

1

Basis and Limitations of Typical Current Reliability Methods and Metrics

Reliability cannot be achieved by adhering to detailed specifications. Reliability cannot be achieved by formula or by analysis. Some of these may help to some extent, but there is only one road to reliability. Build it, test it and fix the things that go wrong. Repeat the process until the desired reliability is achieved. *It is a feedback process and there is no other way.*

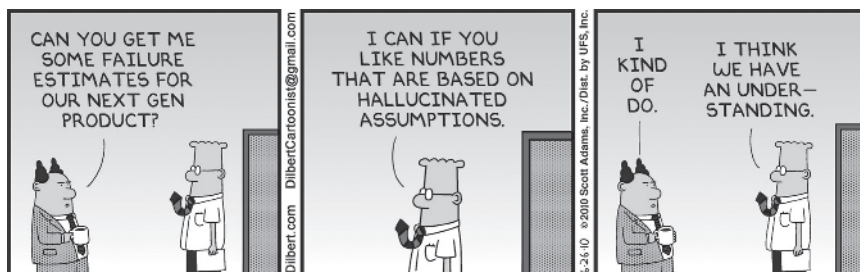
David Packard, 1972

In the field of electronics reliability, it is still very much a Dilbert world as we see in the comic from Scott Adams, Figure 1.1. Reliability Engineers are still making reliability predictions based on dubious assumptions about the future and management not really caring if they are valid. Management just needs a 'number' for reliability, regardless of the fact it may have no basis in reality.

Next Generation HALT and HASS: Robust Design of Electronics and Systems, First Edition.

Kirk A. Gray and John J. Paschkewitz.

© 2016 John Wiley & Sons, Ltd. Published 2016 by John Wiley & Sons, Ltd.



DILBERT © 2010 Scott Adams. Used By permission of UNIVERSAL UCLICK. All rights reserved.

Figure 1.1 Dilbert, management and reliability. Source: DILBERT © 2010 Scott Adams. Reproduced with permission of UNIVERSAL UCLICK

The classical definition of reliability is the probability that a component, subassembly, instrument, or system will perform its specified function for a specified period of time under specified environmental and use conditions. In the history of electronics reliability engineering, a central activity and deliverable from reliability engineers has been to make reliability predictions that provide a quantification of the lifetime of an electronics system.

Even though the assumptions of causes of unreliability used to make reliability predictions have not been shown to be based on data from common causes of field failures, and there has been no data showing a correlation to field failure rates, it still continues for many electronics systems companies due to the sheer momentum of decades of belief. Many traditional reliability engineers argue that even though they do not provide an accurate prediction of life, they can be used for comparisons of alternative designs. Unfortunately, prediction models that are not based on valid causes of field failures, or valid models, cannot provide valid comparisons of reliability predictions.

Of course there is a value if predictions, valid or invalid, are required to retain one's employment as a reliability engineer, but the benefit for continued employment pales in comparison to the potential misleading assumptions that may result in forcing invalid design changes that may result in higher field failures and warranty costs.

For most electronics systems the specific environments and use conditions are widely distributed. It is very difficult if not impossible

to know specific values and distributions of the environmental conditions and use conditions that future electronics systems will be subjected to. Compounding the challenge of not knowing the distribution of stresses in the end - use environments is that the numbers of potential physical interactions and the strength or weaknesses of potential failure mechanisms in systems of hundreds or thousands of components is phenomenologically complex.

Tracing back to the first electronics prediction guide, we find the RCA release of TR-1100 titled *Reliability Stress Analysis for Electronic Equipment*, in 1956, which presented models for computing rates of component failures. It was the first of the electronics prediction 'cook-books' that became formalized with the publishing of reliability handbook MIL-HDBK-217A and continued to 1991, with the last version MIL-HDBK-217F released in December of that year. It was formally removed as a government reference document in 1995.

1.1 The Life Cycle Bathtub Curve

A classic diagram used to show the life cycle of electronics devices is the life cycle bathtub curve. The bathtub curve is a graph of time versus the number of units failing.

Just as medical science has done much to extend our lives in the past century, electronic components and assemblies have also had a significant increase in expected life since the beginning of electronics when vacuum tube technologies were used. Vacuum tubes had inherent wear-out failure modes, such as filaments burning out and vacuum seal leakage, that were a significant limiting factor in the life of an electronics system.

The life cycle bathtub curve, which is modeled after human life cycle death rates and is shown in Figure 1.2., is actually a combination of two curves. The first curve is the initial declining failure rate, traditionally referred to as the period of 'infant mortality', and the second curve is the increasing failure rates from wear-out failures. The intersection of the two curves is a more or less flat area of the curve, which may appear to be a constant failure rate region. It is actually very rare that electronics components fail at a constant rate, and so the 'flat' portion

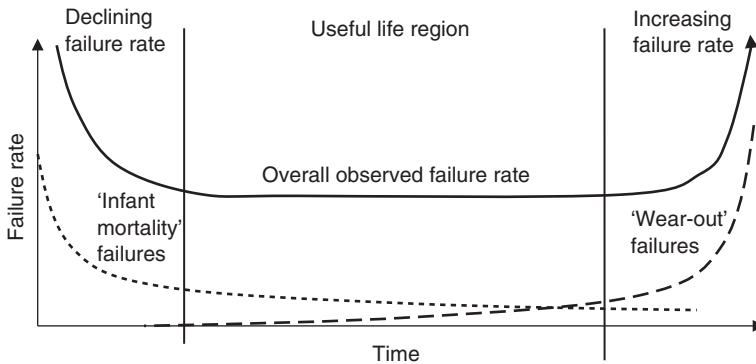


Figure 1.2 The life cycle bathtub curve

of the curve is not really flat but instead a low rate of failure with some peaks and valleys due to variations in use and manufacturing quality.

The electronics life cycle bathtub curve was derived from human the life cycle curves and may have been more relevant back in the day of vacuum tube electronics systems. In human life cycles we have a high rate of death due to the risks of birth and the fragility of life during human infancy. As we age, the rates of death decline to a steady state level until we age and our bodies start to fail. Human infant mortality is defined as the number of deaths in the first year of life. Infant mortality in electronics has been the term used for the failures that occur after shipping or in the first months or first year of use.

The term 'infant mortality' applied to the life of electronics is a misnomer. The vast majority of human infant mortality occurs in poorer third world countries, and the main cause is dehydration from diarrhea, which is a preventable disease. There are many other factors that contribute to the rate of infant deaths, such as limit access to health services, education of the mother and access to clean drinking water. The lack of healthcare facilities or skilled health workers is also a contributing factor.

An electronic component or system is not weaker when fabricated; instead, if manufactured correctly, components have the highest inherent life and strength when manufactured, then they decline in strength, or total fatigue life during use.

The term 'infant mortality', which is used to describe failures of electronics or systems that occurs in the early part of the use life cycle,

seems to imply that the failure of some devices and systems is intrinsic to the manufacturing process and should be expected. Many traditional reliability engineers dismiss these early life failures, or 'infant mortality' failures as due to 'quality control' and therefore do not see them as the responsibility of the reliability engineering department. Manufacturing quality variations are likely to be the largest cause of early life failures, especially for designs with narrow environmental stress capabilities that could be found in HALT. But it makes little difference to the customer or end-user, they lose use of the product, and the company whose name is on it is ultimately to blame.

So why use the dismissive term infant mortality to describe failures from latent defects in electronics as if they were intrinsic to manufacturing? The time period that is used to define the region of infant mortality in electronics is arbitrary. It could be the first 30 days or the first 18 months or longer. Since the vast majority of latent (hidden) defects are from unintentional process excursions or misapplications, and since they are not controlled, they are likely to have a wide distribution of times to failure. Many times the same failure mechanism in which the weakest distributions may occur within 30 to 90 days will continue for the stronger latent defects to contribute to the failure rate throughout the entire period of use before technological obsolescence.

1.1.1 Real Electronics Life Cycle Curves

Of course the life cycle bathtub curves are represented as idealistic and simplistic smooth curves. In reality, monitoring the field reliability would result in a dynamically changing curve with many variations in the failure rates for each type of electronics system over time as shown in Figure 1.3. As failing units are removed from the population, the remaining field population failure rate decreases and may appear to reach a low steady state or appear as a constant or steady state failure rate in a large population.

In the real tracking of failure rates, the peaks and valleys of the curve extend to the wear-out portion of the life cycle curve. For most electronics, the wear-out portion of the curve extends well beyond technological obsolescence and will be never actually significantly contribute to unreliability of the product.

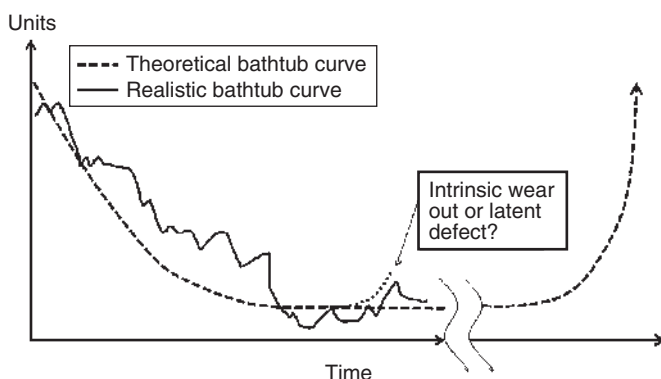


Figure 1.3 Realistic field life cycle bathtub curve

Without detailed root cause analysis of failures that make up the peaks of the middle portion of the bathtub curve, or what is termed the useful life period, any increase in failure rates can be mistaken as the intrinsic wear-out phase of a system's life cycle. It may be discovered in failure analysis that what at first appears to be an wear out mode in a component, is actually due to it being overstressed from a misapplication in circuit or unknown high voltage transients.

The traditional approach to electronics reliability engineering has been to focus on probabilistic wear-out mode of electronics. Failures that are due to the wear-out mode are represented by the exponentially increasing failure rate or back end of the bathtub curve.

Mathematical models of intrinsic wear-out mechanisms in components and assemblies must assume that all the manufacturing processes – from IC die fabrication to packaging, mounting on a printed wiring board assembly (PWBA) and then final assembly in a system – are in control and are consistent through the production life cycle.

Mathematical models must also include specific values of environmental stress cycles that drive the inherent device degradation mechanisms for each device, which may include voltage and temperature cycles and shock and vibration, which can interact to modify rates of degradation. The sum of all the stresses that a whole product is expected to be subjected to during its use is the life cycle environmental profile (LCEP).

The cost of failures for a company introducing a new electronics product to market are much more significant at the front end of the bathtub curve, the 'infant mortality' period, rather than the 'useful life' or 'wear-out' period in the bathtub curve. This includes the tangible and quantifiable cost of service and warranty replacements, and less tangible but real costs in lost sales due to perceptions of poor reliability in a competitive market.

There is little data or supporting evidence that in general electronics systems intrinsic life can be modeled and predicted, and this is especially true for the early life failures. The misleading approach of using traditional reliability predictions for reliability development will be discussed further in Chapter 2.

1.2 HALT and HASS Approach

The frame of reference for the HALT and HASS approach, reliability testing is as simple as the old adage that 'a chain is only as strong as its weakest link'. A complex electronics system is only as strong as its weakest or least tolerant or capable component or subsystem. Just like pulling on a chain until the weakest link breaks, HALT methods apply a wide range of relevant stresses, both individually and in combinations, at increasing levels in order to expose the least capable element in the system. If the failure mechanism causes catastrophic damage to a component, when a destruct limit is reached in HALT, makes it easier to isolate a weak link, identifying the weak link is easier to isolate. Operational weakness causing soft failures can be more challenging to isolate.

HALT (highly accelerated life test) is a process that requires specific adaptation when it is applied to almost any system and assembly. Because HALT is a highly adaptive process, the information given in this book will be general guidelines on how to apply HALT. How HALT is adapted to each type of product or assembly is unique to each, and presents a learning process for each different type of electronic and electromechanical system. It is advised that a company that plans to adopt HALT as a new process or a new user of HALT will have a significantly faster adoption and success in implementation if they have the guidance of an experienced HALT consultant. As in any newly introduced adoption of test new methods and techniques, there are

many engineers and managers that will have misunderstandings of the process and the goals of HALT and HASS (highly accelerated stress screening). An experienced HALT consultant will have the data and knowledge to keep the focus on the adaptive application and relevance of the HALT process and future benefits of creating a robust, but not “over-designed” system. The period between the HALT application for reliability development of a new product and the observation of the actual reliability performance in the field with the lower failure rates as a result of HALT may take many months or longer. An experienced HALT consultant can be the champion of HALT during the additional expense of HALT during product development and before the actual benefits increased reliability due to HALT are realized in the field, as reduced warranty and early life field failures.

The same principles of testing to operational or destruct limits used for HALT of electronics circuit boards can be applied to electromechanical and mechanical systems for purpose of again finding the weakest link in the system applied to electromechanical and some mechanical systems. The main difference is in what stress stimuli are used. HALT for systems other than electronics is discussed further in Chapter 11.

The goal of HALT is to develop the stress margin capability and system strength to the fundamental limits of the current technologies during product development. The fundamental limit of the technology (FLT) is the stress level that cannot be exceeded without using non-standard electronics materials or methods.

HALT is used to find stress limits and design weaknesses that could decrease field reliability, and is best performed during design and development phase. HASS is an ongoing application of combinations of stresses, defined from stress limits found empirically during HALT to detect any latent defects or reduction in the design’s strength introduced during mass manufacturing.

Only after a system weakness is discovered can it be investigated and its significance and relevance to reliability be determined. Occasionally a weakness found in HALT is evaluated and not considered a risk of causing field failures. The opportunity to evaluate a weakness only comes when you find the stress limits. If the product is not tested to stress limits or failure, there is nothing to evaluate for potential reliability improvement.

HALT is becoming more widely adopted by electronics companies in the 21st century, although it is also more a current industry buzzword that may be used for marketing promotion than a process for actual improvement of electronics systems by increasing stress-strength margins. Suppliers of some subsystems in the IT hardware industry, such as power supplies, memory, or graphics display devices may use HALT, but the specifics of what is called a HALT can vary widely. It has been the author's experience that many purportedly using HALT may do stress tests, but only stress to a predetermined stress level that someone has arbitrarily determined is 'good enough'. One valuable result of HALT is the comparison of stress limits found between samples of the same product in HALT. Without finding empirical limits they will not be able to compare limits between samples of the same product. Wide distributions of strength seen as large differences in empirical operation or destruct limits can be an indication of inconsistent manufacturing at some level of the product.

One of the author's consulting clients had been performing HALT for many years on their products, yet when asked what the thermal operational limit was for one product of concern they admitted that they did not know because the HALT was stopped at 80°C because that was 'good enough'. Without finding a thermal operational limit, they missed discovering an important and revealing comparison of the operational limits between samples.

1.3 The Future of Electronics: Higher Density and Speed and Lower Power

Moore's Law, the projection that Gordon Moore made in 1965 that the number of components on an integrated circuit would approximately double every two years, has become an industry expectation for new component designs. The increase in densities of integration, reduction of feature sizes in integrated circuits and new packaging technologies introduces new fabrication and use physics that drive failure mechanisms and this is expected to continue for the foreseeable future.

Other changes in electronics materials may be implemented from concerns of the impact of electronics on the earth's environment. The change in going from using leaded solders to lead-free solders, and

restricting the use of flame retardants are two examples of changes required by the directive on the restriction of the use of certain hazardous substances in electrical and electronic equipment. The directive was made by the European Union in 2002 for all electronics sold there and has been adopted worldwide. It is now commonly abbreviated as RoHS (Restriction of Hazardous Substances).

In the design and development of electronics, all of the changes and the rapidly increasing density and complexity of devices and systems make modeling each potential failure mechanism a moving target. Soon after a model of a new technology or failure phenomena is introduced, new materials and new technologies change the underlying physics of the causes of wear-out failures in devices or systems during their use.

The reliability of an electronics system is a phenomenologically complex issue. Prediction models do not include all the potential design and manufacturing errors or process deviations that may affect device and system reliability. Models of electronic component failure mechanisms that are used for reliability predictions are – and must be – based on the assumptions of the manufacturing processes being consistent and capable at all levels of system fabrication and throughout the manufacturing life cycle.

1.3.1 There is a Drain in the Bathtub Curve

The life entitlement of today's electronics components and systems with no moving parts far exceeds the life needed before a system is replaced by a newer more capable system. Technological obsolescence comes faster for today's electronics systems and they are replaced long before their life is consumed. The timescales between intrinsic wear-out modes of active devices and technological obsolescence of a system is significant in the vast majority of electronics. Because of the large difference in the timescales between obsolescence and wear-out of components and assemblies, wear-out mechanisms in electronics systems will never be observed. Also, because of the long life entitlement of electronics, using a small percentage of the fatigue life of electronics during HASS in production in order to find latent defects leaves more than enough life for the system to be shipped as new and to exceed the period of its technological obsolescence.

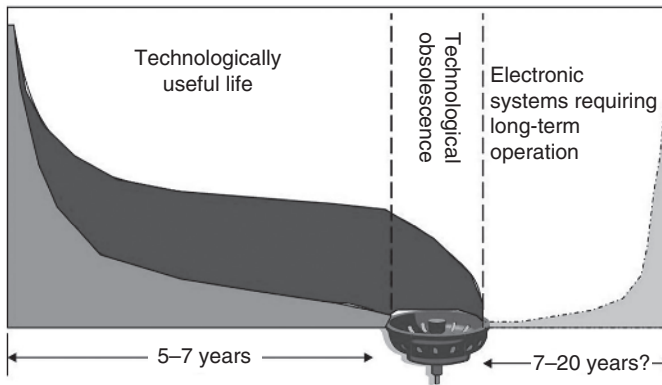


Figure 1.4 The 'drain' of technological obsolescence in the life cycle bathtub curve

There are electronic devices, such as batteries, that still have a relatively limited life compared to most circuit components, and for that reason are typically designed to be easily replaced. With few exceptions, intrinsic wear-out mechanisms on most components have not been shown to contribute to electronics systems unreliability during the first years of operation. Almost all failures of electronics in the first years are due to assignable causes, such as overlooked low design margin, an error in manufacturing, or misapplication or abuse by the customer, among many potential causes. An illustration of the electronics bathtub curve and the 'drain' of technological obsolescence is shown in Figure 1.4.

The vast majority of the costs of failures for almost all electronics manufacturers come in the first few years in the product's life. The left side of the bathtub curve shows a declining failure rate.

It is during the time of warranty coverage that company must pay for system repair or replacement costs for failed units. Failures during and shortly after the warranty period may also be a much greater a contributor to the financial loss for an electronics OEM as a result of lost sales from the market's perception of lower reliability. Loss of market share, and therefore unit sales, may be much greater than the material and service warranty costs, but since it is difficult to quantify lost sales, the actual monetary losses may never be known.

Technological obsolescence occurs at a much faster rate than intrinsic mechanisms in electronics and systems that lead to wear-out for most electronics systems and it is especially true in consumer electronics. It can be argued that most failures of electronics systems are due to errors in the design, manufacturing, or customer misuse or abuse.

Failure Prediction Methodologies (FPM) are more relevant for mechanical systems (i.e. motors, gears, switches) which can have a more limited life due intrinsic to friction and fatigue damage wear out mechanisms. In mechanical systems wear-out, the lifetime use can be modeled from physical measurements of material consumed, change in torque resistance or current draw, or other relevant measurements. The models can then be used to determine whether the wear-out duration of the mechanical device is adequate for the required life or mission requirement. If the intrinsic life is limited by the consumption of material (as in mechanical bearings) the reservoir of material can be increased meet the life requirements.

Technology has changed significantly in electronics in the past decades as IC densities and metallization line widths have continued to shrink, and lower voltage, faster ICs with more functionality are introduced every year. Yet, in the field of electronics reliability engineering, little has changed. The concepts and theories based on MIL HDBK-217 are still widely used, even though MIL HDBK-217 was removed as a government reference document and has not been updated or republished since the last revision ('F', notice 2) in 1991. Much of the data on failure rates of components, such as fans, is outdated by decades and has little relevance to today's electronics. Because of decades of reliance on handbook-based or 'cookbook' reliability predictions and invalid assumptions regarding temperature and component life, there is a continued perception that the higher the temperature at which electronics are operated, the faster the system will use up its 'life entitlement' and the sooner it will fail – regardless of well-documented evidence to the contrary.

1.4 Use of MTBF as a Reliability Metric

Traditional reliability engineering methods have focused on producing a quantitative reliability prediction based on time. The most widely used metrics in reliability are the terms 'mean time between

failures' (MTBF) for repairable systems and 'mean time to failure' (MTTF) for non-repairable systems. MTBF is a single average of the total number of hours a set group of systems have operated between repairs or with MTTR until the first failure. Historically traditional reliability predictions use this single number to describe what can be very different distributions of failure rates. Because it is an average number, without more information it is not very useful for understanding the probability of failures based on use or age of the product. It is a broad statistic that should not be used as a metric for defining reliability design goals or for field analysis of failures and warranty returns.

MTBF is a poor metric for providing information on the reliability of any system. It is derived from a very simple equation:

$$MTBF = \theta = \int_0^{\infty} t \times f(t) dt \quad (1.1)$$

If we have 40 units that all run for 100 hours and right at the end of 100 hours one of the units fails, we can calculate the MTBF as follows. First determine the total hours that all the units operated. It is a very simple calculation, 40 units times 100 hours is 4000 hours. Next divide the total operating hours by the number of failures. One failure makes for a simple example: dividing by one the resulting MTBF is equal to 4000 hours.

The following section 1.5, written by Andrew Rowland who is a Certified Reliability Engineer (CRE), explains how the same MTBF number is calculated for three significantly different distributions and reliability risks.

1.5 MTBF: What is it Good For?

1.5.1 Introduction

The mean time between failure (MTBF) is arguably the most prolific metric in the field of reliability engineering. It is used as a metric throughout a product's life cycle, from requirements, to validation, to operational assessment. Unfortunately, MTBF alone doesn't tell us too much.

It's not that MTBF is a bad metric; it is just an incomplete metric, and as an incomplete metric it doesn't lend itself to risk-informed decision-making. The real problem is not with the MTBF, but with the implicit assumption that failure times are exponentially distributed.

1.5.2 Examples

To illustrate how relying on the MTBF can be misleading, let's look at two examples. In these examples we will assume that the failure times are Weibull distributed. The Weibull distribution is popular in reliability engineering and the exponential is a special case of the Weibull. From the literature, we know that the probability density function and survival (or reliability) function of the Weibull can be expressed as:

$$f(t) = \left(\frac{\beta}{\eta}\right) \left(\frac{t}{\eta}\right)^{\beta-1} e^{-\left(\frac{t}{\eta}\right)^\beta} \quad (1.2)$$

$$S(t) = e^{-\left(\frac{t}{\eta}\right)^\beta} \quad (1.3)$$

We also recall that the mean of a Weibull distributed variable can be estimated as:

$$MTBF = \eta \Gamma\left(1 + \frac{1}{\beta}\right) \quad (1.4)$$

In these functions, η is referred to as the scale parameter and β the shape parameter.

1.5.2.1 Example 1

Consider three items; item A, item B and item C. Perhaps the goal is to select one of these items for our design, and the requirement is to have a 90 hour MTBF or greater. All three items have an MTBF of 100 hours. So, from a reliability perspective, which is the item to choose?

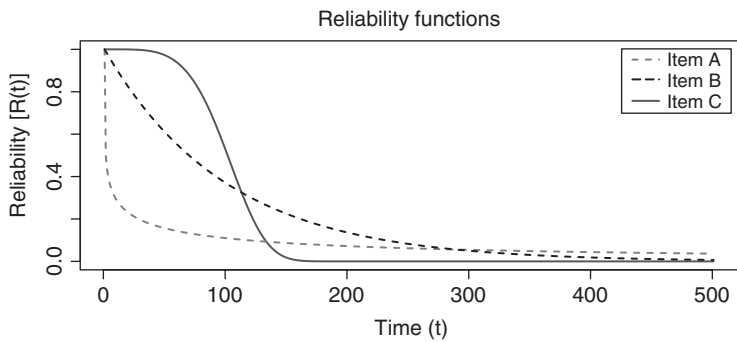


Figure 1.5 Reliability functions for item A, item B and item C

Table 1.1 Reliability at 100 Hours for item A, item B and item C

Item	R(100)
Item A	0.109 (10.9%)
Item B	0.367 (36.7%)
Item C	0.521 (52.1%)

Under the implicit assumption that failure times are exponentially distributed, we might conclude that any of the three is acceptable, reliability-wise. All three satisfy the 90 hours MTBF requirement. However, let’s look a little deeper into the 100 hour MTBF and see if we still agree that any of the three is acceptable.

Let’s take a look at the reliability over time of each item. Figure 1.5 shows the reliability function over 500 hours for each of these items. Clearly, the reliability of these items is not the same. Given that each item has an MTBF of 100 hours, what is the reliability at 100 hours? Table 1.1 summarizes the 100 hour reliability for each item. Once again, we can see a large difference between the three items.

Another way to compare these three items is via the hazard, or failure, rate. Figure 1.6 shows the hazard function for each item. The ‘bathtub’ curve is a plot of hazard rate versus time. Thus, Figure 1.6 shows the ‘bathtub’ curve for each item. Clearly the hazard rate behaviors are very different for these items.

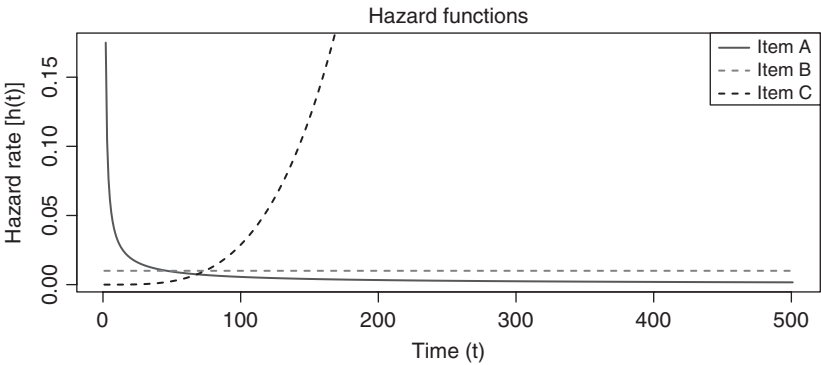


Figure 1.6 Hazard functions for item A, item B and item C

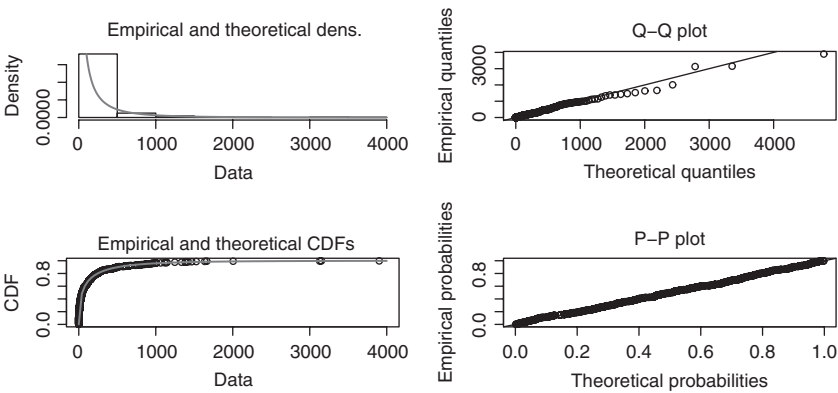


Figure 1.7 Item D: Graphical analysis of survival data

1.5.2.2 Example 2

Consider another situation where we have three items; item D, item E and item F. Presume for a moment that we have all of the data used to derive the MTBF statistic for each item. The first thing we might do is graphically explore the data. Figure 1.7 shows a set of plots commonly used in graphical analysis of survival data for item D. Let's look at the histogram in the upper left corner. We see that the distribution is heavy-tailed, indicating failure times are not exponentially distributed.

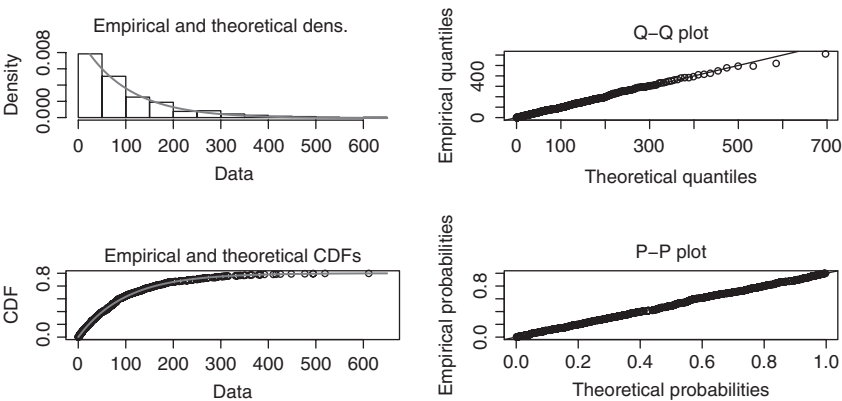


Figure 1.8 Item E: Graphical analysis of survival data

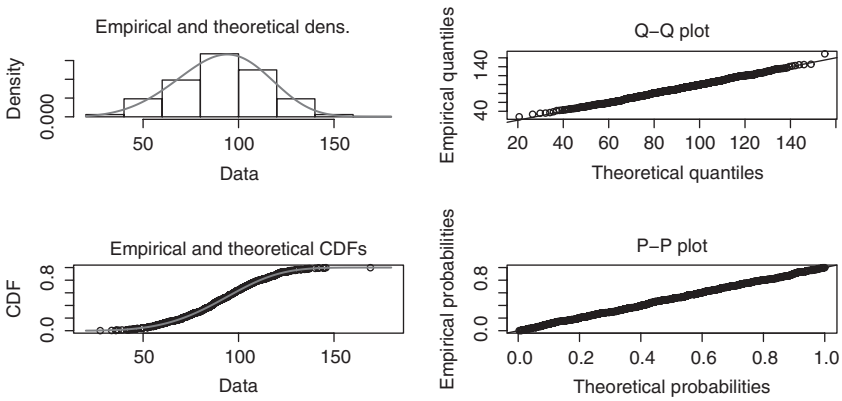


Figure 1.9 Item F: Graphical analysis of survival data

Compare the histogram in Figure 1.7 to that in Figure 1.8 for item E and Figure 1.9 for item F. Clearly the distribution of failure times differs among these three items. Yet all three items have the same MTBF. Perhaps we need to look a bit closer at the data! Now that we’ve graphically analyzed the data and concluded that we may be looking at different populations, we decide to fit the data to a distribution and estimate the parameters.

Table 1.2 Estimated parameters for item D, item E and item F

Item	Eta	Beta	MTBF
Item D	101.42	0.478	220.7
Item E	107.73	1.000	107.7
Item F	100.84	4.524	92.0

Our goal, then, is to estimate the value of β and η for each item. We use the `fitdist` function from the R [1] package, `fitdistrplus` [2] which uses maximum likelihood to estimate the parameters. The results for these three populations are summarized in Table 1.2. We can see from these results that the populations are not the same, although all three items satisfy our 90 hours MTBF requirement.

Now that we’re confident that we’re dealing with three different populations, all with the same MTBF, what is the implication of selecting one item over another? Since we fit the data to a Weibull distribution, we know the shape parameter (β) determines the region of the ‘bathtub’ curve. With a $\beta < 1$, we are in the early life region, a $\beta = 1$ puts us in the useful life region, and a $\beta > 1$ indicates wear-out. In other words, item D is dominated by early-life failure mechanisms, item E is dominated by useful life failure mechanisms, and item F by wear-out.

As we did with the first example, let’s look at the reliability function for these three items.

Figure 1.10 shows the reliability functions. Similar to the first example, we see the reliability functions are not the same as we would expect from our assessment of Figures 1.4, 1.5 and 1.6.

Let’s assume we are interested in the reliability at 50 hours. The reliability at 50 hours for the three items can be found in Table 1.3. We see a dramatic difference in the reliabilities and, interestingly, the item with the highest 50 hour reliability is the item with the lowest MTBF.

We can also look at plots of the hazard function for these three items. These hazard functions are plotted in Figure 1.11 over 500 hours. We see different hazard rate behaviors as we expected from our assessment of the β values we estimated earlier.

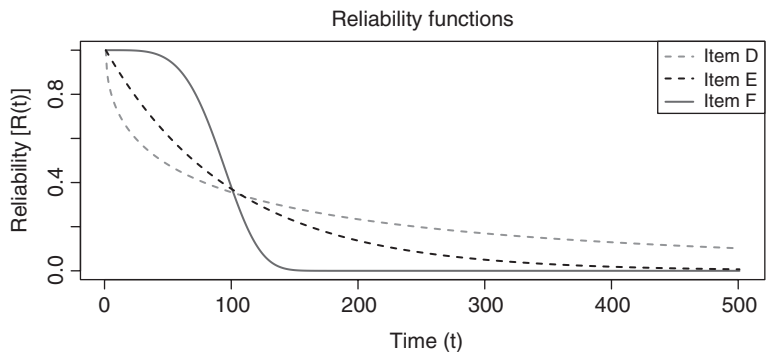


Figure 1.10 Reliability functions for item D, item E, and item F

Table 1.3 Reliability at 50 hours for item D, item E, and item F

Item	R(50)
Item D	0.490 (49.0%)
Item E	0.645 (64.5%)
Item F	0.959 (95.9%)

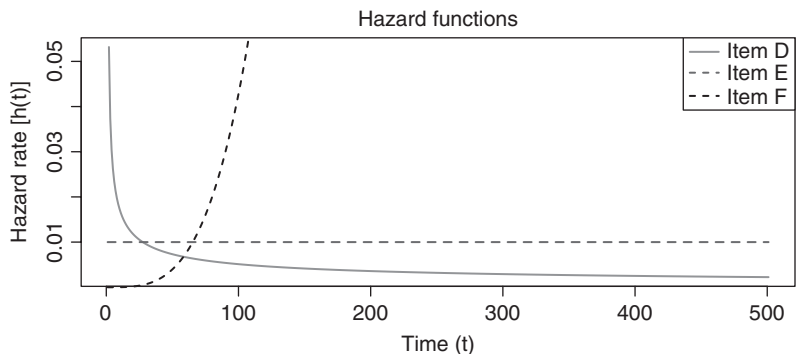


Figure 1.11 Hazard functions for item D, item E, and item F

1.5.3 Conclusion

Hopefully we've come to understand that stating an MTBF value with no other information doesn't really tell us much about the reliability of an item. Neither does it tell us if the item truly satisfies our reliability needs. We saw in one example three items with the same MTBF, but most definitely with different reliability behaviors.

In the second example, we looked at three items with different MTBF. Once again, we saw the reliability behaviors of these items were different. In this example, we saw the item with the largest MTBF having a 50 hour reliability almost half that of the item with the lowest MTBF.

Without an understanding of the reliability characteristics that is more complete than simply MTBF are we making good, risk-informed decisions? Selecting item A or item D, we can expect to see high rates of failure during validation, reliability growth testing or, worse yet, early in customer ownership. If we warrant our product, we can expect large warranty costs associated with items A or D. Given the competing requirements we need to satisfy, we may need to select item A or item D. If we only know the MTBF will we put the necessary barriers in place, such as screening, to minimize the risk?

At the other end of the 'bathtub' curve, if we select item C or item F, our validation or reliability growth testing may not test far enough into wear-out to surface failures. Will we develop a preventive maintenance program for these items to minimize the risk?

MTBF is ingrained in the reliability community as well as throughout most companies. It is unlikely that we will ever see the end of MTBF. Ultimately it comes down to us, as reliability engineers, to understand the limitations of MTBF and educate those around us to its shortcomings. If the reliability community gets in lock-step, we can be the tugboats that change the ship's heading.

~~~~~

The use of MTBF will likely continue along with other misunderstandings of the realities of actual field unreliability since real reliability information that is needed to clarify the rates and causes of field unreliability of most electronics products will never be disclosed. The reason is that publishing the real causes of unreliability of electronics risks

potentially very costly liability and litigation and market share loss for a electronics producer in a competitive marketplace. The change that is needed in electronics reliability will largely come from engineers who have observed and understand the root causes of field failures, not theoretical component failure rates or assumptions of wear-out mechanisms, to change the fundamental approach from developing reliable systems using theoretical assumptions to an approach using deterministic empirical discovery of weaknesses in an electronics and electromechanical system.

#### *1.5.4 Alternatives to MTBF for Specifying Reliability*

Fred Schenkelberg, an experienced reliability engineering and management consultant, is so passionate and determined to help remove the term MTBF from reliability engineering that he has created a website, 'No MTBF' that is dedicated to using better metrics than MTBF to define reliability requirements. Fred has written the following regarding the use of MTBF as a reliability metric.

'MTBF is often used to represent product life. It is neither complete nor sufficient. Product life or reliability has four elements: function, probability, duration and environment. MTBF is only the probability and assumes (in most cases) the duration does not matter, or worse is not even stated.

As an alternative, use reliability directly. State the probability of success over a specified time frame, along with the functions (leads to understanding of product failure definition) and environment. The function and environment are often abbreviated, i.e. a respirator provides life support breathing in North American intensive care facilities. The details of the functions and environment are often well stated in product development and marketing documents.

The probability and duration may include multiple statements. One statement might be for the critical period of the product life. For example, since products that experience failure during first use damage the product brand significantly, we may want to have a very high probability of success during the first 3 months of product use. Say, 99.99% reliability over first 3 months of use.

The warranty period may be duration of interest. In that case the statement for that period would be 98% reliability over the 1 year warranty period. And, the design life (how long the product should last and provide value to the customer) might be stated as 90% reliability over 5 years.

The early failures focus on component, assembly, shipping and installation sources of product failure. The warrant period and reliability is of interest as a business liability. The design life focuses on the longer term failure mechanisms.

Therefore, move away from a partial statement concerning product reliability. Make full use of clear statements of expectations (goals) and measures.' [3]

## **1.6 Reliability of Systems is Complex**

The overall reliability of electronic assemblies and systems is a phenomenologically complex interaction of materials, manufacturing processes and end user applications and the broad potential variations in each of these factors.

If all the functions of design and assembly are performed correctly and if the system is used as intended, it will likely operate without failure until it is technologically obsolete. The pace of electronics technology is increasing and there is no reason to believe that it will slow down. The time for developing reliability in new electronics systems has become and will continue to be shorter. A faster method of ensuring the reliability of electronics systems is needed and will be required for meeting the market expectations and demand.

Gregg Hobbs, with his development of HALT and HASS, derived a much more efficient approach to reliability development using empirical limits under step stress testing to discover elements of a new design that could become a field reliability risk.

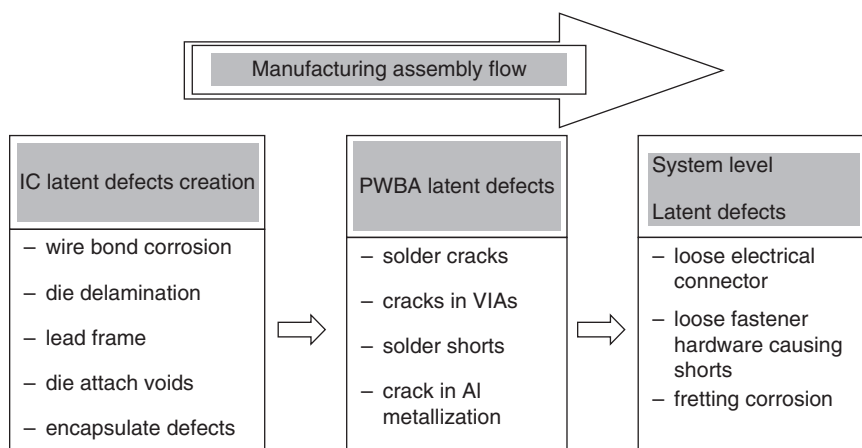
The most valuable time for the creation of a reliable new electronics system is during the design phase when the costs of changes are the lowest. A robust and reliable design provides a higher tolerance to extremes of environmental stress and potential abuse of the product,

as well as creating margins that allow a higher tolerance to variations in the manufacturing processes.

At any point in the manufacturing process a latent defect can be introduced unknowingly and take a product that had been reliable to one that has poor reliability. There are some exceptions, but for most electronics components and systems the life entitlement – that is the length of time it functions before inherent wear-out mechanisms driven by fatigue or chemical reactions result in failure – is much longer than the time at which it is retired because it is technologically obsolete. Most electronics systems have a significant margin between the life entitlement of a properly designed and properly manufactured electronic system relative to that the product is technologically obsolete.

At each manufacturing level of an electronic system there can be variations in the quality and consistency of materials and processes used in the production of systems. Some common latent defects that cause electronics systems to fail can be introduced at each subsequent level of assembly, as shown in Figure 1.12.

For the vast majority of electronics systems, it can be very difficult, if not impossible, to know the life cycle environmental profile



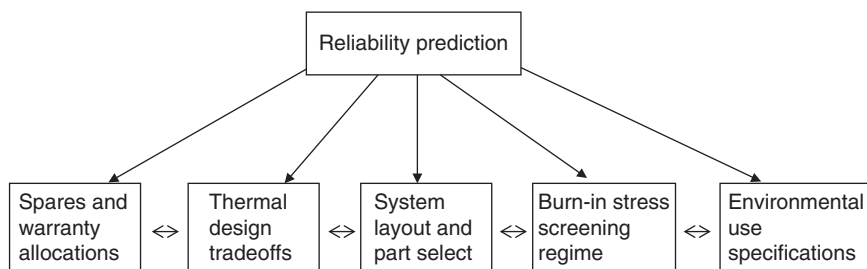
**Figure 1.12** Examples of where latent defects are introduced during assembly fabrication

(LCEP) that any particular system will be exposed to during its use. Even if the LCEP is determined for a system, there may be a new use or application that was not considered during product development and that has significantly different environmental conditions. A good example would be a portable video projector. One population of a particular system may be attached to a room ceiling and have much less shock, vibration, and thermal cycling environmental stress. Another portion of the projectors purchased will be transported regularly by the user to various locations and will have many more mechanical shocks and vibration events, as well as thermal stress variations, compared to ceiling mounted projectors, yet the warranty and reliability expectations of the end user will be the same. The projector LCEP will have a wide distribution of conditions between environments yet the expectations for reliability and warranty coverage are the same regardless of the end use environmental conditions.

## 1.7 Reliability Testing

Reliability testing and assessment has been strongly influenced by FPM as shown in Figure 1.13.

Reliability predictions from FPM guides such as MIL-HDBK-217F, are based on the invalid assumption that the Arrhenius equation applies for many wear-out modes in semiconductors and other electronics components and has resulted in unnecessary costs in additional cooling and the belief that thermal derating during design



**Figure 1.13** Impact of reliability tasks on electronics. Source: Adapted from Pecht and Nash, 1994

provides longer life. It also influences testing regimes with the belief that testing with steady state elevated temperature can provide a quantifiable acceleration of intrinsic wear-out mechanisms in electronics assemblies. There has been no data or evidence to support these beliefs.

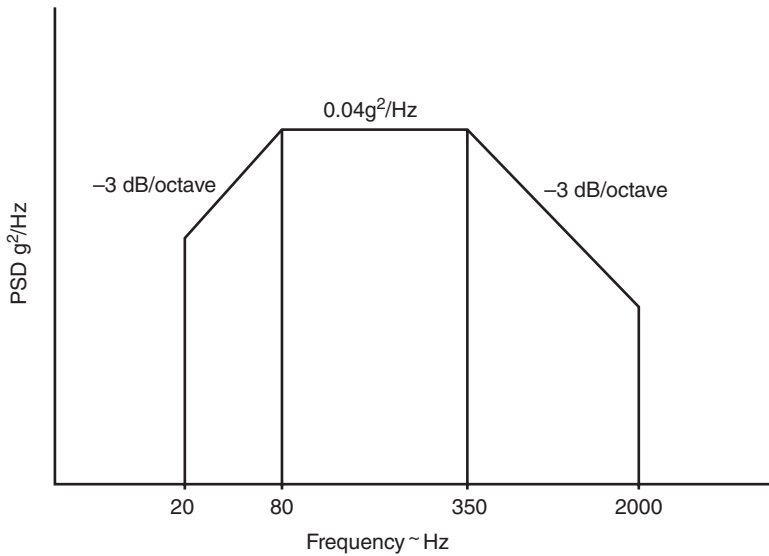
Thermal and vibration stress has long been known to be a very useful stress to find latent defects in electronic hardware. In 1982, Hughes Aircraft published a guide entitled *Stress Screening of Electronic Hardware*, which was an early guide on using environmental stress screening (ESS). The objective of the guide was 'to develop methodologies and techniques for planning, monitoring and evaluating stress screening programs during electronic equipment development and production.'

One very interesting aspect of the development of the environmental stress screening curves shown in the Hughes Aircraft guide was the comparisons of the effectiveness of different stress stimuli used to precipitate the latent defects to patent or detectable defects.

In the Hughes ESS guide they confirm that thermal cycling stress screens and random vibration screens were generally the most effective screens for finding latent defects in electronics systems. They also acknowledge that the industry consensus was that the effectiveness of thermal cycling screens increases with wider temperature ranges and greater rates of change. Additionally it illustrated the industry knowledge that random or broadband vibration is more effective than single or sweep frequency sine vibration.

The vibration regime of a 6 Grms (gravity root mean squared) ESS profile presented in the government publication *Navy Manufacturing Screening Program* (see Figure 1.14) was intended to be a guideline. The 6 Grms vibration profile became the de facto standard auto spectral density (ASD) profile and was applied generically to all systems. Although ESS was a useful new method for finding latent defects, it may have been ineffective for some systems by using too low of stress levels to find defects, and for other systems it may have used stresses severe enough to shorten the products usable life.

HASS processes, like ESS processes, have the identical goal of finding latent defects. The most significant difference between HASS and ESS is how stress levels for a production stress screening process are determined.



**Figure 1.14** The ESS vibration power spectral density spectrum guideline from NAVMAT 9492 (US Navy)

The levels of stress for ESS were determined by ‘stress screening strength’ curves derived from industry consensus regarding the levels of stress needed to precipitate to detection a percentage of latent defects that would be expected per number of components in an assembly or subsystem being screened. In comparison the levels of stress used for HASS encompass a variety of stresses before product is shipped and is uniquely developed based on the product’s empirical strength limits found in the HALT process.

In fact, the UUT (unit under test) in an ESS regime was not typically powered or monitored during the application of stress. Powering and functionally monitoring the UUT is another significant difference between ESS and HALT and HASS. In HALT and HASS, the product may be power cycled and briefly off during the stress application, but should be operating and its function monitored as much as possible during the process.

Many types of latent defects in electronics systems that are likely to become field failures may only be detectable during the application of stress. An example could be a ball grid array (BGA) solder joint that

may have a 100 per cent fracture across the ball, but the surfaces without stress make contact, completing a conduction path that allows the product to operate normally. Only when the surfaces separate under thermomechanical stress or vibration is the conduction path open, which results in a detectable failure if the circuit operation is being monitored at the same time. If tested before and after a HALT without operational monitoring during the application of stress, many of the latent defects and weaknesses could go undetected. In some cases may be necessary to stress a product beyond operational levels for it to provide sufficient acceleration for a latent defect, followed by a lower stress level to operate the UUT for detection of the defects for an effective HASS.

HALT and HASS methods have provided documented cases of detecting operational reliability issues in the field. Many times a marginal system may have a degraded operational reliability from intermittent 'soft failures'. Soft failures are defined as the system failing but recovering normal operation when reset or power cycled. Soft failures may be more prevalent than catastrophic failures in the field, but unless they occur frequently, they may not be recorded, since no hardware needs to be replaced to return the unit to operation.

Many readers may have experienced a screen 'lock up' or 'blue screen of death' operational failure on a personal computer or other personal digital hardware. It can be an annoyance or worse, but it is usually a reason to return the device if it recovers and functions normally when we reboot or power cycle the system. If these 'soft' failures occur frequently enough, the user may return the unit to the manufacturer. It is often that due to the intermittent nature of the failure, the manufacturer will likely declare it 'no defect found' from the limited failure analysis it may have when returned. But the user's perception of overall poor reliability or quality will likely be told to others and may result in the purchase of a different brand when it comes time to upgrade.

As digital systems have been pushing up bus speeds to the gigahertz range and beyond, thermal stress, stepping up the clock frequency and voltage margining to limits will provide more sensitive discriminators to increase the probability of finding software and marginal signal integrity issues that result in operational reliability issues.

Variation in manufacturing causes variations in parametric performance from sample to sample, or lot to lot of electronic components and assemblies. The parametric variations at each assembly level stack up and can lead to timing and signal integrity failures. If the signal integrity is near the margin of failure at room temperature, it may become an intermittent soft failure if operated at higher or lower temperature, but still within the specifications of a design.

Soft failures due to marginal signal integrity can be some of the most challenging to find. It may take hundreds of operational cycles on many samples to reproduce the fault at nominal room conditions. For many engineers performing reliability testing, the potential benefit of stimulating variations of signal propagation and timing may never be realized because the fear that failures from HALT are due to stress levels that the system will never experience in its end use environment, and this is irrelevant and therefore will lead to “over-design”.

If the fears of over-engineering a system are set aside long enough to perform a HALT on a new product, the HALT may demonstrate that a design is very robust and has significant margins. When a weakness is found in a properly run HALT, its relevance to field reliability can be determined and, in most cases, it is relevant. Finding the stress limits provides an opportunity to find and improve the weaknesses that may result in field unreliability, and to establish benchmarks for similar products. Testing to environmental specifications, or expected worst-case conditions, will not accelerate or provide a faster rate of cumulative fatigue over the fielded products that end up being used in a worst case environment. The point of accelerated testing is to find latent defects in electronics that result in failures, so that your customers do not find them. Worst case stress testing will find weaknesses and latent defects in the same time period for products being subject to worst-case end-use environments.

The only way to confirm if a weakness found in HALT is relevant to the field is to ship the units without improving the weakness and wait for failures. Of course this is a significant economic risk for most companies, and for most users of HALT the additional expense of improving weaknesses and possibly “over-designing” a product is much smaller than the potential costs of field failures if the weakness is not addressed.

## 1.8 Traditional Reliability Development

Since the early days of solid state electronics, reliability engineers have been taught that the dominant cause of hardware unreliability comes from component failures and that the reliability of components can be as much as doubled for each 10°C reduction in temperature. This belief was a fundamental tenet of the *U.S. Military Handbook 217* (MIL-HDBK-217),<sup>1</sup> the first document on the reliability prediction of electronic components [5]. While there is no empirical data to support this belief, the concept has persisted and has made its way into other reliability prediction handbooks, such as Telcordia SR-332 (formerly Bellcore), PRISM, FIDES and the Chinese GJB-299. These prediction methods rely on the analysis of insufficient failure data collected from the field, and they assume that the components of a system have inherently constant failure rates that can be derived from the collected data. These methods assume that such constant failure rates could be tailored by independent ‘modifiers’ to account for various quality, operating and temperature parameters.

In the 1990s, with a host of studies conducted by the National Institute of Standards and Technology (NIST) [6], Bell Northern Research [7], the U.S. Army [8], Boeing [9], Honeywell [4], Delco [10], Ford Motor Co. [11], and British Aerospace [12], it became clear that the approach propagated by these handbooks has been damaging to the industry and that a change was needed. The consensus is now that these methods and this type of approach should never be used, because they are inaccurate for predicting actual field failures and they provide highly misleading predictions, which can result in poor designs and poor logistics decisions [13]. Although most of these handbooks have been discontinued and are no longer used by the U.S. military, a few manufacturers of electronic components, printed wiring and circuit boards, and electronic equipment and systems even today still subscribe to the traditional reliability prediction techniques (e.g. MIL-HDBK-217 and its progeny) in some manner, although sometimes unknowingly.

---

<sup>1</sup> The last version of Mil-HDBK 217 was revision ‘F’, in 1995. Since then the document has been cancelled and not updated. Regardless of the fact that the predictions are inaccurate and misleading, it continues to be used have an influential role in reliability engineering.

Electronics systems, especially in the consumer products, have undergone a relatively rapid increase in technological features and benefits. For example, in less than 10 years, the cellular phone industry has gone from a simple portable unit that makes and receives calls to the current smart phones, which are small handheld computers.

When using models to estimate the life entitlement of a component or system certain assumptions must be made that the manufacturing processes are consistent with little variation in its fit or function. Properly manufactured components that are not in a marginal circuit are generally not the cause of the vast majority of hardware failures.

The 'life entitlement' of today's microelectronic components is not known and may never be known, but for most applications it is long beyond any required use time and almost always will reach far beyond the time when the component becomes obsolete.

## Bibliography

- [1] R Development Core Team, *R: A Language and Environment for Statistical Computing*. Vienna, Austria: R Foundation for Statistical Computing, 2009
- [2] Delignette-Muller, M.L, Pouillot, R., Denis, J. and Dutang, C. *fitdistrplus: help to fit of a parametric distribution to censored or non-censored data*. 2013
- [3] Schenkelberg, F. No MTBF-Actions. *NOMTBF.COM*. [Online] 2014. [Cited: 5 May 2014.] <http://nomtbf.com/actions/>.
- [4] Pecht, M.G., Nash, F.R. *Predicting the Reliability of Electronic Equipment*, 1994, Proceedings of the IEEE, pp. 992–1004.
- [5] US Department of Defense. *MIL-HDBK-217: Military Handbook for Reliability Prediction of Electronic Equipment, Version A*. 1965.
- [6] Kopanksi, J., Blackburn, D.L., Harman, G.G., Berning, D.W. *Assessment of Reliability Concerns for Wide-Temperature Operation of Semiconductor Device Circuits*, Albuquerque, NM: Transactions of the First International High Temperature Electronics Conference, 1991.
- [7] Witzmann, S. and Giroux, Y. *Mechanical Integrity of the IC Device Package: A Key Factor in Achieving Failure Free Product Performance*. Albuquerque, NM: Transactions of the First International High Temperature Electronics Conference, pp. 137–142, 1991.
- [8] Cushing, M.J. US Army Reliability Standardization Improvement Policy and its Impact. *IEEE Transactions on Components, Packaging, and Manufacturing Technology*. 1996, Vol. 19.

- [9] Pecht, M. Issues Affecting Early Affordable Access to Leading Electronics Technologies by the US Military and Government. *Circuit World*. 1996, Vol. 22, 2.
- [10] Kleyner, A., Bender, M. s.l.: *Enhanced reliability prediction method based on merging military standards approach with Manufacturers Warranty data*. Annual Reliability and Maintainability Symposium, 2003.
- [11] Derr, James H. Reliability Prediction of Automotive Electronics – How Well Does MIL-HDBK-217-D Stack Up? [Online] 1985. <http://papers.sae.org/850531>.
- [12] O'Conner, P.D.T., *Undue Faith in US MIL-HDBK-217 for Reliability Prediction*. *IEEE Transaction on Reliability*, Vol. 37, p. 468.
- [13] Wong, K.L. What is Wrong with Existing Reliability Prediction Methods? *Quality and Reliability Engineering International*. 1990, Vol. 6.