

SECTION 1

In Chapter I we discussed the difficulties entailed in attempting to describe military events using recorded or experimentally obtained data.

There is a third possible method of arriving at the desired parameters for the underlying time distributions we desire to describe. This method entails the use of professional estimates. Professionals, utilizing their background and experience in a particular field should logically be able to give a reasonable estimate of the time to complete a specific event under certain conditions.

This method offers advantages over the two mentioned in Chapter I in that it is cheap, and with the availability of professionals in the field of interest, quite rapid. It suffers from two obvious shortcomings:

First, what estimates are you going to ask the professionals to make?

Second, how valid will these estimates be?

There is really no good way to check on the validity of subjective estimates short of running a series of the same experiment and comparing the mean time of the experiment with the subjectively estimated time for the same event.

The answer to neither of these questions is readily apparent.

It would logically seem that in answer to the question, "What estimates do you want from a professional?" One might reply, "What estimates can he be expected to make with reasonable consistency?" Obviously, if we are to request an individual to furnish answers based on his experience, we must ask questions he has frequently asked and answered for himself.

In this light it is believed that the professional military man can interpolate or extrapolate from the domain of his experience to furnish the following four estimates for many events.

1. The shortest possible time in which the event can be completed, call it $\hat{t}_o$.
2. The most likely time required to complete the event, say $\hat{t}_m$.
3. The most pessimistic time to complete the event, $\hat{t}_p$.
4. The percentage of times, $\hat{P}$, the event takes place within some given time, centered about t_m .

Where $\hat{t}_o$, $\hat{t}_m$, $\hat{t}_p$, and $\hat{P}$ are estimates of the real values t_o , t_m , t_p and P .

Before we attempt to answer the second question, let us examine subjective estimates in greater detail. Subjective estimates are not new. The military commander always has, and probably always will be required to use his estimates

and those of his staff to effectively plan an operation. How does he check his estimates? If the operation comes off somewhat as planned, he has done what we are unable to do. He has estimated the time for an event and then run an experiment which furnishes an immediate check on his estimate of this time. From each operation he learns. It is this equity in experience from which we must draw our estimates for events we cannot otherwise practically measure. One would generally feel that if a time estimate from one knowledgeable individual was good, that the average of the estimates of ten of these professionals would be better. Ideally it would be hoped that as the number of qualified persons estimating the time for some event was increased, the average of these estimates would approach the true value of the parameter being estimated.

It is under the above assumption that the use of subjective estimates is justified.

How are subjective estimates obtained? The pollsters have long realized that the obtaining of subjective estimates from a population is a delicate and difficult problem. Through experience, they have come to realize that the phrasing of the question, the occupation of the individual, and even the time of day can affect the answer received to a particular question. It may well happen that a wide variation in the response to a particular question is the result of a poorly

defined or ambiguous question rather than a demonstration of the recipient's inability to furnish the hoped for response.

In general, subjective estimates are obtained through either an interview or through response to a written questionnaire. Both methods have their shortcomings. The written questionnaire generally suffers from poor response, misinterpretations due to ambiguous wording, and the failure to distinguish between those persons unable to answer the question and those who are able, but unwilling to answer. The interview (while probably furnishing more accurate information) is more costly, frequently inconsistent due to variations in the delivery of the question, and difficult to quantify.

Since in our particular case point estimates of specific times are desired, the use of written questionnaires would appear to be the more practical approach.

In an attempt to overcome some of the objections to the written questionnaire, the following rules to be used in obtaining subjective time estimates of military events have been set down.

1. Select a group of military professionals capable of furnishing time estimates for the requested events. It is worth noting, that since we are after $\hat{t}_o$, $\hat{t}_m$, $\hat{t}_p$ and $\hat{P}$ for particular events we must choose the group to fit the events, not the events to fit the group. Where possible, a homogeneous group comprised

of individuals with various specialties should be chosen. This will allow a comparison of the stratification within the group, to guard against the consistent, but biased, answer.

2. Assure response to the questionnaire by follow-up. Many persons who fail to respond to a questionnaire through indifference are capable of producing an acceptable answer.
3. The questions to be answered must objective; fill-in the-blank type questions will facilitate the ultimate reduction and comparison of data.
4. The information desired must be defined carefully and precisely on the first page of each questionnaire. Only an introduction should be required from the administrator of the questionnaire. Every effort must be made to keep the questionnaire as short as possible.
5. The questionnaire should be carefully pretested by a trial group to check the adequacy of event descriptions and guard against ambiguities, etc., prior to general use.
6. The data resulting from the experiment must be carefully evaluated.

Following the above listed guide, a questionnaire covering ten hypothetical military events has been prepared,

(Appendix (2)). This questionnaire is directed to the experienced Marine Corps officer. The officers are assumed to vary in rank from Captain to Lieutenant Colonel, all with service in excess of ten years.

Let us take the definitions for t_o , t_m and t_p from our sample questionnaire. This more precise definition will be required when techniques are discussed for analyzing the data from an experiment.

$t_o \equiv$ most optimistic time = The minimum time required to accomplish the action if all circumstances, personnel reaction, weather, terrain, etc., combine to work in your favor.

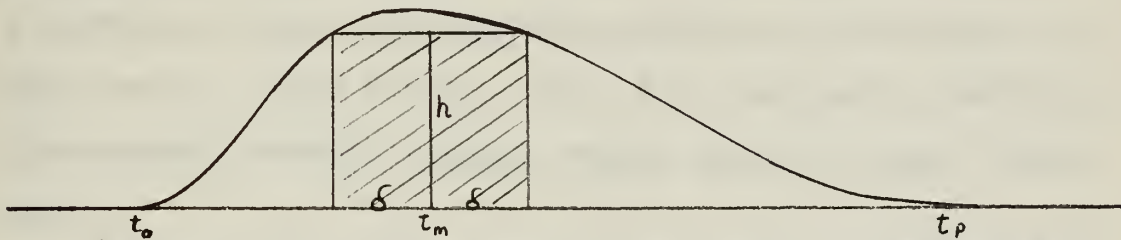
$t_m \equiv$ most likely time = The time required to accomplish the action under "normal" circumstances. "Normal" to be construed as those day to day conditions you would expect to encounter under circumstances such that the described action could logically be assigned.

$t_p \equiv$ most pessimistic time = The time required to accomplish the action under extreme adverse conditions, and poor response on the part of participants. The conditions are not so severe as to preclude accomplishment of the mission, but weather, fatigue, human error, etc., all exercise strong degrading effects.

Inspection of the questionnaire will indicate that each question is followed by the statement:

"Percentage (%) of times the event will occur within 6 hours/minutes of the most likely time."

This question is designed to obtain a feel for the amount of the distribution located within a certain range of the most likely time. Answers to this question afford the experimenter a method of obtaining a rather crude, but nevertheless necessary estimate of the central tendency of the distribution.



In the figure above one can see that $h = \frac{\text{Percentage}}{2\delta}$,

where the shaded area under the arbitrary curve is the percentage estimate of the times an event will occur within δ units of the most likely time.

It is apparent that the choice of δ is important to the outcome of the experiment in that the ultimate selection of an approximating distribution may depend on h . As can be seen from the figure, the smaller our choice of δ , the more accurate our estimate of h . It is therefore desirable to choose the smallest interval (2δ), for which the experimenter can obtain consistent estimates of the percentage of times an event will fall in this interval. The question to be answered is probably more psychological than statistical. How small an interval can be chosen such that reasonably consistent estimates can be anticipated?

In the questionnaire (Appendix (2)) δ was chosen such that $2\delta \doteq \frac{R}{4}$ or 25 per cent of the range.

In most cases the experimenter can himself furnish a crude estimate of the range of the event to be considered. It would then be a relatively simple matter to prepare several questionnaires using the same questions, but varying δ as a function of the experimenters estimate of the range. In this way one could obtain a feel for a subject's ability to discriminate between various ranges about the most likely time.

Before we proceed to describe methods for evaluating data, it is best that we discuss some of the assumptions we are forced to make in conducting our experiment in this manner.

When we decide to use subjective estimates as a means of securing t_o , t_m , t_p and P , we are assuming that our sample values are distributed in some fashion about the true values of these parameters t_o , t_m , t_p and P .

We further assume that these estimates are so distributed that the central limit theorem holds. Thus, where N is the number of estimates, $\bar{\hat{t}}_o$, the mean of our estimates for t_o , is normally distributed with mean t_o and variance decreasing to zero as $N \rightarrow \infty$, irrespective of the underlying distribution of $\hat{t}_o$. The same holds for $\bar{\hat{t}}_m$, $\bar{\hat{t}}_p$ and $\bar{\hat{p}}$.

Therefore, under the foregoing assumptions, and for fairly large N we feel that we know t_o , t_m , t_p and P quite accurately.

We are fortunate that our method of sampling by questionnaire is relatively inexpensive. This allows us to choose a sample size large enough to insure that our experiment falls within the domain of large sampling statistics as defined above. This permits us to consider our estimates for t_o , etc., as being approximately normally distributed.

Let us assume that 100 of the questionnaires in (Appendix (2)) have been duly completed and returned. What will we do with the data? How will we decide whether these estimates are adequate for our needs? What tests should we run?

We want estimates for t_o , t_m , t_p and h. Let us first compute these for each event.

$$\bar{t}_o = \frac{1}{100} \sum_{i=1}^{100} \hat{t}_{oi};$$

$$\begin{aligned} i &= 1, 2, \dots, 100 \\ j &= 1, 2, \dots, 10 \end{aligned}$$

where $\hat{t}_{oi}$ is the most optimistic time

estimate from the i th questionnaire for the j th event.

$\bar{t}_{pj}$ and $\bar{t}_{mj}$ would be calculated in a like manner.

As previously defined, $\hat{h} = \frac{\text{Percentage}}{2\delta}$. Thus, $\hat{h}_j = \frac{\bar{P}_j}{2\delta_j}$
 $= \frac{1}{200\delta_j} \sum_{i=1}^{100} P_{ij}$. Where P_{ij} is the estimate from the i th questionnaire for the percentage of times the j th event would fall within $\pm \delta$ of the most likely time.

Taking say the j th event, let us make a linear plot for each of our parameters $\hat{t}_{oj}$, $\hat{t}_{mj}$, $\hat{t}_{pj}$ and $\hat{P}_j$, to investigate

the possibility of outliers or sample stratification.

When we take our sample of the time estimate for a particular event we are assuming that all of our observations in the set, comprised of our sample constitute a random sample from the same population. For various reasons there will occasionally occur in the sample a small number of observations which differ considerably from the others. These observations may cause us to doubt that all of our observations are truly from the population we intended to sample.

As an example let us assume we have plotted $\hat{t}_0$ for the j th event. Let us further assume that this linear plot indicated that 99 of our 100 estimates of $\hat{t}_0$ were reasonably grouped, but that the remaining, say the k th observation was significantly different from the rest. We cannot automatically assume that this k th observation is an outlier and thus delete it from the sample.

We can however delete it for cause. If on checking the k th questionnaire we noted that it had been completed by a twenty-one year old, 2nd Lieutenant, Disbursing officer, while the event being estimated was described as, say, the time to conduct some facet of an air operation, we would probably be safe in assuming that our k th observation was an outlier. In this event we could delete the k th question and recompute our estimates based on a reduced sample N , where in

this case $N = 100$ - (samples rejected as outliers).

In addition to the rejection of outliers for subjective cause there are several objective statistical tests designed to reject outliers at a certain confidence level. Since we are dealing with a rather large sample, the number of outliers that would be rejected statistically that could not be rejected for subjective cause is considered not to be statistically significant.

Again consider the linear plot of say $\hat{t}_0$ for our j th event. We might well observe that our sample estimates were in say two or three distinct groups. In this case it might pay to stratify our group by specialty say, Infantry officers, Aviators, Supply officers etc. If our event to be estimated was described as an infantry action we would probably profit by taking our estimate of $\hat{t}_0$ from the sample of officers most closely associated with this type of an event.

Again this reduction in sample size for cause would require a recomputation of all of our estimates.

For the remainder of the discussion on sample evaluation let us define:

N = the sample sizes as reduced by outliers and/or stratification thus for example

$$\hat{t}_0 = \frac{1}{N} \sum_{i=1}^N t_{0,i}$$

It will be noted that throughout this discussion of sample evaluation we have not talked of the accuracy of an estimate. As stated earlier, there is really no way we can

judge the accuracy of an estimate. We can however discuss the consistency of our estimates. Probably the best measure of the consistency of an estimate, say $\hat{t}_{in}$, is the measure of its variance.

$$\text{Var } \hat{t}_{m,j} = \frac{1}{N} \sum_{i=1}^N \left(\hat{t}_{m,j} - \bar{\hat{t}}_{m,j} \right)^2$$

When would our sample estimates be rejected? The answer to this question will depend entirely on the user. Of the ten events described in our questionnaire, very likely some will produce variations which would make one hesitate to use these estimates, say $\text{Var } \hat{t}_0 > \epsilon$, where ϵ is some number established by the user prior to conducting the experiment. In this case, the only conclusion we can draw is that either the question was ambiguously worded, or that the professionals answering our questions were unable to give a satisfactory response.

What can we do if our estimates for a particular event have such wide variation as to be considered inconsistent or unusable. If the event is of relatively long range and it is still felt that random time delays should be used, there remains little recourse but to re-word the offending question and distribute it to another group of professionals. Where the event is of relatively short range and the variations in the estimates for the required parameters are large, one might feel justified in the use of fixed or arbitrary time delays in assessing the time to complete this event.

SECTION 2

We have previously seen that in order to generate the random times we seek, we must have some approximation to the probability distribution underlying these times. Since data of the type usually used to derive such approximations is not available, we have suggested use of data of a type which is available. In this section, we shall examine what effect possession of such data would have on our knowledge of certain probability distribution laws.

Table II-1 contains some common probability density functions. An examination of the table indicates that in each case, $f(x)$ is dependent upon some constants listed in the column headed "Parameters." In order to specify the probability laws best describing real life events, it is necessary to know both the form of the probability density function, and the value of its parameters. Hence, if we wish to generate a random time to simulate the time required to conduct a specific operation, we must have knowledge of the form of the probability density function best describing the distribution of times, and we must be able to determine the parameters involved.

For the present, we will assume that we know the form of the probability density function, and that it is one of those listed in Table II-1. We shall further assume that we have precise information regarding certain aspects of our event,

PROBABILITY DENSITY FUNCTIONS

NAME	$f(x)$	PARAMETERS
Uniform	$\frac{1}{b-a}$ for $a \leq x \leq b$; 0 elsewhere	a, b
Triangular	$\frac{2x}{ab}$ for $0 \leq x \leq a$; $\frac{2(b-x)}{b(b-a)}$ for $a \leq x \leq b$; 0 elsewhere	a, b
Normal	$\frac{1}{\sqrt{2\pi}\sigma} \exp\left\{-\frac{1}{2\sigma^2}(x-\mu)^2\right\}$ $-\infty < x < \infty$	$\mu, \sigma > 0$
Gamma	$\frac{1}{\alpha! \beta^{\alpha+1}} x^\alpha \exp\left\{-\frac{x}{\beta}\right\}$ for $0 \leq x < \infty$; 0 elsewhere	$\alpha > -1$ $\beta > 0$
Log Normal	$\frac{1}{\sqrt{2\pi}\sigma x} \exp\left\{-\frac{1}{2\sigma^2}(\ln x - \mu)^2\right\}$ for $x > 0$; 0 elsewhere	$\mu, \sigma > 0$
Beta	$\frac{(\alpha+\beta+1)!}{\alpha! \beta!} x^\alpha (1-x)^\beta$ for $0 \leq x \leq 1$; 0 elsewhere	$\alpha > -1$ $\beta > -1$
Truncated Normal	$\frac{1}{\sqrt{2\pi}\sigma A} \exp\left\{-\frac{1}{2\sigma^2}(x-\mu)^2\right\}$ for $a \leq x \leq b$; 0 elsewhere $A = \int_a^b \frac{1}{\sqrt{2\pi}\sigma} \exp\left\{-\frac{1}{2\sigma^2}(x-\mu)^2\right\} dx$	$A, \mu, \sigma > 0$

Table II - 1

which will be the conduct in a finite time of some military operation.

Specifically, we shall assume that the following are exactly known;

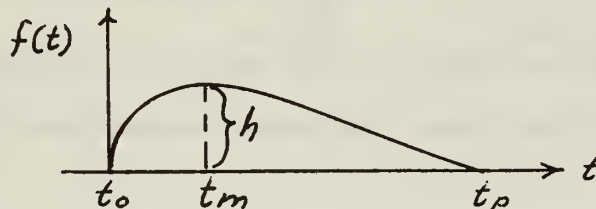
- t_o - The minimum time which could be required to complete the operation.
- t_m - The most likely (most frequently encountered) time required to complete the operation.
- t_p - The maximum time which could be required to complete the operation.

Henceforth, since our random variable is time, we shall denote it by T , and its density function by $f(t)$. Therefore,

$$F(\sigma) = P[T \leq \sigma] = \int_{t_o}^{\sigma} f(t) dt \quad \sigma > t_o$$

where the lower limit of integration is the minimum time.

Times less than t_o or greater than t_p have no significance for us. If we make a graphical representation of $f(t)$, it would appear approximately as



where t_o , the minimum time is the zero of our t axis. t_m is the point at which $f(t)$ has its maximum value. This is denoted as the mode of the density function.

The last item about which we shall assume to have exact

knowledge is

h - the height of the density function at its mode.

With the assumed information, we shall define

$$R = t_p - t_o \quad \text{as the range}$$

$$R_o = t_m - t_o \quad \text{as the premodal range}$$

$$R_p = t_p - t_m \quad \text{as the post modal range}$$

$$R = R_o + R_p$$

Recalling that each probability function has associated with it certain parameters we shall now proceed to develop knowledge of the parameters from the information assumed above. The seven density functions from Table II-1 will be dealt with in turn.

a. The Uniform Probability Density Function

$$f(t) = \frac{1}{b-a} \quad a \leq t \leq b$$

clearly
$$= 0 \quad \text{elsewhere}$$

$$t_o = a \quad \text{and} \quad t_p = b$$

$$\therefore f(t) = \frac{1}{R} \quad t_o \leq t \leq t_p$$
$$= 0 \quad \text{elsewhere}$$

The density function has no mode as such, and its height, i.e., h , is everywhere constant, and completely determined by the range.

b. The Triangular Probability Density Function

$$f(x) = \frac{2x}{ab} \quad 0 \leq x \leq a$$

$$= \frac{2(b-x)}{b(b-a)} \quad a \leq x \leq b, \quad \text{and} \quad f(x)=0 \quad \text{elsewhere}$$

Since the range is $0 \leq x \leq b$, we must scale $t_o \leq t \leq t_p$ into that range. If we let

$$x = t - t_0, \text{ then}$$

$$a = t_m - t_0 = R_0 \quad \text{since the mode occurs at } x = a \text{ and } t = t_m$$

$$\text{and } b = t_p - t_0 = R$$

$$\text{then } f(t) = \frac{2(t-t_0)}{R R_0} \quad t_0 \leq t \leq t_m$$

$$= \frac{2(R-t+t_0)}{R R_p} \quad t_m \leq t \leq t_p$$

$$= 0 \quad \text{elsewhere}$$

and h is determined by the range for this density function also.

c. The Normal Probability Density Function

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2\sigma^2} (x - \mu)^2 \right\} \quad -\infty < x < \infty$$

This density function is determined by μ and σ , its mean and standard deviation. The mean occurs at the mode which implies

$$\mu = t_m$$

$$\text{and } f(t_m) = h = \frac{1}{\sqrt{2\pi}\sigma} \quad \text{or } \sigma = \frac{1}{\sqrt{2\pi}h}$$

The value of $f(t)$ at a distance 4σ from the mean is approximately 0.0001. Between $(\mu - 4\sigma)$ and $(\mu + 4\sigma)$, to at least four significant places, the area under $f(t)$ is unity.

In this function, t is permitted to range from $-\infty$ to $+\infty$, and we have limited ourselves to consideration of times with the finite range t_0 to t_p . Two situations can exist,

1. R_0 and R_p can both be greater than or equal to 4σ
or
2. one or both are less than 4σ .

From the preceding discussion of values of $f(t)$ at 4σ , it may be seen that in the first situation, we can closely approximate $f(t)$ by $f(t) \doteq h \exp \{-\pi h^2 (t-t_m)^2\}$
using $\sigma = \frac{1}{\sqrt{2\pi}h}$ and $\mu = t_m$

The second situation definitely implies that although the time required to complete the operation may be normally distributed, the density function must be cut off at one or both ends (truncation). This leads to the rejection of use of the normal density function under these circumstances, and to the truncated normal density function which is presently to be discussed.

d. The Gamma Probability Density Function

$$f(x) = \frac{1}{\alpha! \beta^{\alpha+1}} x^\alpha e^{-x/\beta} \quad 0 \leq x < \infty$$

$$= 0 \quad \text{elsewhere}$$

If we set $\frac{d}{dx} f(x) = 0$ and solve for x , we find that the mode occurs at

$$x = \alpha\beta$$

Since $f(x)$ has its origin at zero, we must rescale our time so that t_0 corresponds to 0. To accomplish this, let

$$x = t - t_0$$

The mode then occurs at

$$x = t_m - t_0 = R_0 = \alpha \beta$$

and

$$f(t) = \frac{1}{\alpha! \beta^{\alpha+1}} (t - t_0)^\alpha e^{-\frac{t-t_0}{\beta}} \quad t_0 \leq t < \infty$$

$$= 0 \quad \text{elsewhere}$$

Using the above, and evaluating the function at t_m ,

$$h = f(t_m) = \frac{1}{(\frac{R_0}{\beta})! \beta^{\frac{R_0}{\beta}+1}} (R_0)^{\frac{R_0}{\beta}} e^{-\frac{R_0}{\beta}}$$

Making use of Stirlings Approximation

$$n! \sim (2\pi)^{1/2} n^{n+1/2} e^{-n},$$

$$h \sim \frac{1}{\sqrt{2\pi} (\frac{R_0}{\beta})^{\frac{R_0}{\beta}+1/2} e^{-\frac{R_0}{\beta}} \beta^{\frac{R_0}{\beta}+1}} (R_0)^{\frac{R_0}{\beta}} e^{-\frac{R_0}{\beta}}$$

which reduces to

$$h \sim \frac{1}{\sqrt{2\pi} \sqrt{R_0 \beta}} \quad \text{or} \quad \beta \sim \frac{1}{2\pi R_0 h^2}$$

and

$$\alpha = R_0/\beta \sim 2\pi R_0^2 h^2$$

Thus if our maximum time had no bound, h and R_0 are sufficient to determine approximations to α and β .

It should be noted that the approximations depend on t_0 , t_m and h , and are not affected by t_p .

Under our conditions however, we do have a finite maximum time. This implies that our density function is truncated at that point, and in order that

$$1 = \int_{\infty} f(t) dt, \quad \text{we must adopt}$$

44

some alternative. There are three readily apparent from which to choose

- (a). We may consider that there exists at t_p a discrete probability mass point where the probability concentrated at t_p is defined by

$$\int_{t_p}^{\infty} \frac{1}{\alpha! \beta^{\alpha+1}} (t-t_0)^{\alpha} e^{-\frac{t-t_0}{\beta}} dt$$

then

$$\begin{aligned} f(t) &= \frac{1}{\alpha! \beta^{\alpha+1}} (t-t_0)^{\alpha} e^{-\frac{t-t_0}{\beta}} \quad t \leq t < t_p \\ &= \int_{t_p}^{\infty} \frac{1}{\alpha! \beta^{\alpha+1}} (t-t_0)^{\alpha} e^{-\frac{t-t_0}{\beta}} dt \quad \text{at } t=t_p \\ &= 0 \quad \text{elsewhere} \end{aligned}$$

- (b). We may define

$$\begin{aligned} f(t) &= \frac{1}{A \alpha! \beta^{\alpha+1}} (t-t_0)^{\alpha} e^{-\frac{t-t_0}{\beta}} \quad t_0 \leq t \leq t_p \\ &= 0 \quad \text{elsewhere} \end{aligned}$$

where
$$A = \int_{t_0}^{t_p} \frac{1}{\alpha! \beta^{\alpha+1}} (t-t_0)^{\alpha} e^{-\frac{t-t_0}{\beta}} dt$$

and is regarded as a constant

(c).
$$f(t) = \frac{1}{\alpha! \beta^{\alpha+1}} (t-t_0)^{\alpha} e^{-\frac{t-t_0}{\beta}} \quad \text{may}$$

be asymptotically zero at t_p , and the probability

$$\int_{t_p}^{\infty} f(t) dt \doteq 0 \quad , \quad \text{in which}$$

case, we regard t_p as being so large that its mathematical effect is that of ∞ .

Under alternative (a), we have in effect established that the probability at t_p may be disproportionately large, but have not invalidated our determination of an approximation to the parameters α and β .

Under alternative (b), the probability of a time greater than t_p has been redistributed over the interval $[t_0, t_p]$, but it has presented us with a third parameter in $f(t)$. The effect of this is to invalidate our approximation to α and β . The mode will still be located at $t - t_0 = \alpha\beta$, but with only the information at our disposal, we can no longer obtain a closed expression for α , β and A . Numerical methods with the assistance of the Tables of the Incomplete Gamma Function (16) must be applied to evaluate the parameters.

Under alternative (c), we recognize that there exists some finite

$$\xi = \int_{t_p}^{\infty} f(t) dt \quad \text{where } \xi \text{ is}$$

the measure of the probability that a time will exceed the maximum, but regard ξ as being so small that it may be neglected. This implies that

$$f(t_p) = \frac{1}{\alpha! \beta^{\alpha+1}} R^{\alpha} e^{-R/\beta} \doteq 0$$

Since $\alpha > -1$ and $\beta > 0$, the constant term $\frac{1}{\alpha! \beta^{\alpha+1}}$

can only have the effect of reducing $f(t_p)$, and hence, if

$$R^\alpha e^{-R/\beta} \doteq 0, \implies f(t_p) \doteq 0$$

This will occur if

$$e^{R/\beta} \gg R^\alpha \quad \text{or}$$

$$R/\beta \gg \alpha \ln R$$

$$\frac{R}{\ln R} \gg \alpha \beta, \quad \text{but} \quad \alpha \beta = R_0$$

so that if

$$\frac{R}{\ln R} \gg R_0, \quad \text{we may disregard } \xi.$$

e. The Beta Probability Density Function

$$(i.) f(x) = \frac{(\alpha + \beta + 1)!}{\alpha! \beta!} x^\alpha (1-x)^\beta \quad 0 \leq x \leq 1$$

$$= 0 \quad \text{elsewhere}$$

The mode occurs at

$$x = \alpha / (\alpha + \beta)$$

Since the range is 0 - 1, we must scale our times as follows

$$\text{Let } x = \frac{t - t_0}{t_p - t_0} = \frac{t - t_0}{R}$$

$$t = Rx + t_0$$

The mode occurs at t_m or $x = \frac{t_m - t_0}{R} = \frac{R_0}{R}$

$$\therefore \frac{\alpha}{\alpha + \beta} = \frac{R_0}{R} = \frac{R_0}{R_0 + R_p}$$

$$\Rightarrow \frac{\alpha}{R_0} = \frac{\beta}{R_p} = k$$

$$(ii) \quad \alpha = k R_0$$

$$\beta = k R_p$$

where k is some constant

to be determined.

Evaluating (i) at $\frac{R_0}{R}$,

$$f\left(\frac{R_0}{R}\right) = h = \frac{1}{R} \frac{(k R_0 + k R_p + 1)!}{(k R_0)! (k R_p)!} \left(\frac{R_0}{R}\right)^{k R_0} \left(1 - \frac{R_0}{R}\right)^{k R_p}$$

where the factor $\frac{1}{R}$ is required by the scaling.

Simplifying

$$R h = \frac{(k R + 1)!}{(k R_0)! (k R_p)!} \left(\frac{R_0}{R}\right)^{k R_0} \left(\frac{R_p}{R}\right)^{k R_p}$$

Again making use of Stirlings Approximation,

$$R h \sim \frac{(k R + 1) \sqrt{2\pi} (k R)^{k R + \frac{1}{2}} e^{-k R}}{\sqrt{2\pi} (k R_0)^{k R_0 + \frac{1}{2}} e^{-k R_0} \sqrt{2\pi} (k R_p)^{k R_p + \frac{1}{2}} e^{-k R_p}} \frac{R_0^{k R_0} R_p^{k R_p}}{R^{k R}}$$

which reduces to

$$R h \sim \frac{k R + 1}{\sqrt{2\pi}} \sqrt{\frac{R}{k R_0 R_p}} \quad \text{or}$$

$$(iii) \quad \frac{k R + 1}{\sqrt{k}} = h \sqrt{2\pi R R_0 R_p}$$

a relationship which can be solved to provide a value of

h . Since from (ii) $\alpha = k R_0$

$\beta = k R_p$, the solution to

(iii) completely determines α & β in terms of R_0 & R_p

it should be noted, that the equation (iii) in solution places a bound upon h .

$$h^2 \geq \frac{1}{R_0 R_p} \quad \text{as lower values}$$

of h would give imaginary roots in solving for k .

This implies that although R_0 and R_p do not determine h for the Beta distribution, they do provide a lower bound.

f. The Log Normal Probability Density Function

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma x} \exp \left\{ -\frac{1}{2\sigma^2} (\ln x - \mu)^2 \right\} \quad x > 0$$

$$= 0 \quad \text{elsewhere}$$

The mode of the density function found by setting

$$\frac{d}{dx} f(x) = 0 \quad \text{and solving for } x \text{ is}$$

$$x = e^{\mu - \sigma^2}$$

Postponing discussion of the upper limit of the range we scale

$$\text{so that} \quad t - t_0 = x.$$

$$\text{At} \quad x = t_m - t_0, \quad R_0 = e^{\mu - \sigma^2} \quad \text{or}$$

$$\sigma^2 = \mu - \ln R_0$$

$$f(R_0) = h = \frac{1}{\sqrt{2\pi}\sigma R_0} \exp \left\{ -\frac{1}{2\sigma^2} (\ln R_0 - \mu)^2 \right\}$$

$$\text{or} \quad \sqrt{2\pi} h R_0 = \frac{1}{\sigma} \exp \left\{ -\frac{\sigma^2}{2} \right\}$$

$$2\pi h^2 R_0^2 = \frac{1}{\sigma^2} e^{-\sigma^2} \quad \text{upon squaring.}$$

By means of a simple graph of $\frac{1}{n} e^{-n}$ a value of n can be selected such that $\frac{1}{n} e^{-n} = 2\pi h^2 R_0^2$. $\sqrt{n}$ will then be σ and $n + \ln R_0$ will equal μ .

As was the case with the Gamma density function, we have so far been considering a density function whose variate is without upper limit, as applied to a time variate with a finite upper bound. Also as was the previous case, our determination of the parameters was independent of this upper bound.

In general, we may choose from three alternatives to resolve this problem.

- a. Regard t_p as a point possessing the following probability mass

$$(1) \int_{t_p}^{\infty} \frac{1}{\sqrt{2\pi} \sigma(t-t_0)} \exp\left\{-\frac{1}{2\sigma^2}(\ln(t-t_0)-\mu)^2\right\} dt$$

which leads to

$$f(t) = \frac{1}{\sqrt{2\pi} \sigma(t-t_0)} \exp\left\{-\frac{1}{2\sigma^2}(\ln(t-t_0)-\mu)^2\right\}$$

for $t_0 \leq t < t_p$

= (1) above for $t = t_p$

= 0 elsewhere

As long as $t_p > t_m$, the method of evaluating μ and σ remains valid.

b. Regard

$$f(t) = \frac{1}{A\sqrt{2\pi}\sigma(t-t_0)} \exp\left\{-\frac{1}{2\sigma^2}(\ln(t-t_0)-u)^2\right\}$$

for $t_0 \leq t \leq t_p$

= 0 elsewhere

where $A = \int_{t_0}^{t_p} \frac{1}{\sqrt{2\pi}\sigma(t-t_0)} \exp\left\{-\frac{1}{2\sigma^2}(\ln(t-t_0)-u)^2\right\} dt$

This has the effect of introducing an additional parameter A to consideration, and imposes the requirement for at least one item of information other than t_0 , t_m , t_p and h for a unique determination of u , σ and A.

c. $f(t) = \frac{1}{\sqrt{2\pi}\sigma(t-t_0)} \exp\left\{-\frac{1}{2\sigma^2}(\ln(t-t_0)-u)^2\right\}$

may be asymptotically zero at t_p which would imply that

$$\int_{t_p}^{\infty} f(t) dt \doteq 0$$

t_p as being so large that the probability of exceeding it is zero. This requires that

$$\frac{1}{e^{\frac{1}{2}\sigma^2[\ln(t_p-t_0)-u]^2}} \doteq 0 \quad \text{or}$$

$$e^{\frac{1}{2}\sigma^2(\ln R - u)^2} \gg 1$$

$$\Rightarrow \frac{1}{2}\sigma^2(\ln R - u)^2 \gg 1 \quad \text{which reduces to}$$

$$\ln R \gg u + \sigma\sqrt{2}$$

g. The Truncated Normal Distribution

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma A} \exp \left\{ -\frac{1}{2\sigma^2}(x-\mu)^2 \right\}$$

$$A = \int_a^b \frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2\sigma^2}(x-\mu)^2 \right\} dx$$

$$a \leq x \leq b$$

$$f(x) = 0 \quad \text{elsewhere}$$

The mode occurs at μ , and our assumed information

leads to

$$f(t) = \frac{1}{\sqrt{2\pi}\sigma A} \exp \left\{ -\frac{1}{2\sigma^2}(t-t_m)^2 \right\}$$

$$A = \int_{t_0}^{t_p} \frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2\sigma^2}(t-t_m)^2 \right\} dt$$

$$\text{for } t_0 \leq t \leq t_p$$

$$f(t) = 0 \quad \text{elsewhere}$$

It should be noted that A plays the part of a constant, but it is another parameter which must be evaluated.

Now,

$$i.) \quad f(t_m) = h = \frac{1}{A\sqrt{2\pi}\sigma} \quad \text{is a}$$

constraint on the two parameters A and σ , but with no more than this, we could find any number of A 's and σ 's and thus any number of possible distribution functions.

What is required is one further piece of information restricting a point on the curve traced by $f(t)$. The value

of $f(t)$, the height of the density function at either end point (t_o or t_p) would be sufficient.

If such information were in our possession, it could be used as follows:

$$f(t_p) = h_{t_p} = \frac{1}{A\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\sigma^2(t_p - t_m)^2}$$

and from i) $\frac{h_{t_p}}{h} = e^{-\frac{R_p^2}{2\sigma^2}}$
which leads to

$$ii) \sigma^2 = \frac{R_p^2}{2|\ln h_{t_p} + \ln h|}$$

could uniquely determine σ

If however, all that were known was that $f(t)$ was asymptotically zero at t_p , then the above would provide no assistance in determining σ^2 . Under the assumption though that $f(t) = 0$ at t_p , and that $f(t) > 0$ at t_o (i.e., $R_o < R_p$), then $.5 < A < 1$. This follows since $t_m > t_o$ and we have not truncated more than 1/2 the area under the basic normal distribution. Further, the basic normal distribution

$$f_n(t) = \frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2}\sigma^2(t - t_m)^2 \right\} \quad \text{will}$$

approach zero sooner than $f(t)$ since

$$f(t) = \frac{f_n(t)}{A} \quad \text{and} \quad A \leq 1$$

implies $f(t) \geq f_n(t)$ for any value of t

We also know from properties of the normal distribution that

$$f_n(t) \Big|_{4\sigma} \doteq 0 \quad \text{so that}$$

$$f(t) \Big|_{4\sigma} = \frac{.0001}{A} \quad \text{or} \quad .0001 < f(t) \Big|_{4\sigma} < .0002$$

We can proceed to approximate a value for σ by assuming that $f(t) \doteq 0$ at t_p implies

$$.0001 < f(t) \Big|_{t_p} < .001$$

Under this assumption, using tables of the normal distribution it can be inferred that

$$3.5\sigma < (t_p - t_m) < 4\sigma$$

$$3.5\sigma < R_p < 4\sigma, \quad \text{and that}$$

R_p can be used to obtain an approximation to σ .

SECTION 3

Section 1 outlines the general statistical procedures to be followed in attempting the normal pattern to determine the probability laws underlying a real-life event. Under many circumstances, these procedures cannot be followed.

It may be that the cost of carrying out the necessary repetitions of the event is prohibitive, or it may be that the event is an infrequent occurrence of nature, happening too rarely to obtain a sufficient sample. At any rate, the data necessary to follow standard procedures is not available.

In the large majority of cases, the times required for the completion of various military operations fall into this category.

What we shall next attempt is to examine some procedures which might be indicated in such cases, but where certain specific minimal data is available.

Given that we have

$\hat{t}_0$ an estimate of t_0 , the minimum time required to conduct the operation,

$\hat{t}_m$ an estimate of t_m , the most likely time required,

$\hat{t}_p$ an estimate of t_p , the maximum time required, and

$\hat{h}$ an estimate of h , the height of the underlying density function at t_m ,

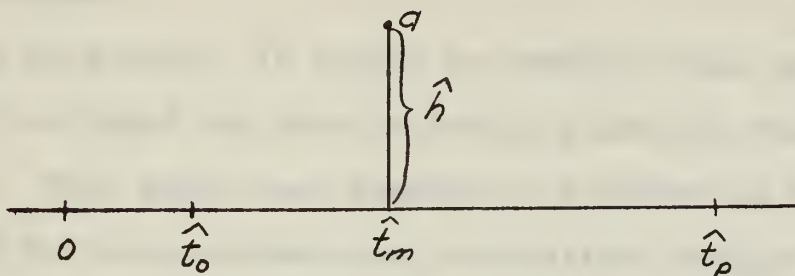
What knowledge can we deduce of the underlying probability law?

In Section 2 of this Chapter, it was shown how t_o , t_m , t_p and h could be used to determine the parameters of a set of Probability Density Functions. It should be recognized that at that point, the discussion was conducted from the position that there was no search for the form of the density function, and that t_o , t_m , t_p and h were regarded as known constants of the problem. Here we are dealing with estimates of the values of those constants, and have no real knowledge of the form of the density function.

If for the moment, we assume that we have accomplished one half our task, and have determined the form of the density function which we are seeking, we then have only to estimate its parameters to complete our approximation to the probability law underlying the event. Since in Section 2, we related these parameters to t_o , t_m , t_p and h for seven probability density functions, if the form with which we are dealing is a member of that set, we can use $\hat{t}_o$, $\hat{t}_m$, $\hat{t}_p$ and $\hat{h}$ to obtain estimators for the parameters. If the form of the function is not from that set, it will be assumed that manipulations of the type performed in Section 2 will produce some approximation to its parameters.

Having thus rather summarily disposed of the problem of parameter estimation, we turn to the much more difficult

problem of determining the form of the probability density function. The basic information we possess is that we have estimators for the following points,



and are attempting to define a unique curve which has its origin at $\hat{t}_0$, passes through the point a , which is at a height $\hat{h}$ above $\hat{t}_m$, and terminates at $\hat{t}_p$. The curve is further restricted in that its highest point must be at a , and that it represent a probability density function $f(t)$. The latter implies

$$f(t) \geq 0 \quad \text{and}$$

$$1 = \int_{\hat{t}_0}^{\hat{t}_p} f(t) dt$$

Under certain circumstances, almost any of the seven probability density functions discussed in Section 2 will generate such a curve. It must be stressed that these seven were selected for discussion only because they are among the most frequently encountered, and easy to work with. It would be more likely that there are an infinity of such curves, each with an associated $f(t)$. With no further information available, we cannot, in general hope to cope with any such

function dependent upon three or more parameters. Thus, we are forced to restrict our attention to examination of those functions for which we possess enough information to enable us to proceed.

At this point, it should be admitted that we may have already excluded the true probability density function, if it exists. This should not however be a matter of major concern. Even in the more conventional statistical problem (Chapter I), no claim is made that the "true" probability density function results. What is a result, is a density function which is a member of a set of known, in general tabulated, and convenient density functions, and which "best fits" the data. Even that much depends on far more data than is available under our present assumptions.

Since the seven functions in our set are mathematically well defined, and in general, well tabulated, we shall seek a "best fit" from that set. We shall first proceed to examine each of the various functions in turn to see what effect the data can have upon their applicability.

The uniform probability density function is dependent only upon its range $t_p - t_o = R$, and therefore having $t_o - t_p = R$ as an estimate of the range, the function would be determined. However, the value of $h = \frac{1}{R}$ is also determined by R . This would infer that we should reject the possibility that the uniform is our function if h is such that it differs

considerably from $\frac{1}{R}$.

A very much similar situation exists for the triangular probability density function where $h = \frac{2}{R}$. If h is considerably different from $\frac{2}{R}$, we would reject the triangular from consideration.

As developed in Section 2, the Beta probability density function also has a restriction in that $h^2 > \frac{1}{R_0 R_p}$. Therefore, we would reject it from consideration if $\hat{h}^2$ were less than $\frac{1}{\hat{R}_0 \hat{R}_p}$.

Consideration of the normal density function involves several factors. To begin with, it is a strictly symmetrical distribution. Unless it is to be truncated, a case presently to be discussed, the value of h fixes its standard deviation σ . If this is the case, and either $\hat{R}_0$ or $\hat{R}_p$ are less than 4σ , then there are contradictions to our assumptions in Section 2, Chapter II which rule out its use.

In Section 2, it was seen that unless certain assumptions are made, our assumed information is insufficient to determine the parameters of the Truncated normal probability density function. Hence, if we consider this density function, it must be with the assumption that the truncation is one-sided. Even after making this assumption, we are left with only an approximation to σ where

$$3.5 \sigma \leq R_p \leq 4\sigma \quad \text{or} \\ \hat{\sigma} \doteq \frac{\hat{R}_p}{3.5}, \text{ and unless}$$

$\hat{\sigma}$ is such that $\hat{h} \doteq \frac{1}{\sqrt{2\pi}\hat{A}\hat{\sigma}}$ where

$$\hat{A} = \int_{\hat{t}_0}^{\hat{t}_p} \frac{1}{\sqrt{2\pi}\hat{\sigma}} \exp\left\{-\frac{1}{2\hat{\sigma}^2}(t-\hat{t}_m)^2\right\} dt, \text{ then}$$

our assumptions will lead to contradictions with our data.

The differences from the above described equalities and inequalities upon which to base rejection is subject to the users latitude. In all cases, we are now discussing estimators rather than the true nature of R_0 , R_p , and h . The magnitude of a difference sufficient to cause rejection can be established only upon consideration of the accuracy of the estimates and the desired bounds for approximation.

For the remaining two density functions, given our assumed estimators, a probability density function can be cranked out which will fit the three desired points, but it will have a discrete mass point at t_p . Should the height of this mass point exceed h , then we fail to meet the criteria.

Any further attempts to determine which provides a "best fit" run head-on into the discouraging fact that we possess no other points with which to measure "fit". At this point, our excursions have reached a stage beyond which future progress can result only with more data. Under our hypothesized conditions, we cannot extract more data.

Since further theoretical progress seems barred, we shall now examine the origins of the problem to determine

whether the requirements of the problem and the bounds upon an acceptable solution can provide instruction as to future routes to pursue.

Basically, we are searching for a method of producing random time delays to assess against participants in a war game. These time delays represent the time which might be required for the accomplishment of the real life event which is being simulated. The only tools considered available to the assessor are a table of $\hat{t}_o$, $\hat{t}_m$, $\hat{t}_p$ for the various allowed events, and a uniform random number generator. This latter is a device, which upon activation, will generate a random number drawn from a population with a uniform distribution on the interval $[0 - 99]$. Specifically, we are attempting to provide for the assessor another tool--one which will convert such a random number to a random time T where $t_o \leq T \leq t_p$ and where T has a distribution roughly corresponding to the distribution of occurrence times for the real-life event.

There are three requirements for any method to be generally acceptable. It must be simple, rapid, and a reasonable approximation to reality.

War games, computer, manual or composite are extremely complicated models intended to simulate a most complex real-life situation. Any complication introduced into a component of play of the game must have its effect measured in the overall game atmosphere. The components do not exist in vacuo,

they exist to fulfill specific requirements for simulation. Unless relatively simple, the complication they introduce can over-balance the benefit of accuracy of simulation.

Present day sophisticated war games are lengthy affairs. A manual game may take as much as six months for completion. Computer simulations are accomplished in minutes or hours, but behind the play of the game is a programming effort that is measured in man-years. Any refinement of methods or introduction of new components can ill afford to further extend time requirements.

Most present war games utilize the method of fixed time delays, i.e., each event has associated with it a fixed time to be assessed for its accomplishment. Those that differ from this method generally apply uniformly distributed random time delays, with the implication that all times within the range for an event are equally likely. On no more than intuition, neither seems a very close approximation of reality. Any refinement of method should have some closer agreement to the way in which real times vary.

Having thus reviewed what it is we are trying to do, and the bounds on our method for doing it, we proceed to examine the effect of these bounds upon selection of a distribution function.

The uniform distribution is the simplest of our seven. The conversion of a uniformly generated random number

requires only a scaling to produce a random time. This ease of conversion makes any use of a uniform distribution quite rapid. It has the important disadvantage that in many circumstances it is not even a reasonable approximation to reality. One hardly expects that in very many events the minimum or maximum time will be observed with equal likelihood to any other time.

The triangular distribution is the next simplest. It possesses several features expediting work with it. The density function is completely determined by the extremes of the range and the position of the mode. For a given triangular distribution of T , $\hat{t}_o \leq T \leq \hat{t}_p$ (determined by the position of the mode), there is a corresponding distribution of x , $0 \leq x \leq 1$. By means of a table or graph of the applicable $[0 - 1]$ triangular cumulative probability distribution function, a uniform random number $[0 - 99]$ can be converted to a triangular random number $[0 - 1]$. This can then be scaled to the interval $[\hat{t}_o - \hat{t}_p]$ to provide a random time. This should give a closer approximation to reality than the uniform distribution. At least, times at the extremes of the range would have the least frequency of occurrence, and according to the position of t_m , any degree of skewness may be introduced. If we assume that the mode will occur at some discrete point $[0 - .99]$ we need only concern ourselves with 100 triangular distributions to

have a set of graphs or tables which can be used to handle any case. In fact, only 50 graphs will suffice since for distributions skewed to the right, we can use a corresponding distribution skewed to the left, taking 1 - the triangular random number, and scaling it to $\hat{t}_o \leq \tau \leq \hat{t}_p$.

The Beta distribution involves a greater number of calculations than the previous two distributions. Its parameters must be calculated from the data, and for a fixed set of values of $\hat{t}_o$, $\hat{t}_m$ and $\hat{t}_p$, there are a variety of curves which must be restricted to one unique curve by use of $\hat{h}$. That is to say that for a beta density function on $[0 - 1]$ and a given mode, there are many curves possible depending on the parameters. The distribution, once the parameters have been determined is conveniently handled with the assistance of Tables of the Incomplete Beta Function, (17). Again, by use of a table or graph of the cumulative probability distribution, a uniform random number $[0 - 99]$ may be converted to a beta number $[0 - 1]$ which can be scaled to $[\hat{t}_o - \hat{t}_p]$. Quite different from the preceding case is the number of tables or graphs required to handle the possible distributions. Since there is no unique distribution associated with each modal position, there are many graphs required for each such position. Intuitively, the Beta distribution may more closely approximate reality than the two previous distributions. Its height at the mode can

vary independently of R_0 , R_p and R , at least as long as $h^2 > \frac{1}{R_0 R_p}$. Depending on its parameters α and β the frequencies associated with the extremes of the range, particularly the maximum end will be less than with a triangular distribution over the same range.

Although not simple, the normal distribution is probably the best known of our seven. There exists a greater body of literature about it, and more readily available tables than for any other. Unfortunately, we have no assurance that we can use it in other than truncated form, which to some extent degrades the utility of existent literature and tables. For any event which has a skewed distribution, non-truncated use of it would represent a distortion of reality rather than an approximation to reality. For that reason, any further discussion of it will be continued in the Truncated Normal where the normal may be considered a special case.

The gamma and the lognormal distributions both have features which tend to complicate their use for our purposes. Both density functions are unbounded at their upper end. As was discussed in Section 2, there are three courses available in terminating the functions at t_p , truncation, consideration of a discrete probability mass point at t_p , and treatment of the functions as asymptotically zero at t_p . If truncation is adopted, we can no longer make convenient approximation to the parameters. If we adopt the idea of a discrete mass point at t_p , in generating random times, we will produce t_p more

frequently than some times less than t_p . The third alternative will, in general, produce the same effect. Since the idea of truncation seriously impairs our ability to produce the parameters, its use must be rejected. Under either of the two others we can proceed as follows. Prepare graphs or tables of the cumulative probability distribution function on a scale $[0-1]$. Use these to convert our random number to a random number on the appropriate $[0-1]$ lognormal or gamma distribution. Scale this number to $[\hat{t}_0 - \hat{t}_p]$ to produce T . For both distributions, the parameters depend only upon R_0 and h . If we consider that the mode will occur at some discrete point $\frac{\hat{R}_0}{\hat{R}}$ in $[0-1]$ we must still prepare many graphs for each such point to consider different possible values of $\hat{h}$. This, of course magnifies considerably the effort required for the use of either distribution. Both distributions have the disadvantage that they cannot readily be applied to any event which would seem to have a symmetrical distribution.

The Truncated normal distribution has the advantage that its use is facilitated by the extensively tabled normal distribution. If we consider for the present only truncations on the lower side, with the truncation at 0 and the height of the upper tail as approximately 0 at 1, then for the interval $[0-.99]$, only 50 tables or graphs are required to provide cumulative probability density functions for the 50 discrete

points $[0-.49]$ at which the mode could occur. This is due to the fact that at each position, we have but one curve to consider. Using these graphs, uniform random numbers can be converted to the appropriate truncated normal distribution, and scaled for use as T. For truncations at 1, with the lower tail height $\hat{=}$ 0 at 0, the same tables or graphs can be used, but the random number to be scaled would be $1 -$ the random number generated by the conversion.

Grouped in order of simplicity and rapidity in use, our seven functions would appear as

Uniform	}	Most simple and rapid
Triangular		
Truncated Normal (the normal a special case)		Intermediate
Lognormal	}	Least simple and rapid
Gamma		
Beta		

Comparison as to ability to approximate reality is not so easily done. In the first place what is reality? For all but a limited number of the simplest type events, no one knows. It is most likely that no one distribution best represents all classes of events. One form of distribution may best represent certain classes, but other classes would require other distributions. Experimentation has not been conducted nor is it likely to be conducted so as to comprehensively determine the distributions associated with the classes of events under discussion. We can however be sure that under certain conditions several distributions may all

be very close approximations. Even in the standard statistical problem with adequate sampling this frequently occurs. Using "goodness of fit" tests, two or more distributions are found to "fit" sample data.

Although we have no factual data as to reality, certain generalizations and intuitive conjectures can be made from observation and experience. In dealing with times required to complete a military operation, it seems plausible that many times (if not the majority of the time) one is working with an event requiring a skewed distribution. That is, the minimum time is closer to the most likely than is the maximum. There are only so many ways in which to expedite the operation, but it frequently seems that there are a million ways to degrade it. Also intuitive is a feeling that the probabilities of times near either extreme of the range should be very low. That is, at either extreme, $f(t)$ should approach zero asymptotically.

Excluding the uniform and the triangular, all of our density functions if generated by the same t_o , t_m , t_p and h tend to plot close together from the mode to the upper end point. From the lower end point to the mode, the truncated normal will, in general, plot higher than the remainder. All of them also conform to our previous conjectures to reality, with the exception of the truncated normal at the truncation end.

From the foregoing, we are reduced to selecting a distribution function on the basis of simplicity, speed in use, and general applicability. Those that are simplest and fastest are the least close to what we assume as reality. Those most approximate to our notions are also the more complex and laborious to use.

The next chapter of this thesis presents an algorithm for use in generating the random times desired. In this algorithm, we have elected to make use of the Truncated Normal Distribution as the most satisfactory compromise of speed and simplicity versus approximation to reality.

CHAPTER III

AN ALGORITHM TO GENERATE RANDOM TIMES

Section 1

This chapter describes an algorithm for manual use to generate random times. The input to the algorithm consists of a uniformly distributed random number, and the user's estimates of $\hat{t}_o$, $\hat{t}_m$ and $\hat{t}_p$. The output consists of a random time T such that

a. $\hat{t}_o \leq T \leq \hat{t}_p$ and

- b. T is derived from a truncated normal distribution, where truncation occurs on only one side and

$$3.5 \sigma = \text{Max} [\hat{R}_o, \hat{R}_p] \text{ with the mode at } \hat{t}_m.$$

If $\hat{R}_o = \hat{R}_p$ the distribution is normal, $\sigma = \hat{t}_m$.

$$3.5 \sigma = \hat{R}_o = \hat{R}_p.$$

In accomplishing the above, the algorithm performs the following functional steps:

1. Selects an appropriate truncated normal distribution scaled to the interval $[0 - 1]$
2. Converts the uniform random number to a truncated normal random number in the selected distribution.
3. Scales the truncated random number to $[\hat{t}_o, \hat{t}_p]$ to produce T .

The necessary elements to accomplish this are:

- a. A Uniform Random Number Generator

This is a device which upon activation produces a

random number drawn from a population which is uniformly distributed from 0 to 99. It may consist of no more than use of two columns from a table of random numbers, or it may be a mechanical/electrical device to produce such numbers.

b. A Table of Event Estimates

This is a user furnished table, listing alphabetically, by class, etc., the events occurring in simulation with their associated $\hat{t}_o$, $\hat{t}_m$, and $\hat{t}_p$.

c. Cumulative Probability Density Functions

These may be either graphs or tables, one for each possible case of the truncated normal distribution to be considered. Each graph represents a different possible position of the mode for discrete points in $[0.00 - 0.50]$. If the maximum (51) are used, there will be little significant change from graph to graph. In Section 4, it will be shown that the desired degree of approximation may be obtained by considering some number less than the maximum.

d. The Selector Scaler

This is a mechanical means for selecting the appropriate cumulative probability distribution function graph, and scaling a random number from that distribution to the interval $[t_o - t_p]$.

Construction of a selector-scaler should be within the production capability of any reasonably large training aids facility. If no selector-scaler is used, the necessary instructions for manually carrying out its functions will presently be explained. The selector-scaler consists of the following:

- (1) The 0-1 Bar - a horizontal bar with divisions from 0-99.
- (2) The Time Bar - a bar mounted by bracket, at a convenient angle to the 0-1 Bar, with 150 divisions, and such that it may slide through the bracket to place any desired dividing point coincident with the zero of the 0-1 Bar.
- (3) The Converter Grid - a ruled sheet of acetate or plexiglass with equidistant parallel lines whose spacing is equal to that of the divisions of the 0-1 Bar. The right most line is termed the Base Line and labelled B. Lines on the left half are black, and those on the right are red. Each line has a numerical designator.

A sketch of the components appears as Figure III-1.

The functioning of the algorithm is as follows:

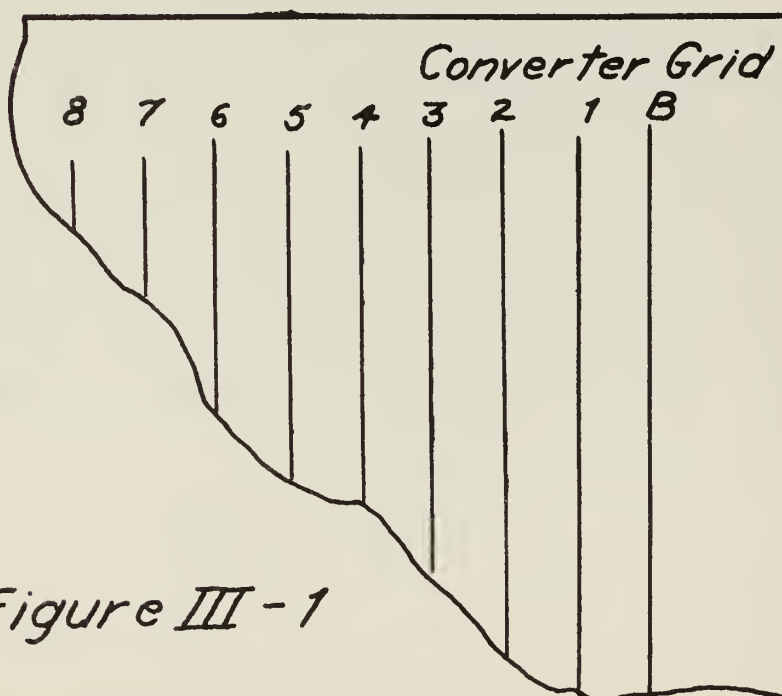
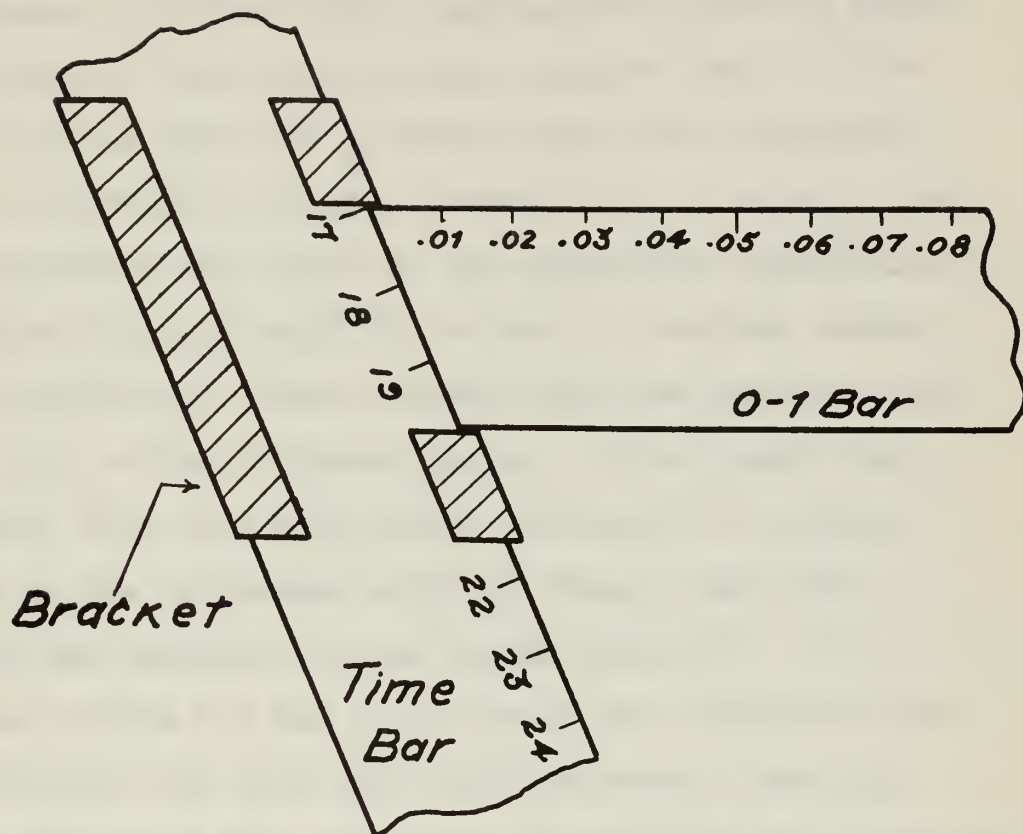


Figure III-1

When a random time is desired for an event simulation, the event is located in the Table of Event Estimates. The Time Bar is moved until t_0 on the Time Bar is coincident with the zero of the 0-1 Bar. The Converter Grid is placed over the bars so that the base line connects the 1 of the 0-1 Bar with the value of t_p on the Time Bar. The grid line on the converter passing closest to t_m is noted. Its number represents the number of the Cumulative Probability Distribution Function Graph to be used. A uniform random number is generated. Entry is made with this number on the vertical axis of the indicated graph. At the point where a horizontal from this entry point intersects the graph, the value on the horizontal scale is read. This value represents the truncated normal random number $[0 - 1]$. At this number on the 0-1 Bar some line on the Converter Grid is superimposed (the base line still connects 1 and t_p). The value that this line intersects on the Time Bar represents T.

SECTION 2

Instructions for Use of the Algorithm with the Selector Scaler

When an event occurs in simulation for which a random time delay is desired,

1. Locate the event in the Table of Event Estimates.
Note $\hat{t}_o$, $\hat{t}_m$ and $\hat{t}_p$.
2. Slide the Time Bar until $\hat{t}_o$ is coincident with zero on the 0-1 Bar.
3. Overlay the Converter Grid so that line B connects the end point (1) of the 0-1 Bar with the value of $\hat{t}_p$ on the Time Bar.
4. Locate $\hat{t}_m$ on the Time Bar, and note the color and number designator of the Converter Grid line crossing it.
5. Locate the C.P.D.F. graph corresponding to this designator.
6. Generate a uniform random number.
7. With this number enter the selected graph on the vertical axis. Where a horizontal from the entry point intersects the graph, read the value of the horizontal axis. If the Converter Grid line crossing $\hat{t}_m$ was red, subtract this value from 1. If the line was black, use the value as read.

8. Locate the number resulting from Step 7 on the 0-1 Bar. Follow the grid line crossing this point back down to the Time Bar, and read the nearest value. This is the desired random time.

Example: Suppose a random time delay is desired for an engineer platoon to install a six row mine belt across a valley 1500 meters wide. The Table of Event Estimates yields

$$t_o = 2 \text{ hours}$$

$$t_m = 8 \text{ hours}$$

$$t_p = 22 \text{ hours}$$

The number 2 on the Time Bar is placed coincident with 0 of the 0-1 Bar. The Converter Grid is placed so that line B connects the 1 of the 0-1 Bar and 22 of the Time Bar. The black line crossing t_m has the designator .30. Figure III-2 is C.P.D.F. graph #30. A uniform random number is generated, say .57. At .57 on the vertical axis of Figure III-2, a horizontal line intersects the graph where the nearest horizontal scale value is .35. At 35 on the 0-1 Bar, the grid line is traced back to yield 9 on the Time Bar. Since we are using hours, 9 hours represents the random time delay to be assessed.

SECTION 3

Instruction for Use of the Algorithm

Without the Selector Scaler

When an event occurs in simulation for which a random time delay is desired,

1. Locate the event in the Table of Event Estimates.
2. Compute $\frac{t_m - t_o}{t_p - t_o} = n$, $\frac{t_p - t_m}{t_p - t_o} = n^1$
3. Locate C.P.D.F. graph # $\min [n, n^1]$.
4. Generate a uniform random number.
5. Enter the vertical axis of the selected graph with the generated random number. Where a horizontal from the entry point intersects the graph, read the value of the horizontal axis. If the graph # was n^1 , subtract this value from 1. If it was n , use the value as read.
6. Calling the number resulting from Step 5, t , compute

$T = (t_p - t_o) t + t_o$ to determine a random time delay.

Using the example of Section 2,

$$\frac{t_m - t_o}{t_p - t_o} = \frac{6}{20} = .3 \quad \frac{t_p - t_m}{t_p - t_o} = .7$$

Again using graph # 30, Figure III-2 with a uniform

random number assumed to be 57 the value of t
is .35

$$T = .35(22 - 2) + 2 = 9 \text{ hours.}$$

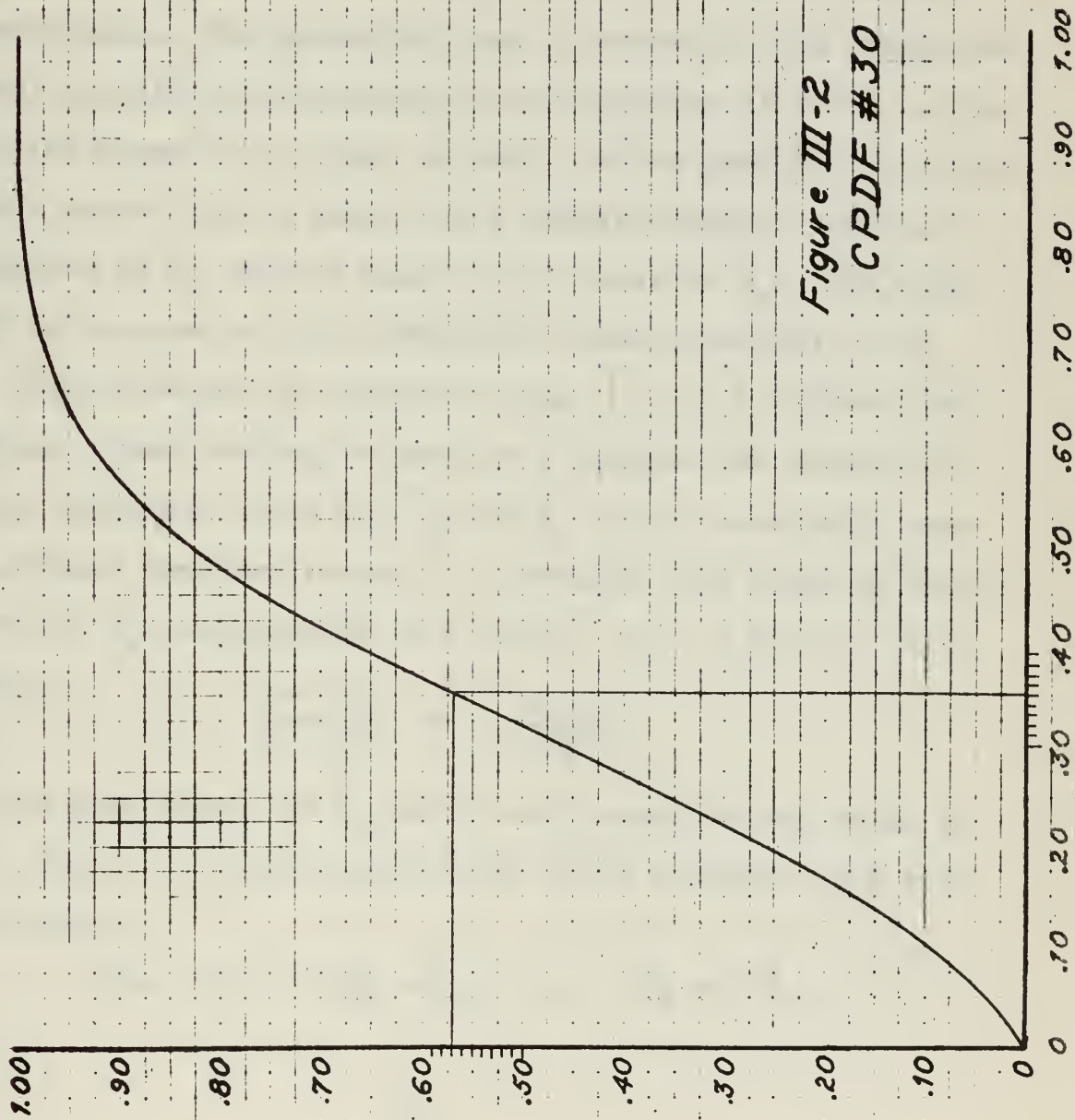


Figure III-2
CPDF #30

SECTION 4

Construction of Cumulative Probability

Distribution Function Graphs

This algorithm makes use of graphs of the cumulative probability density function (C.P.D.F.) of various truncated normal distributions to convert a uniform random number (0-99) to a random number from a desired truncated normal distribution. For generality and convenience, the truncated normal density is considered to be truncated at zero, and to have its upper 3.5 σ point at one. We are actually concerned with a random time T which has a similar density function truncated at $\hat{t}_0$ and its upper 3.5 σ point at $\hat{t}_p$. The mode (and so the mean of the underlying normal density) is at $\hat{t}_m$. This case can be reduced to the $[0 - 1]$ interval by scaling. Such scaling reduces to a minimum the number of graphs required, since $\hat{t}_0$, $\hat{t}_m$ and $\hat{t}_p$ could conceivably take on any real positive values. In scaling, any point t , where $\hat{t}_0 \leq t \leq \hat{t}_p$, corresponds to a point x , $0 \leq x \leq 1$, by letting

$$x = \frac{t - \hat{t}_0}{\hat{t}_p - \hat{t}_0} = \frac{t - \hat{t}_0}{\hat{R}}$$

Now the mode occurs at $\hat{t}_m$ which could occur at any value in $[\hat{t}_0, \hat{t}_p]$. It will suffice for us to consider only such cases where

$$(\hat{t}_m - \hat{t}_0) \leq (\hat{t}_p - \hat{t}_m) \quad \text{or} \quad \hat{R}_0 \leq \hat{R}_p$$

i.e., where truncation occurs at $\hat{t}_0$. Cases where $\hat{t}_0$ is the lower 3.5 σ point, and $\hat{t}_p$ is the point of truncation can be handled using the same set of graphs, but with slightly different application.

If we desire a set of graphs to cover all possible cases, we then require a graph for each possible position of the mode. In the interval $[0 - 1]$, this would imply a graph for each possible value of

$$\frac{\hat{t}_m - \hat{t}_0}{R} = \frac{\hat{R}_0}{R} \quad \text{where } \hat{R}_0 \leq \hat{R}$$

Our uniform random numbers can only take on values in (0-99), and the truncated normal random numbers we will use will be from the discrete set of points in (0.00 - .99), so that really we are only interested in values of $\frac{\hat{R}_0}{R}$ which are discrete points in (0.00 - 0.49). This would imply that a maximum of 51 graphs would be needed to cover all possible cases in which we have an interest.

Figure III-3 is a set of graphs of the C.P.D.F. for truncated normal distributions as described, and for $\frac{\hat{R}_0}{R}$, the mode, equal to 0.00, 0.10, 0.20, 0.30, 0.40, and 0.50. An examination of the curves indicates that as the mode approaches zero, the successive curves lie closer together. As an example, if we had a uniform random number of .5, and converted it using curve #6, we would get a truncated random number of .19. If we used graph #5, we would get .22. Now graphs #5 and 6 have very near maximum separation of the point used

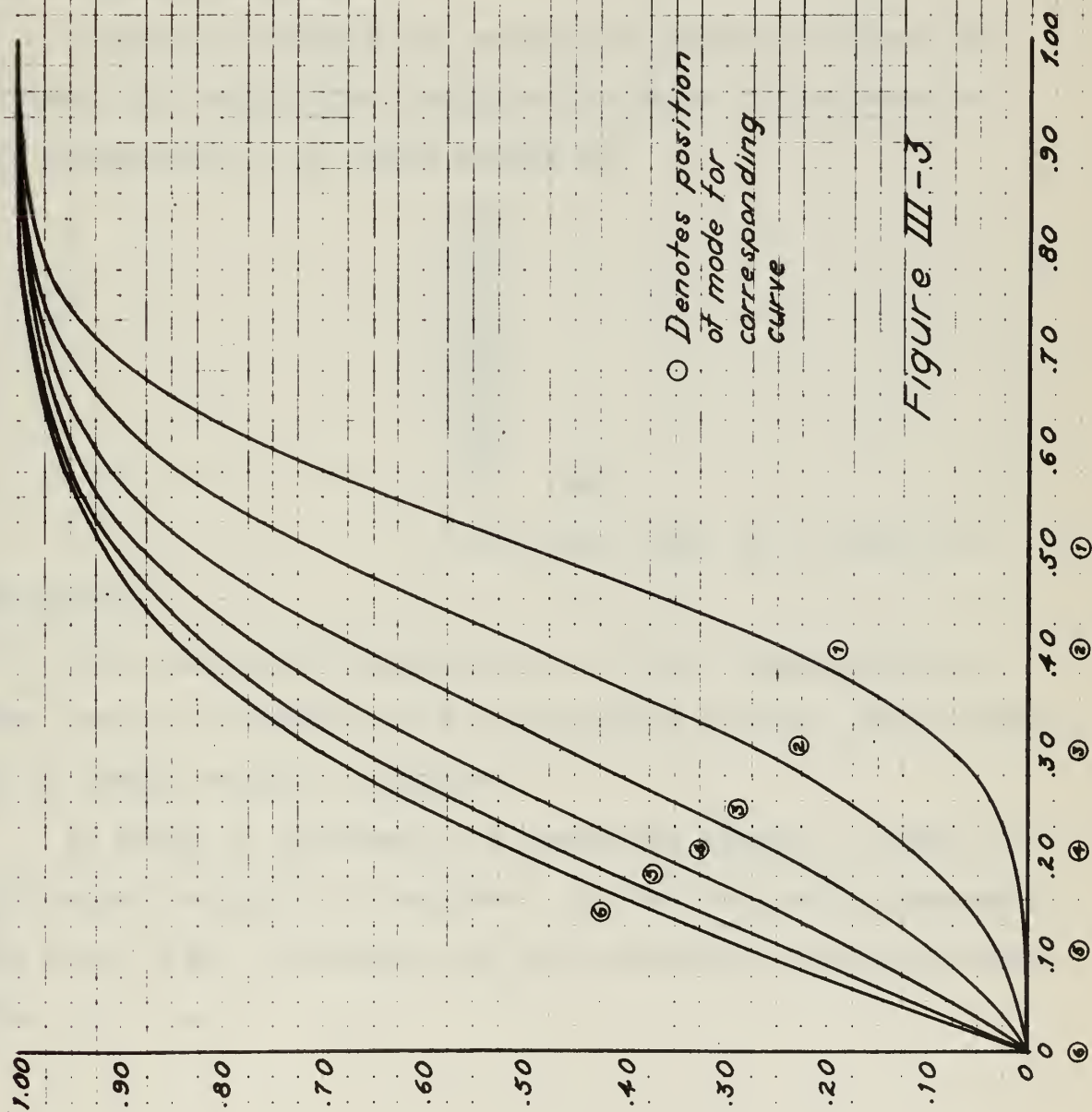


Figure III-3

in the example. This means that for an interval that covers 10 possible positions of the mode, the resolution in random number conversion is 3. Apparently, for this interval, we do not need graphs for all possible positions of the mode. In fact, the graphs for 0.00, 0.05 and 0.10 should be adequate for our approximations.

A similar analysis for successive pairs of curves indicates that we can get adequate coverage of the cases we are interested in by using graphs for

0.00
0.05
0.10
0.13
0.16
0.20
0.22
0.24
0.26
0.28 and

0.30 (.01) 0.40, or a total of 30 graphs.

If we are only interested in a "fair" approximation, then the use of every other curve should suffice, and a total of 15 graphs would be adequate.

In order to construct the necessary graphs, a table of the normal integral is required. In the following discussion, the term $\Phi(a)$ represents the area under the normal integral from $-\infty$ to a , i.e.

$$\Phi(a) = \int_{-\infty}^a \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt \quad \text{and}$$

$$\int_a^b \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt = \Phi(b) - \Phi(a)$$

A property of the normal integral is that

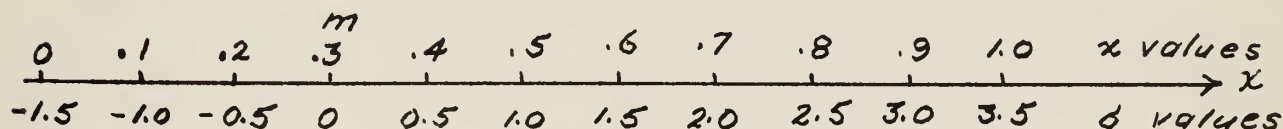
$$\Phi(-a) = 1 - \Phi(a)$$

The mean (μ) of the normal integral must be zero to use the tables in customary form. Therefore, we will perform our computations in such a manner as to effectively consider $\mu = 0$. Values less than the mode (which is equal to μ) will therefore be considered as negative. We can accomplish this by measuring all our distances in terms of σ and measuring from the point which is the mode. Under our assumptions, we determine σ by

$$3.5\sigma = 1 - m \quad \text{where} \quad m = \frac{R_0}{R}$$

$$\text{or} \quad \sigma = \frac{1-m}{3.5}$$

If our horizontal axis is X , then all values on X can be scaled to some value of σ , with 0 at m on X . If for example, $m = .3$ on X we would have



Next, we determine A, the area under the normal density from our lowest σ value (equal to zero on the χ scale) to our highest σ value (equal to 1 on the χ scale). In the example above

$$A = \int_{-1.5\sigma}^{3.5\sigma} \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt = \Phi(3.5\sigma) - \Phi(-1.5\sigma)$$

$$A = \Phi(3.5\sigma) + \Phi(1.5\sigma) - 1$$

Using a table of the normal distribution,

$$A = .9330$$

For successive points $t \cdot \sigma$ corresponding to 0.00, 0.05, 0.10, etc., on the χ scale, compute using the table,

$$\frac{\Phi(t\sigma) - \Phi(-1.5\sigma)}{A} = \frac{\Phi(t\sigma) + \Phi(1.5\sigma) - 1}{A}$$

The values obtained are points on the C.P.D.F. curve at the corresponding value of χ , and can be used to draw the desired curve.

SECTION 5

Extensions to the Algorithm

Throughout the development of the algorithm continued efforts have been exerted to develop a tool with sufficient flexibility that the user can best fit to the individual needs of his organization. Obviously there will be cases where it is impractical or impossible to secure professional estimates of t_o , t_m and t_p .

Many war games evaluate weapons systems that are still in the process of development. In these cases little more than a vendor's estimate of the equipment's performance times is available. In other cases new techniques and tactics are evaluated in areas where there is no reservoir of experience to offer assistance in establishing the appropriate delays for performance or reaction times.

On the surface it would seem that in this situation the war gamer has little choice but to resort to arbitrary or fixed time delays.

With the advent of ever more sophisticated weapons, the ability or inability of an opponent in a war game to use a particular piece of equipment in a given situation can conceivably affect the outcome of the entire game. Under such circumstances, the controllers of the game must consider every avenue that affords a closer approximation to reality.

To restate the underlying theme of this thesis; "there

exist situations where the use of random time delays is not a technical nicety, but a practical necessity to insure validity."

How could the algorithm of this thesis be used in situations such as those described above. Let us devise an example.

Suppose that we as war gamblers are conducting a game in which one of our perimeter defense weapons is an as yet unevaluated mobile rocket launcher. Let us further suppose we have only the contractor's estimate of the average time to put this weapon into operation; say it is six hours. It would indeed be a rare event if this estimate was accurate. But still, some estimate must be accepted as the best available at this stage of development.

One can rather easily decide on some most optimistic time short of which, due to physical limitations, the set-up of this equipment is virtually impossible. Let us assume this time (t_o) is four hours. In a similar fashion we can establish our most pessimistic time (t_p) at say 12 hours. Lastly, we can accept the contractor's estimate of t_m as six hours or adjust it as our common sense dictates. We now have the necessary estimates to produce a random time delay of between 4 and 12 hours where, of course, the bulk of our generated times will tend to fall near t_m .

SELECTED BIBLIOGRAPHY

WAR GAMING, HISTORY AND BACKGROUND

1. J. P. Young, A Survey of Historical Development in War Games, ORO, APL, The Johns Hopkins University, Washington, D. C., 1959.
2. H. Kahn and I. Mann, War Gaming, Rand Corporation, Santa Monica, California, 1957.
3. R. D. Sprech, War Games, Rand Corporation, Santa Monica, California, 1957.
4. A. M. Mood and R. D. Sprech, Gaming as a Technique of Analysis, Rand Corporation, Santa Monica, California, 1954.

WAR GAMING, SPECIFIC APPLICATIONS

5. M. G. Weiner, Use of War Games in Command and Control Analysis, Rand Corporation, Santa Monica, California, 1961.
6. M. G. Weiner, War Gaming Methodology, Rand Corporation, Santa Monica, California, 1959.
7. H. E. Adams, Carmonette, ORO, APL, The Johns Hopkins University, Washington, D. C., 1961.
8. J. O. Harrison and E. M. Lee, Stratspiel Pilot Model, ORO, APL, The Johns Hopkins University, Washington, D. C., 1960.
9. L. H. Sell and C. M. Hynden, First Status Report - Amphibious Warfare Simulation Model, NWL, 1959.
10. O. Helmer and R. E. Bickner, How to Play Safe, Rand Corporation, Santa Monica, California, 1961.
11. M. C. Waddel, Air Battle Analyzer, APL, The Johns Hopkins University, Silver Spring, Maryland.

WAR GAMING, BIBLIOGRAPHY

12. War Games and Gaming. Selected References. Air University Library, Maxwell AFB, 1959.

13. History and Bibliography of War Gaming, ORO, APL, The Johns Hopkins University, Washington, D. C., 1957.

PROBABILITY AND STATISTICS

14. Parzen, Emanuel. Modern Probability Theory and its Applications. New York: John Wiley and Sons, Inc., 1960.
15. Mood, A. M. Introduction to the Theory of Statistics. New York: McGraw-Hill Book Company, Inc., 1950.
16. Pearson, Karl. Tables of the Incomplete Gamma Function. London: Cambridge University Press, 1922.
17. Pearson, Karl. Tables of the Incomplete Beta Function. London: Cambridge University Press, 1932.
18. Kenney and Keeping. Mathematics of Statistics. Part II. New York: D. Van Nostrand Company, Inc., 1951.
19. Lingren, B. W. Statistical Theory. New York: The Macmillan Company, 1962.
20. Kendall, M. G. Advanced Theory of Statistics. Vol. II. Third edition. New York: Hafner Publishing Company, 1947.
21. Siegel, S. Nonparametric Statistics. New York: McGraw-Hill Book Company, Inc., 1956.
22. Elderton, W. P. Frequency Curves and Correlation. Fourth edition. Washington, D. C.: Harren Press, 1953.

APPENDIX 1

THE AIR BATTLE ANALYZER

The Air Battle Analyzer was devised by Dr. M. C. Waddel, Applied Physics Laboratory, Silver Spring, Maryland. It has recently been used by the United States Marine Corps in a study undertaken to evaluate the Air Defense capabilities of the Marine Corps Expeditionary Force.

The Air Battle Analyzer is a scenario type, free play war game in which each opponent is given a specific capability and plan of action. It attempts to provide, and display, a readily accessible means for recording and displaying chronologically the principle movements and operations of the different surface and airborne units involved in a battle. This display points up the interactions between different units, in particular their ordering in time.

The physical equipment associated with the Air Battle Analyzer consists of a plotting and display chart, several tools for measuring and plotting on the displays, and several nomographs.

The chart combines three display plots, a range-azimuth plot, a range-altitude plot, and a range-time plot. The three plots have a common (horizontal) range scale with range origins on a common vertical, to facilitate reference from one plot to another. This combination of plots allows for the simultaneous visualization of the elements of a

hypothetical air/ground or sea battle in the four dimensions of range, bearing, altitude and time.

Of the various tools and nomographs associated with the Air Battle Analyzer only three will be described as being germane to this discussion; the Fire Power Analyzer, the Range of First Detection nomograph and the Target Speed Scaling nomograph.

FIRE POWER ANALYZER

A precise estimate of firepower of a surface-to-air missile system requires detailed examination of the particular system operation. The cycle time between successive intercepts, and the intercept envelope depend, often in a rather involved manner, upon a large number of system parameters. However, a reasonable understanding of missile system firepower can be obtained from consideration of a hypothetical system. To this end, a simple missile system employing one launcher and one director has been defined.

The Firepower Analyzer is a tool for estimating the firepower of a launcher defined as above, employing a 100 n.m. missile, and engaging Mach 1.0 aircraft in a wave attack. This speed was chosen arbitrarily. When the Analyzer is placed on the range-azimuth plot of the display chart with center at the missile ship and axis parallel to the attack path, the various curves of the Analyzer intersect the attack path at successive intercept points. To determine

the firepower for any maximum intercept range of 100 n.m. or less, first determine the points of first and last intercept along the attack path. The first of these will be limited by maximum missile range or by range of detection, etc. The latter point is determined by weapon release point. To estimate the single launcher system firepower, place the Analyzer as indicated above and count the intercepts from 100 mile range in to the last intercept point, and then subtract those occurring beyond the assumed first intercept point. In computing the difference, that is, in counting intercepts, it is appropriate to interpolate for fractional intercepts, inasmuch as the computation is directed toward estimating the expected number of intercepts. To the above difference, one more intercept should be added; the one occurring at maximum intercept range, to obtain the total firepower.

FIRST INTERCEPT RANGE

The range of first intercept by a surface-to-air missile system may be limited by target detection range and subsequent delay for decision, etc., rather than by the maximum range of the missile. In this event the range of first intercept is the detection range less the distance traveled by the target during the delay and the missile time of flight, and so is dependent on target speed.

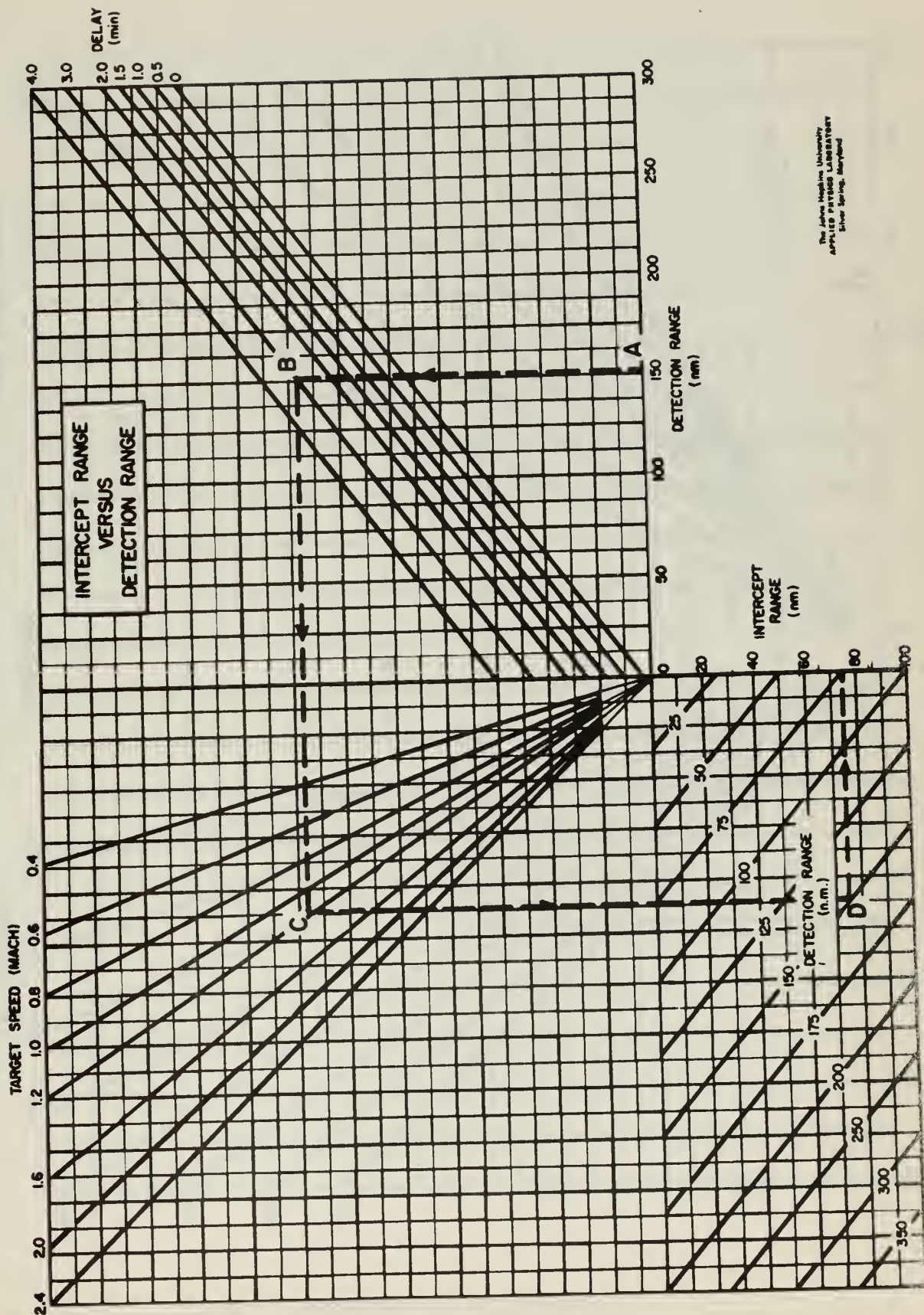
This game uses a nomograph, enclosure 1, to facilitate the calculation of first intercept range. Enter the right

hand abscissa with the detection range, move vertically upward to the appropriate delay; then move horizontally to the left to the proper target and speed; next move vertically downward to the same detection range, and read off first intercept range from the lower ordinate. The dashed lines of the figure trace the above path for a detection range of 150 n.m. (A), delay of 3 min. (B), and target speed of Mach 1.2 (C), giving first intercept range of 74 n.m. (D and E).

TARGET SPEED SCALING

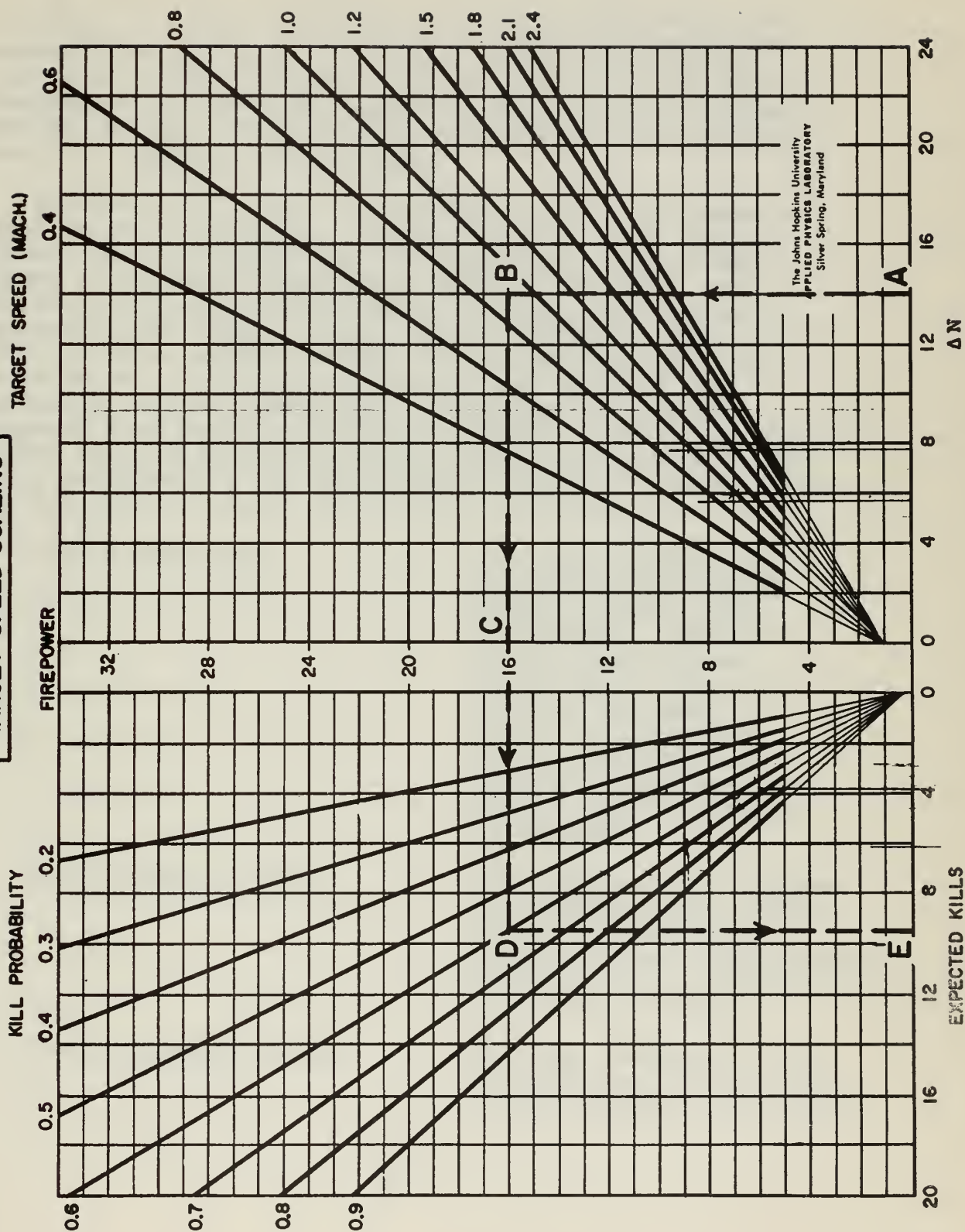
The firepower (expected number of intercepts) of a surface-to-air missile battery depends critically upon target speed inasmuch as fast targets will pass through the missile envelope in shorter time and so be faced with fewer missiles than slower targets. To facilitate scaling to various target speeds the Air Battle Analyzer uses a nomograph, enclosure 2. Enter the right hand abscissa with the number of intercepts (ΔN) obtained from the Firepower Analyzer, move vertically upward to the desired target speed; then move horizontally to the left, reading off the firepower from the ordinate. The dashed lines of the figure trace the path for a difference $\Delta N = 14$ (A), and target speed of Mach 0.9 (B), giving firepower of 16 (C). The left hand portion of the nomograph permits ready conversion of firepower to an estimate of expected number of kills. From the firepower reading on the ordinate, continue to the left to the appropriate kill

probability, then move vertically downward, reading off the expected kills from the left hand abscissa. Again referring to the dashed lines of the figure, firepower of 16 (C) and a kill probability of 0.6 (D) give expected kills of 9.6 (E).



The Johns Hopkins University
APPLIED PHYSICS LABORATORY
Silver Spring, Maryland

TARGET SPEED SCALING



APPENDIX 2

SAMPLE QUESTIONNAIRE

You are asked to furnish your best independent, professional estimate of the time it takes to accomplish certain assignments. The result as furnished by you and other Marine officers will be used in an attempt to determine whether such a method can yield more accurate estimates than are currently provided by PM's, TM's and various planning publications.

Enclosure (1), page 2, lists the events and space to fill in your estimates. Should you feel that certain events are so specialized in nature that you, in the light of personal experience, duty assignments, etc., cannot give a reasonably valid estimate, you may omit estimates for those events. It is recognized that each of these events depends on a wide variety of circumstances. No attempt to completely describe the situation can be successful. Therefore, it is requested that you make the best estimate possible in the light of your interpretation of each problem posed. The time required for the following is offered in explanation of the various times you are asked to estimate.

- | | | |
|-----------------------|---|---|
| Most likely time | - | The time required to accomplish the action under "normal" circumstances. "Normal" to be construed as those day to day conditions you would expect to encounter under circumstances such that the described action could logically be assigned |
| Most optimistic time | - | The minimum time required to accomplish the action of all circumstances, personnel reaction, weather, terrain, etc., combined to work in your favor. |
| Most pessimistic time | - | The time required to accomplish the action under extreme adverse conditions, |

and poor response on the part of the participants. The conditions are not so severe as to preclude accomplishment of the mission, but weather, fatigue, human error, etc., all exercise strong degrading effects.

Note: Unless otherwise indicated, times are to be from receipt of assignment until accomplishment, and should reflect planning, reconnaissance, and coordination requirements.

Estimates are desired to the nearest 15 minutes unless otherwise indicated.

Percentage of times the event will occur within _____ hours/minutes of the most likely time.

If the event were repeated a large number of times, this is your estimate of the per cent that require a time within +/- hours/minutes of your estimate of the most likely time.

1st Event

Reserve battalion to organize and occupy combat outpost to include preparation of hasty defense and preregistration of supporting arms. Assume foot march time from reserve position to OP site is 2 hours.

Most likely time _____

Most optimistic time _____

Most pessimistic time _____

Percentage of times event would occur within 30 minutes of the most likely time _____

2nd Event

To mount out "ready" B LT, from Marine Corps Base to afloat status.

Most likely time _____

Most optimistic time _____

Most pessimistic time _____

Percentage of times event would occur within 60 minutes of
the most likely time _____

3rd Event

Infantry company to conduct physical fitness testing using
current test. Time from first man starts until last man
finishes. Assume no parallel facilities.

Most likely time _____

Most optimistic time _____

Most pessimistic time _____

Percentage of times event would occur within 30 minutes of
the most likely time _____

4th Event

To heli-lift rifle company in vicinity of Helipad. Five
men heliteams. Pad capacity - 3 helicopters. Sufficient
helicopters airborne in vicinity to move company. Time from
touch down of first flight until last flight is airborne.
(Time to nearest 5 minutes)

Most likely time _____

Most optimistic time _____

Most pessimistic time _____

Percentage of times event would occur within 5 minutes of the
most likely time _____

5th Event

Reserve rifle company to provide squad size daylight
reconnaissance patrol. Time from receipt of order to departure
of patrol, but to include briefing, rehearsal, and check.

Most likely time _____

Most optimistic time _____

Most pessimistic time _____

Percentage of times event would occur within 15 minutes of
the most likely time _____

6th Event

Time to establish Infantry Battalion CP. Time from arrival
of first elements at CP site until CP is established and
functioning adequately.

Most likely time _____

Most optimistic time _____

Most pessimistic time _____

Percentage of times the event will occur within 15 minutes
of the most likely time _____

7th Event

Time to prepare one platoon from an FMP unit to conduct
formal guard mount, to include as you feel might be required,
classroom training, rehearsals, and preliminary inspection.

Most likely time _____

Most optimistic time _____

Most pessimistic time _____

Percentage of times the event will occur within 30 minutes
of the most likely time _____

8th Event

Time to establish one platoon road block in hasty defense,
800 meters forward of parent company's MLR position.

Most likely time _____

Most optimistic time _____

Most pessimistic time _____

Percentage of times the event will occur within 15 minutes
of most likely time _____

9th Event

Time to establish field mess. Time from arrival of personnel and equipment until mess is set up and ready to feed "A" type rations.

Most likely time _____

Most optimistic time _____

Most pessimistic time _____

Percentage of times the event will occur within 15 minutes
of most likely time _____

10th Event

You are a Forward Air Controller directing a live ordnance air exercise on Browns Island. You are also acting as Division Air Officer. Assuming that you are booked in by radio net to a group headquarters at Cherry Point. What is your estimate of the time required from time of radio contact with the a/c group headquarters until a flight of four jet a/c are on target under your direction as FAC. Assume the group has alerted a squadron to have pilots and armed a/c standing by. Assume a ten minute pilot briefing by the squadron, and consider ordnance plug in. (To nearest 5 minutes)

Most likely time _____

Most optimistic time _____

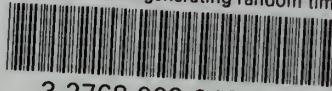
Most pessimistic time _____

Percentage of times the event will occur within 15 minutes of
the most likely time _____

Your rank _____ Specialty _____ Years of Service _____

thesM6685

An algorithm for generating random time



3 2768 002 04686 4

DUDLEY KNOX LIBRARY