

PAPER NUMBER

60-MD-1

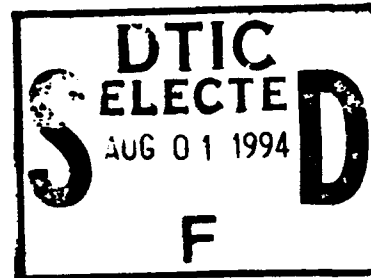
03  
08  
4  
08

COPY 1

# Reliability Prediction—Its Validity and Application as a Design Tool

1-84831  
THOMAS C. REEVES

Defense Electronic Products  
Division, Radio Corporation  
of America, New York, N. Y.



Introductory survey of reliability prediction. This paper is developed in three main parts: Introduction to reliability prediction, what it consists of and the conditions which must accompany its intelligent use; validity of such predictions; and uses of reliability predictions in operations analysis, maintenance and logistic studies, design and value analysis.

This document has been approved  
for public release and sale; its  
distribution is unlimited.

1960

DTIC QUALITY INSPECTED 8

94 7 18 036

Contributed by the Machine Design Division for presentation at the Design Engineering Conference and Show, New York, N. Y., May 23-26, 1960, of The American Society of Mechanical Engineers. Manuscript received at ASME Headquarters, February 10, 1960.

Written discussion on this paper will be accepted up to June 27, 1960.

Copies will be available until March 1, 1961.

# Reliability Prediction—Its Validity and Application as a Design Tool

THOMAS C. REEVES

Although the Greeks are not known to have advanced very far in the fields of electronics and space technology, they did - as usual - have a word for what this paper is about. The word is prognosis, meaning to know beforehand, to predict the outcome of an event before it takes place. This paper deals with the prognostication of the reliability of complex equipments before these equipments are built or even designed.

For the time being, reliability will be loosely defined as the measure of our certainty that the equipment being developed will ultimately do what it is supposed to do when called upon to do it. When so defined, one can appreciate that every product engineer - regardless of his field - faces the problem of reliability prediction.

In principle, therefore, reliability prediction is not new nor is it uncommon. Every product engineer in a sense makes a reliability prediction every time he signs off a drawing; in effect, he predicts that the design he is releasing has a high probability of doing its intended job. Any engineer who has ever made a stress analysis was actually making a very old and basic kind of reliability prediction. So reliability prediction is not new. What is new would seem to be the increased emphasis on time as the critical, dependent variable in a stress relationship. And even this perspective is not new to anyone who has had to design structures or machine elements subject to fatigue or wear.

As a result, the product engineer and especially the mechanical designer and stress analyst should find themselves comfortably at home - philosophically at least - with reliability prediction as discussed in this paper. They will also recognize that all we are doing that is basically new is dealing with large and complex systems made up of electrical and electromechanical elements and that we are looking at these parts from the point of view of their life expectancy under electrical stresses, such as voltage, as well as under mechanical stresses, such as created by thermal environment.

It is hoped, therefore, that by the time we reach the end of this paper, the reader (regardless of his field of engineering or product line) will not only understand what reliability prediction consists of and what it does, but that he will also see its place in the bundle of engineering design tools needed today.

To reach this objective, the paper has been developed in three main parts. The first part is an introduction to reliability prediction; what it consists of, and the conditions which must accompany its intelligent use. In the second section the validity of such predictions will be examined. In the final section will be outlined some of the

uses of reliability predictions in operations analysis, maintenance and logistic studies, design and value engineering.

While these examples will be presented in the context of the complex systems common to space technology and weapon systems, it is the hope that the reader will be able to draw fruitful analogies with his own fields and similar problems and requirements.

This paper is frankly an introductory survey of reliability prediction intended for the engineer who is not acquainted with the subject. In the interests of brevity, the mathematics and statistics will be minimized and offered without proof. Those who desire to actually apply these techniques to their own fields will certainly want to examine these aspects more rigorously in the source documents cited.

## 1 - ELEMENTS OF RELIABILITY PREDICTION THE GENERAL RELIABILITY PROBLEM

It is evident why it is essential to have reliable devices in the fields of space technology and weapon systems and what happens when these devices are not reliable. To know that the failure of a 50-cent part can lead to loss of control and subsequent destruction of a vehicle costing millions in time and skill is to understand the painstaking detail related to reliability programs.

To achieve reliability in such complex systems requires extremely high reliabilities for all parts which can cause system failure. But success depends on more than simply having more reliable parts; it depends on being able to design the system so that the inherent reliability of good parts is not compromised by misapplication. It also depends on building the system using processes and workmanship which will not degrade the parts as they are integrated into equipments and subsystems. Success depends also on exhaustive quality control and tests, in-house and out, to minimize defects and to permit their prompt diagnosis and remedy. Success also depends on intelligently planned maintenance up to the time of use and certainly during use.

### Necessity For Designing Reliability Into Products

The achievement of reliability is thus not just a matter of design and certainly not just a matter of being able to measure or predict reliability. On the other hand, the initial design phase does largely determine the shape of things to come. It is necessary, in the design phase, to evaluate reliability so that if necessary something constructive can be done about it before design release or cer-

A-1

tainly before construction and field test of prototypes. There is just no time in today's programs for achieving reliability by trial and error. Not only is there no time but the semi-public demonstration of space failures carries penalties beyond the technical ones in the sense of damaging national prestige. It is essential that there be a high level of confidence that they will work first time out.

#### Necessity For Continuous Reliability Evaluation

As a result, the design of such systems is usually conducted within the framework of a comprehensive reliability discipline which calls for continuous evaluation of reliability. One usually starts to design the system to realize a certain specified probability of operational success and this requirement is continually compared with the reliability expectations of the evolving design. Design changes are then made as necessary to reconcile the two. In a sense, it is still a process of identifying errors and of correcting for them but in this process of control the feedback is continuous throughout design rather than the one-shot feedback after design that characterizes trial-and-error reliability improvement.

Reliability control in design thus calls for the ability to make continual evaluation of the reliability of the product throughout the design cycle. The earlier such an evaluation can be made the more valuable it is in terms of permitting corrective action with minimum disturbance. The earliest reliability evaluation should be made on the proposed design as it exists on paper or even as a gleam in the designer's eye. This is why reliability prediction is so important a part of the over-all reliability program.

Procedures for predicting the reliability of complex systems have been developed for the most part within the past 10 years. While procedures in current use vary somewhat from one design organization to another, they have basic similarities in their premises and rationale, in the computational routines and in the end results and further uses of the results. The following description will reflect RCA's procedure, because, first of all, the author is most familiar with it and especially with its validation to date, and then because working-level handbooks on the RCA procedure are more readily available than those for the other procedures.

#### Unreliability and Failure

Before launching into a description of the prediction procedure itself, a few basic concepts and characteristics will be presented as groundwork.

First of all, what is unreliability? It is a measure of a lack of dependability to perform properly when needed. When a device doesn't perform as it should, it is said to have failed. A failure does not have to be catastrophic in the sense of meaning complete, irreparable destruction of a part; a failure may comprise only minor performance degradation which requires only a slight readjustment for the part to be restored to service.

Failure is therefore defined as an occasion when per-

formance is no longer within specifications and usually requires some adjustment, maintenance or replacement to restore performance. Failures also reflect an element of embarrassment and surprise. Instances of planned maintenance and periodic adjustment are not considered to constitute failures if the interim performance is within specifications.

This definition of failure is very general. It is not surprising that many categories of failure are recognized in the study of reliability and we must be careful of what kinds of failure we are talking about -- particularly in reliability prediction. Six categories will be cited here, according to the cause of failure:

1 Parts fail because they wear out, as a process of deterioration in use; for instance, brushes on a motor. But such parts can be replaced before they wear far enough to cause failure. Therefore, one can minimize the probability of wear-out failures by adequate inspection and preventive maintenance.

2 Parts fail because they are initially defective; that is, incoming inspection has not been keen enough to catch all defective parts and some of them get into the product. Such parts are not always defective to the point of not working at all. Frequently they are just weak; good enough to pass inspection but weak enough to fail just after the product is accepted and gets into service. Rigid incoming inspection combined with proof-stressing or burn-in can weed out such parts and minimize the failures they cause.

3 Otherwise good parts fail because they have been damaged by poor workmanship during installation into the end product. Adequate quality control and acceptance testing can detect such workmanship defects.

4 Good parts even if properly assembled into the product can fail because of improper application, because of being overstressed or called upon to perform tasks they were never intended to do. This is a design error which can be caught by design review and which will, in any event, be revealed by adequate testing in terms of repetitive failure of the part in question.

5 Failures can be caused by gradual performance deterioration. In this case the part is not wearing out but is drifting out of initial setting and requires adjustment. Like wearout, such failures can be minimized by properly scheduled inspection and adjustment and should not result in unplanned failures.

6 Some parts fail as a direct consequence of the failure of other parts. In such cases the failure of the first part has imposed greater stresses and damage on the second part causing it either to fail immediately or later. Secondary failures also can be caused by accidental damage to other parts when repairing the primary failure. Such secondary failures can be minimized by intelligent inspection and intelligent replacement.

Note that in each of these categories of failure, there was some means of detecting its incipency by inspection or implication and some means of prevention by replacement and adjustment. Clearly if this were true of all

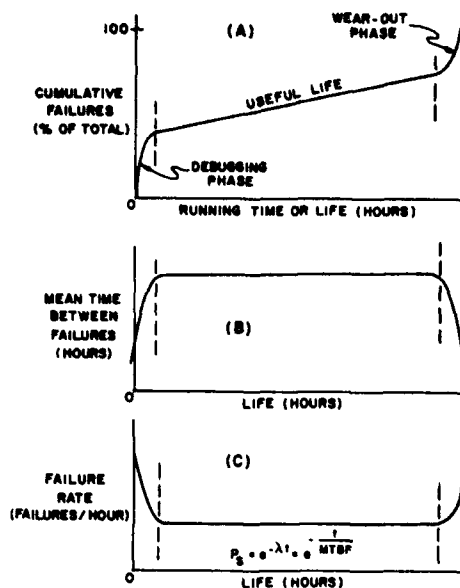


Fig. 1 Generalized life model for reliability prediction

failures there would be no unexpected failures and one would achieve reliability merely by inspection and maintenance.

Unfortunately, there is one category of failure which cannot be eliminated by inspection or tested for by any present means because it occurs without the warnings and clues present in all the rest. Something certainly causes it and the cause can often be traced after failure but its incipency cannot be detected before failure. This characterizes the true random failure and by definition, it is the type of failure that cannot be anticipated or prevented by inspection because it has no recognizable symptoms before the fact. Clearly, this random-failure category, unforeseeable and unpreventable as it is, is the most pernicious of all. It receives the bulk of reliability emphasis and takes the blame for the bulk of failures.

#### Life Model For Reliability Prediction

The foregoing failure categories are also characteristic of various phases in the life cycle of a product. For prediction purposes, the life cycle or model most generally assumed as applying to large complex systems is depicted in Fig. 1. Fig. 1(a) shows what might be expected if the total number of failures for a large number of identical systems were plotted against the hours such systems accumulate in life, with life time beginning at the end of the assembly line. Shortly after time zero, as various performance and acceptance tests are run, failures pile up in relatively rapid order. This reflects the identification and correction of workmanship errors, serious design errors, necessary realignments and adjustments, and so on. As these initial defects are remedied and the equipment becomes "debugged," the rate at which failures occur drops off to a lower rate which represents the normal operating situation. In this phase,

failures occur mainly as a result of residual design defects, degradation requiring part replacement or adjustment, secondary failures, and of the true random failures. This lower rate of accumulation of failures continues for a relatively long time until finally those parts subject to gradual deterioration or wearout start reaching their normal life spans and begin to pop off with regularity, raising the rate of the accumulation curve once more. The system is said to have reached its wearout phase.

If, instead of plotting the accumulation of failures the mean time between failures (MTBF) had been plotted, the life characteristic would appear as shown in Fig. 1(b); MTBF being relatively low in early life, regularly higher in useful life and again lower in wearout.

A third parametric representation of this three-phase life model and the most common is shown in Fig. 1(c) and plots the failure rate in terms of failures per time unit against life.

It should be noted that this three-phase life model - debug, useful life, wearout - is a generalization which need not hold true depending on the point of view of the observer. For instance, as a user of the product, one may never be aware of the high initial failure rate if the manufacturer has thoroughly debugged the product before releasing it. Similarly, as a user, one might never be conscious of product wearout if he follows a policy of replacing the equipment or obsolescing it before expiration of the useful life, or if one uses a thoroughly planned and conscientiously executed program of preventive maintenance.

It is, in fact, highly desirable that the manufacturer should debug the system thoroughly prior to its being used in the field and that wearout be avoided either by replacement at obsolescence or by preventive maintenance. Only in this way will the system be operating in this most reliable central region throughout its useful life. This central region of low, essentially constant failure rate is the most important from the standpoint of operational reliability. Most procedures for reliability prediction, therefore, strive to predict the reliability of systems based on this assumption of a constant failure rate during the useful life.

#### System Survival Probability

At this point, reliability will be restated as the numerical probability that a system will perform within specifications and under the conditions of intended use, for a given period of time.

It is here stated (without proof)<sup>1</sup> that for a system operating in a time region of constant failure rate, this probability is given by an exponential function relating the failure rate of the system to the time period for which the reliability is to be estimated. This "exponential failure law" gives the reliability or survival probability,  $P_s$ , as

$$P_s = e^{-\lambda t}$$

where  $\lambda$  is the failure rate in failures per time unit and  $t$

<sup>1</sup> For proof, see Reference (1) at the end of the paper.

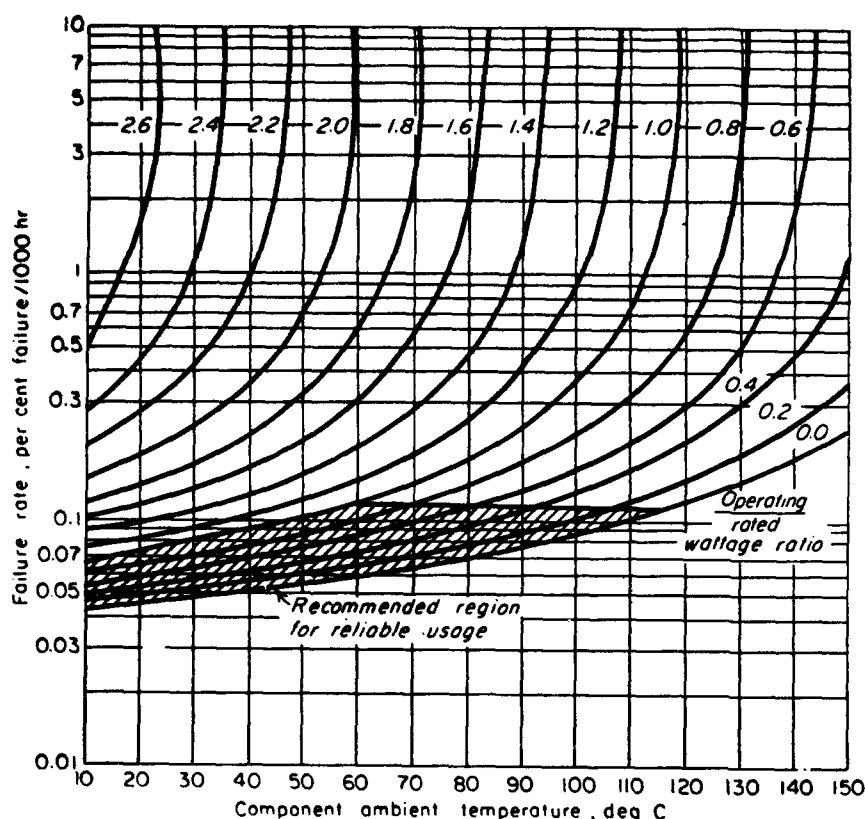


Fig. 2 Part failure rate as a function of application stress. (Source: "A Reliability Approach to Thermal Design and Evaluation," T. C. Reeves, Electrical Manufacturing, February 1957, p. 85.

Predicted failure rates for composition resistors (MIL-R-11A, characteristics GF). Failure rate figures given are the best engineering approximation of reliability characteristics (random failures) for the parts designated

is the period of operation for which the survival probability is sought. Since failure rate is the reciprocal of MTBF, the reliability can also be stated as

$$P_s = e^{-t/m}$$

where  $m$  is the MTBF expressed in the same time units as  $t$ .

Probing the limits of this expression shows that the reliability of a high MTBF unit for short times of operation ( $t/m \rightarrow 0$ ) approaches  $e^0$  or unity. The reliability of a low MTBF unit required for a long period ( $t/m \rightarrow \infty$ ) approaches  $e^{-\infty}$  or zero. It is also notable that where a device is called on to operate for a period equal to its MTBF, its reliability will only be  $e^{-1}$  or about 0.37.

Given the exponential failure law, one can predict the reliability of a system provided he can predict the failure rate of the same system.

#### System Failure Rate

In a system so configured that any part failure will result in a system failure, i. e., a chain or series system, it is apparent that the number of system failures over a period of time will be equal to the sum of the individual part failures causing system failure. This is intuitive, however, and it will be offered (again without proof) that the failure rate of such a system is equal to the sums of the failure rates of the individual parts making up that system (1).

In other words, if the failure rates of all of the parts going into a system can be estimated, the reliability of the

system can be predicted by summation. This is a simple but extremely powerful relationship. Because of it, a great deal of research, test and data-processing effort has been invested over the past 10 years to determine, to a usable degree of accuracy, the failure rates of the typical building blocks used in complex systems. These building blocks are mechanical, electromechanical, and electronic. In terms of population density, the latter two categories predominate in most systems and have received the lion's share of attention and data accumulation.

#### Part Failure Rates

A great deal of useful data on part failure rates can be derived from examination of field maintenance records. Originally, reliability prediction was based on the use of part failure rates derived from prior experience in earlier equipments. However, part application conditions and stresses vary so much from one design to another that the failure rates, even for the same part, will vary from one design to another. There is a limit then as to how much confidence can be placed in the use of past equipment history as a basis for future prediction on new equipments.

Hence, many test programs were conducted on the parts themselves to determine what failure rates obtained for various combinations of electrical and mechanical stresses. These tests, conducted by the parts manufacturers as well as by the parts users, while seldom exhaustive, did establish certain end-points which with interpolation and

much "engineering judgment" provided a basis for assessing the reliability of a part in terms of the design stresses expected in the new design.

Thus, based on data drawn from part tests under controlled conditions and from field histories on complete equipments, so-called failure rate curves for parts have been put together. Typical of these failure-rate curves is Fig. 2 relating to a carbon-composition resistor, a common and high population part in electronic systems. The ordinate is failure rate in per cent failing per 1000 hr. The abscissa is ambient temperature in deg C; i. e., the temperature measured or anticipated in the immediate vicinity of the part as used in the equipment. There is a family of curves, each one representing a different electrical stress in terms of the ratio of usage stress to the part's rated stress. In this case the stress factor is the ratio of actual wattage, as used, to nominal or rated wattage. As the electrical wattage stress and/or the thermal stress are increased, higher failure rates are indicated for that intended application.

Counterpart curves have been established for most other high-population-systems parts such as electron tubes, transistors, capacitors, coils, transformers, motors, relays, connectors, switches and so on (1). With such data it is possible for the designer of an equipment or system to predict the reliability of the design as soon as estimates can be made of the parts to be employed and of the design conditions and stresses under which they are to be employed.

In review, to make such a reliability prediction, the designer need only

- 1 Determine the vital parts making up the system.
- 2 Estimate the stresses imposed on these parts by intended use.
- 3 Determine the applicable failure rates at these stress levels.
- 4 Determine, by summation, the resulting system failure rate.

This is a relatively simple task which any design engineer can learn to do quickly, given the opportunity to familiarize himself with the equipment, given access to failure rate data, and being familiar with the actual computational routine. To be sure, skill is required to determine stress levels and data and judgment are required to assign failure rates.

## 2 - THE VALIDITY OF RELIABILITY PREDICTION

The next question is, does it work? How good are the results? In short, how reliable are reliability predictions?

The test of validity is one of corroboration. Once reliability has been predicted, how closely does the prediction agree with observations made much later after the equipment is in use? The period of uncertainty is a long one since several years usually pass between a final reliability analysis on the prototype design and the collection of sufficient, valid field history on production equipments. It is not surprising then that we do not have a large number

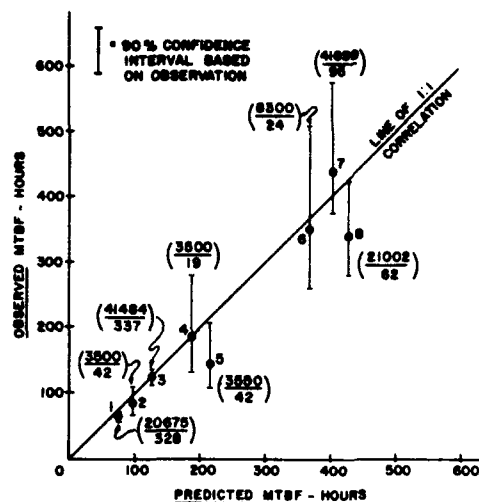


Fig. 3 Correlation of reliability predictions with subsequent observations

of these cases of prediction versus experience on which to base the case for validity. For instance, in RCA, although reliability predictions have been made since 1955, field histories have been accumulated on only about eight systems. On this small sample, the results are encouraging. Needless to say, this has been a great relief to those who have sweated out the past 5 years waiting to see whether a golden egg was laid or just a lead balloon.

These results, which represent case histories of eight military and commercial equipments for ground and air, are shown in the correlation plot, Fig. 3. Observed MTBF, based on a total of over 144,000 hr of operation, are plotted as ordinates against the predicted values as abscissas. In case of perfect correlation, the points would fall on the 45 deg line of 1:1 correlations; displacements from the line thus represent varying degrees of lack of correlation.

Note that the observed mean times to failure shown are based on a limited number of hours of observation. As with any observed average based on a sample period of time, the true average for the "universe" is not known but can only be inferred statistically in terms of confidence intervals. Hence the 90 per cent confidence intervals associated with each of the observation groups are superimposed as vertical ranges on the observed sample MTBF.

These intervals can be interpreted as follows: For equipment No. 7, for instance, the observed sample MTBF is 439 hr. This is not necessarily the true value for all type 7 equipments; it is only the value based on 41,699 hr of observation of some of the type 7 equipments. But based on the distribution observed in this sample, we are 90 per cent sure that the true value for the type 7 universe will lie between 374 and 525 hr. In other words, there is a 5 per cent chance that the true but unknown value is greater than 525 and also a 5 per cent chance that it is less than 374 but 90 per cent of the time we will be correct if we infer the true value to be between 374 and 525 hr.

In relation to the plot, these intervals mean that the true observed value may well fall on the 1:1 line in all but two cases. In the remaining six cases, the fact that the plotted points do not fall closer to the line might well represent sampling errors rather than a defective prediction. Admittedly, in cases 1 and 5, the results are not so good.

Overall, for the eight cases, the greatest error between a prediction and a sampled observation is 50 per cent (case 5) but the average for all cases is only 13 per cent.

On the basis of this degree of agreement between early reliability predictions and later observations, it is concluded that reliability prediction is a valid, dependable design tool which yields results of usable engineering accuracy - not precise and not highly accurate to be sure but useful for a variety of tasks.

### 3 - USES OF RELIABILITY PREDICTION

It should be apparent that armed with a technique for predicting the frequency of equipment failures, one is able to cope not only with the primary problem of designing a reliable equipment but also with a wide range of associated problems touching on maintenance and repair, on availability and reserve requirements and on total cost of operation of the same equipment. These latter problems also require assessment and solution during the equipment-design phase so that decisions can be made as to how many equipments are required to satisfy a given requirement, what support forces are required, what maintenance burden will be encountered and so on.

Following is a brief description of how reliability prediction can be used to furnish early estimates of:

- 1 Maintenance force requirements.
- 2 Availability or up-time and standby requirements.
- 3 Operational force requirements.

#### Maintenance Force Requirements

In estimating maintenance force requirements, one seeks to estimate the number of personnel and skills required to maintain the equipment and also the types and quantities of test equipment and facilities needed. If one estimates the average time and skills required to repair an equipment (based on actual experience with similar equipments or either empirically or synthetically by time study) then, knowing the MTBF, a ratio of average hours of repair per hour of operation can be derived. Knowing how many equipments will be in use in a given location and knowing their planned operating schedule, say in terms of hours per month, one can then proceed to estimate the man-hours of repair time needed per month. From this and assuming a given work-week schedule, one can estimate how many maintenance men should be provided to handle the average maintenance load and also how large a reserve force is required to handle peak loads with given levels of confidence. Test-set requirements and facilities can also be approached in this same way.

Reliability prediction, in giving failure rates for parts and subassemblies, also furnishes useful indications of spare-part requirements and inventory levels. It can be seen, then, that reliability prediction is an important tool for the intelligent planning and control of maintenance cost.

#### Availability and Standby Requirements

When an equipment incurs a failure, it is usually out of service until repairs or adjustments can be made and checked out. This out-of-service time or down-time is a function of the time to get to the site of the failure, isolate it, decide what has to be done, do it, test to see if it now works, clean up and return the equipment to service. This down-time, like repair time itself, can be estimated and averages drawn. Then with such an average down-time and a predicted MTBF, the availability,  $A$ , or up-time is given by

$$A = \frac{MTBF}{MTBF + \text{average down-time}}$$

This expression states the ratio of the average number of hours the equipment is up or available to the total time in commission. Availability can thus be looked on as a quasi-probability that an equipment will be ready when needed.

If availability is low and the cost of being in a down-state is high, standby equipments must be provided ready to take over instantly in event of a failure of the on-line unit. The failed unit is then repaired and becomes the standby for the on-line unit and so on. When such standby is furnished, the availability increases markedly since the probability that no equipment will be available when needed is now really the probability that the standby unit, when switched on to take over, will fail before the original unit can be restored to readiness. This probability  $P_s$  is given by

$$P_s = e^{-r/MTBF} = e^{-r/\lambda}$$

where  $r$  is average down-time.

However, if even this probability is too low, a third standby can be provided to go on in the event the first standby does fail before the original is restored, with even higher resulting availabilities. Very high availabilities are obviously required in early warning and retaliatory defense systems so that determination of adequate standby capacity is a major factor in early systems planning.

#### Operational Force Requirements

An availability estimate permits estimating the number of units actually needed to provide a given degree of mission reliability, provided the units are repairable. However, when the units are not repairable once the mission has begun, as is the case with missiles and most airborne equipment, failed units cannot be restored. In this nonreplacement case, then, the original force undergoes attrition - like the Ten Little Indians - according to the operational failure rate predicted along lines outlined in the first section of this paper. Thus, in order to complete

a mission or time period with a given required number of surviving units, the original starting force must be larger in the same sense that, in a gear train, input power must exceed output to provide for losses. In this case, the efficiency is the mission survival probability and the initial force requirement is given by the final requirement divided by the mission survival probability or

$$N_{\text{initial}} = N_{\text{final}} / e^{-t/\text{MTBF}}$$

where  $t$  is the mission duration.

This estimate of force requirements is critical to system feasibility since it indicates how much hardware is required to carry out a mission, and this in turn determines how much support is needed to keep the hardware in readiness to undertake the mission - both as major elements in the total, lifetime cost of, say, a weapon system. Those who may not consider reliability in the cold, tangible light of dollars and cents will note that it is predicted reliability which determines what portion of the tax budget is earmarked to buy, install and support an adequate force.

#### Optimizing Reliability At Minimum Total Cost

It should be clear that since reliability prediction can furnish useful inputs to the foregoing operational requirement estimates, the same prediction can also help to point out in which equipments and subsystems improvement efforts will be most fruitful. For example, in the typical system, success is defined in terms of a sequence of functions being properly performed, each of the subsystems taking its inputs from the preceding subsystems, performing a function and transferring its own outputs to the following subsystem. With such a chain, the system reliability is given by the product of all of the subsystem reliabilities. Since this over-all reliability can never be higher than that of the least reliable subsystem, reliability prediction serves to identify those subsystems and equipments which limit system achievement.

In addition to identifying the weakest links in a system, reliability prediction also enables estimating the degree of improvement offered by corrective schemes. For example, once the critical subsystem is located, its reliability may usually be improved by several alternatives such as, (a) by redesigning the system as a whole to make this subsystem less critical or even nonessential; (b) by providing standby units for the critical subsystem; (c) by redesigning the subsystem itself to simplify it; (d) by substituting more reliable parts; (e) by further derating the parts; (f) by providing more effective cooling and so on. Each of these alternatives or combinations would probably promise different levels of improvement and also involve different levels and distributions of design, hardware, and support costs. Based on total cost trade studies which weigh the benefits of each approach, one can determine where to put his money to best advantage.

Thus reliability prediction is a necessary part of value engineering. It helps tell us, first, how to spend a

fixed sum of dollars so as to realize minimum total (lifetime) costs for that sum. It can help to tell us how to minimize total cost for a given level of reliability. It can also help indicate where additional investments will yield the highest payoff in terms of additional reductions in total cost. The latter indication applies to product improvements; the former to original development and design.

#### Design Assurance and Discipline

It is in this area of design improvement that reliability prediction finds perhaps its most important role. The prediction itself, while resulting in a very useful number, is nevertheless not as important as the thorough and objective analysis of the design which must be made to derive that number. In the circuit-by-circuit, part-by-part study of performance requirements under assumed stresses and part capabilities, one is performing, a screening which identifies not only outright design oversights but also marginal applications. In both instances, the reliability analysis which precedes the prediction yields recommendations for timely corrective action before design release. Because the reliability analysis thus provides a high degree of design assurance, it is frequently made an integral part of the design review and engineering approval procedure. In order that this design review and approval be as unprejudiced as possible, it is usually made by competent senior engineers and specialists who have not played an active role in the actual design.

This independent review provides an important built-in-element of design discipline. On the other hand, it is sometimes said that if one adopts a procedure for independent design review and analysis prior to engineering release, this will provide indirect encouragement to the design group to minimize its own reliability efforts. The rationale is that since the design group is confident that a review group will pick up and resolve any reliability problems, the design group is free to concentrate on other areas. Human nature being what it is, assurance can be given that it just does not work out this way at all. Original design becomes more self-critical (not less) when a program of independent design review and reliability analysis is undertaken. The reader would be surprised to find how reliability-conscious and competent a design group can become when it sets out to prove that a review of its design by independent experts is just a waste of time! One might say, that the objective of reliability engineering is to bring the reliability problem under such a degree of control that reliability engineers will no longer be needed and will be forced to turn to something more useful, like design engineering. Present indications are that we are not going to meet this objective.

#### Summary and Conclusions

In a paper where simplification has been deliberate, one should not be misled into believing that reliability analysis and prediction is all there is to reliability improvement. Reliability prediction is certainly not a substitute

for the many essential elements of a balanced reliability program but it is a sound first step until other than paper-work measures are possible.

Reliability engineering is a field specialized to a high degree. Sometimes, unfortunately it may seem to be more mysticism and black art than it is down-to-earth engineering. In particular, many engineers look on reliability prediction as a kind of space-age astrology in which failure rate tables have been substituted for the Zodiac.

I hope I have dispelled some of that hocus-pocus and that I have exposed reliability prediction to you as a useful engineering tool not just in the field of complex electronics but in any field where one must seek a high measure of confidence in design.

#### References

1 "Reliability Stress Analysis for Electronic Equipment," RCA Technical Report 1100, 179 pp. Available from the Government Printing Office under above title as NAVSHIPS 900,193 (PB-131678).

2 "Reliability of Military Electronic Equipment," Report by Advisory Group on Reliability of Electronic Equipment, Office of Assistant Secretary of Defense R & D, 378 pp. Available from the Government Printing Office at \$1.75.

3 "NEL Reliability Design Handbook," U. S. Navy Electronics Laboratory. Available from Office of Technical Services, U. S. Department of Commerce, Washington 25, D. C., as PB 121839 at \$3.00.