

Chapter 2

Basic Reliability Mathematics

The basics of mathematical theory that are relevant to the study of reliability and safety engineering are discussed in this chapter. The basic concepts of set theory and probability theory are explained first. Then the elements of component reliability are presented. Different distributions used in reliability and safety studies with suitable examples are explained. The treatment of failure data is given in the last section of the Chapter.

2.1 Classical Set Theory and Boolean Algebra

A set is a collection of elements having certain specific characteristics. A set that contains all elements of interest is known as universal set, denoted by 'U'. A sub set refers to a collection of elements that belong to a universal set. For example, if universal set 'U' represents employees in a company, then female employees is a sub set A of 'U'. For graphical representation of sets within the frame of reference of universal set, Venn diagrams are widely used. They can be very conveniently used to describe various set operations.

The Venn diagram in Fig. 2.1 shows the universal set with a rectangle and subset A with a circle. The complement of a set A (denoted by $\bar{A}$) is a set which consists of the elements of 'U' that do not belong to A.

2.1.1 Operations on Sets

Let A and B be any sub-sets of the universal set U, the union of two sets A and B is a set of all the elements that belong to at least one of the two sets A and B. The union is denoted by ' $\cup$ ' and read as 'OR'. Thus $A \cup B$ is a set that contains all the elements that are in A, B or both A and B. The Venn diagram of $A \cup B$ is shown in Fig. 2.2.

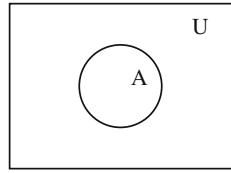


Fig. 2.1 Venn diagram for subset A

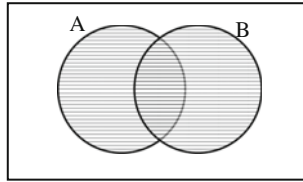


Fig. 2.2 Venn diagram for $A \cup B$

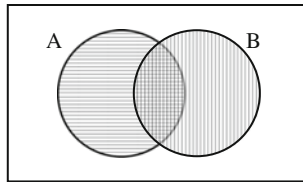


Fig. 2.3 Venn diagram for $A \cap B$

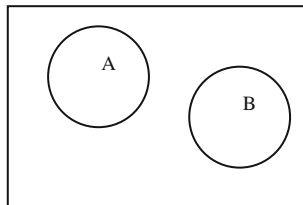


Fig. 2.4 Venn diagram for mutually exclusive events

The intersection of A and B is the set of elements which belong to both sets. The intersection is denoted by ' $\cap$ ' and read as 'AND'. The Venn diagram of $A \cap B$ is shown in Fig. 2.3.

Two sets of A and B are termed mutually exclusive or disjoint sets when A and B have no elements in common i.e., $A \cap B = \emptyset$. This can be represented by Venn diagram as shown in Fig. 2.4.

Table 2.1 Laws of set theory

Name of the law	Description
Identity law	$A \cup \emptyset = A; A \cup U = U$
	$A \cap \emptyset = \emptyset; A \cap U = A$
Idempotency law	$A \cup A = A$
	$A \cap A = A$
Commutative law	$A \cup B = B \cup A$
	$A \cap B = B \cap A$
Associative law	$A \cup (B \cup C) = (A \cup B) \cup C$
	$A \cap (B \cap C) = (A \cap B) \cap C$
Distributive law	$A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$
	$A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$
Complementation law	$A \cup \bar{A} = U$
	$A \cap \bar{A} = \Phi$
	$\bar{\bar{A}} = A$
De Morgan's laws	$\overline{(A \cup B)} = \bar{A} \cap \bar{B}$
	$\overline{(A \cap B)} = \bar{A} \cup \bar{B}$

2.1.2 Laws of Set Theory

Some important laws of set theory are enumerated in the Table 2.1.

2.1.3 Boolean Algebra

Boolean algebra finds its extensive use in evaluation of reliability and safety procedures due to consideration that components and system can present in either success or failure state. Consider a variable 'X' denotes the state of a component and assuming 1 represents success and 0 represents failure state. Then, probability that X is equal to 1 $P(X = 1)$ is called reliability of that particular component. Depending upon the configuration of the system, it will also have success or failure state. Based on this binary state assumption, Boolean algebra can be conveniently used.

In Boolean algebra all the variables must have one of two values, either 1 or 0. There are three Boolean operations, namely, OR, AND and NOT. These operations are denoted by +, . (dot) and $\bar{}$ (super bar over the variable) respectively. A set of postulates and useful theorems are listed in Table 2.2. X denotes a set and x_1, x_2, x_3 denote variables of X.

Consider a function of $f(x_1, x_2, x_3, \dots, x_n)$ of n variables, which are combined by Boolean operations. Depending upon the values of constituent variables $x_1, x_2, \dots, x_n$, function f will be 1 or 0. As these are n variables and each can have two possible values 1 or 0, 2^n combinations of variables will have to be considered for determination of the value of function f. Truth tables are used represent the value of f for

Table 2.2 Boolean algebra theorems

Postulate/Theorem	Remarks
$x + 0 = x$ $x \cdot 1 = x$	Identity
$x + x = x$ $x \cdot x = x$	Idempotence
$\bar{0} = 1$ and $\bar{1} = 0$ $\bar{\bar{x}} = x$	Involution
$x_1 + x_1x_2 = x_1$ $x_1(x_1 + x_2) = x_1$	Absorption
$x_1 + (x_2 + x_3) = (x_1 + x_2) + x_3$ $x_1 \cdot (x_2 \cdot x_3) = (x_1 \cdot x_2) \cdot x_3$	Associative
$\overline{(x_1 + x_2)} = \bar{x}_1 \cdot \bar{x}_2$ $\overline{(x_1 \cdot x_2)} = \bar{x}_1 + \bar{x}_2$	De Morgan's theorem

Table 2.3 Truth table

x_1	x_2	x_3	F
0	0	0	0
0	0	1	0
0	1	0	0
0	1	1	1
1	0	0	0
1	0	1	1
1	1	0	1
1	1	1	1

all these combinations. A truth table is given for a Boolean expression $f(x_1, x_2, x_3) = x_1x_2 + x_2x_3 + x_1x_3$ in the following Table 2.3.

In reliability calculations, it is necessary to minimize the Boolean expression in order to eliminate repetition of the same elements. The premise of all minimization techniques is the set of Boolean algebra theorems mentioned in the Table 2.2. The amount of labor involved in minimization increases as the number of variable increase. Geometric methods and famous Karnaugh's map is applicable only up to six number of variables. Nowadays, sophisticated computerized algorithms are available for calculation with large number of variables.

2.2 Concepts of Probability Theory

The word 'experiment' is used in probability and statistics to describe any process of observation that generates raw data. An experiment becomes 'random experiment' if it satisfies the following conditions: it can be repeatable, outcome is

random (though it is possible to describe all the possible outcomes) and pattern of occurrence is definite if the experiment is repeated large number of times. Examples of random experiment are tossing of coin, rolling die, and failure times of engineering equipment from its life testing. The set of all possible outcomes of a random experiment is known as 'sample space' and is denoted by 'S'. The sample space for random experiment of rolling a die is {1, 2, 3, 4, 5, and 6}. In case of life testing of engineering equipment, sample space is from 0 to ∞ . Any subset of sample space is known as an event 'E'. If the outcome of the random experiment is contained in E then once can say that E has occurred. Probability is used to quantify the likelihood, or chance, that an outcome of a random experiment will occur. Probability is associated with any event E of a sample space S depending upon its chance of occurrence which is obtained from available data or information.

The concept of the probability of a particular event is subject to various meanings or interpretations. There are mainly three interpretations of probability: classical, frequency, and subjective interpretations.

The classical interpretation of probability is based on the notion of equally likely outcomes and was originally developed in the context of games of chance in the early days of probability theory. Here the probability of an event E is equal to the number of outcomes comprising that event (n) divided by the total number of possible outcomes (N). This interpretation is simple, intuitively appealing, and easy to implement, but its applicability is, of course, limited by its restriction to equally likely outcomes. Mathematically, it is expressed as follows:

$$P(E) = \frac{n}{N} \quad (2.1)$$

The relative-frequency interpretation of probability defines the probability of an event in terms of the proportion of times the event occurs in a long series of identical trials. In principle, this interpretation seems quite sensible. In practice, its use requires extensive data, which in many cases are simply not available and in other cases may be questionable in terms of what can be viewed as 'identical trials'. Mathematically, it is expressed as follows;

$$P(E) = \lim_{N \rightarrow \infty} \frac{n}{N} \quad (2.2)$$

The subjective interpretation of probability views probability as a degree of belief, and this notion can be defined operationally by having an individual make certain comparisons among lotteries. By its very nature, a subjective probability is the probability of a particular person. This implies, of course, that different people can have different probabilities for the same event. The fact that subjective probabilities can be manipulated according to the usual mathematical rules of probability is not transparent but can be shown to follow from an underlying axiomatic framework.

Regardless of which interpretation one gives to probability, there is general consensus that the mathematics of probability is the same in all cases.

2.2.1 Axioms of Probability

Probability is a number that is assigned to each member of a collection of events from a random experiment that satisfies the following properties:

If S is the sample space and E is any event in a random experiment,

1. $P(S) = 1$
2. $0 \leq P(E) \leq 1$
3. For two events E_1 and E_2 with $E_1 \cap E_2 = \emptyset$, $P(E_1 \cup E_2) = P(E_1) + P(E_2)$

The property that $0 \leq P(E) \leq 1$ is equivalent to the requirement that a relative frequency must be between 0 and 1. The property that $P(S) = 1$ is a consequence of the fact that an outcome from the sample space occurs on every trial of an experiment. Consequently, the relative frequency of S is 1. Property 3 implies that if the events E_1 and E_2 have no outcomes in common, the relative frequency of outcomes in is the sum of the relative frequencies of the outcomes in E_1 and E_2 .

2.2.2 Calculus of Probability Theory

Independent Events and Mutually Exclusive Events

Two events are said to be ‘independent’ if the occurrence of one does not affect the probability of occurrence of other event. Let us say A and B are two events, if the occurrence of A does not provide any information about occurrence of B then A and B are statistically independent. For example in a process plant, the failure of a pump does not affect the failure of a valve.

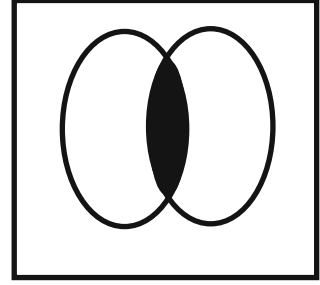
Two events are said to be ‘mutually exclusive’ if the occurrence of one event makes the non-occurrence of other event. If the occurrence of A ensures that B will not happen then A and B are mutually exclusive. If two events are mutually exclusive then they are dependent events. Success and failure events of any component are mutually exclusive. In a given time, if pump is successfully operating implies failure has not taken place.

Conditional Probability

The concept of conditional probability is the most important in all of probability theory. It is often interest to calculate probabilities when some partial information concerning the result of the experiment is available, or in recalculating them in the light of additional information. Let there be two event A and B , the probability of A given that B has occurred is referred as conditional probability and is denoted by $P(A|B) = P(A \cap B)/P(B)$.

If the event B occurs then in order for A to occur it is necessary that the actual occurrence be a point in both A and B , i.e. it must be in $A \cap B$ (Fig. 2.5). Now, since we know that B has occurred, it follows that B becomes our new sample space and hence the probability that the event $A \cap B$ occurs will equal the probability of $A \cap B$ relative to the probability of B . It is mathematical expressed as,

Fig. 2.5 Venn diagram for $A \cap B$



$$P(A|B) = \frac{P(A \cap B)}{P(B)} \quad (2.3)$$

Similarly one can write

$$P(B|A) = \frac{P(A \cap B)}{P(A)} \quad (2.4)$$

Probability for Intersection of Events

From Eq. 2.4, one can write

$$P(A \cap B) = P(A) \times P(B|A) \quad (2.5)$$

If A and B are independent events then the conditional probability $P(B|A)$ is equal to $P(B)$ only. Now Eq. 2.5 becomes, simply the product of probability of A and probability of B.

$$P(A \cap B) = P(A) \times P(B) \quad (2.6)$$

Thus when A and B are independent, the probability that A and B occur together is simply the product of the probabilities that A and B occur individually.

In general the probability of occurrence of n dependent events $E_1, E_2, \dots, E_n$ is calculated by the following expression,

$$P(E_1 \cap E_2 \cap \dots \cap E_n) = P(E_1) \times P(E_2|E_1) \times P(E_3|E_1 \cap E_2) \dots P(E_n|E_1 \cap E_2 \cap \dots \cap E_{n-1})$$

If all the events are independent then probability of joint occurrence is simply the product of individual probabilities of events.

$$P(E_1 \cap E_2 \cap \dots \cap E_n) = P(E_1) \times P(E_2) \times P(E_3) \times \dots \times P(E_n) \quad (2.7)$$

Probability for Union of Events

Let A and B are two events. From the Venn diagram (Fig. 2.6), as the three regions 1, 2 and 3 are mutually exclusive, it follows that

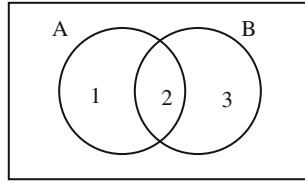


Fig. 2.6 Venn diagram for A and B

$$\begin{aligned}
 P(A \cup B) &= P(1) + P(2) + P(3) \\
 P(A) &= P(1) + P(2) \\
 P(B) &= P(2) + P(3) \\
 \text{which shows that} & \\
 P(A \cup B) &= P(A) + P(B) - P(2) \\
 \text{As } P(2) &= P(A \cap B), \\
 P(A \cup B) &= P(A) + P(B) - P(A \cap B)
 \end{aligned} \tag{2.8}$$

The above expression can be extended to n events $E_1, E_2, \dots, E_n$ by the following equation

$$\begin{aligned}
 P(E_1 \cup E_2 \cup \dots \cup E_n) &= P(E_1) + P(E_2) + \dots + P(E_n) \\
 &\quad - [P(E_1 \cap E_2) + P(E_2 \cap E_3) + \dots + P(E_{n-1} \cap E_n)] + \\
 &\quad + [P(E_1 \cap E_2 \cap E_3) + P(E_2 \cap E_3 \cap E_4) + \dots + P(E_{n-2} \cap E_{n-1} \cap E_n)] - \\
 &\quad \vdots \\
 &\quad (-1)^{n+1} P(E_1 \cap E_2 \cap \dots \cap E_n)
 \end{aligned} \tag{2.9}$$

Total Probability Theorem

Let $A_1, A_2 \dots A_n$ be n mutually exclusive events forming a sample space S and $P(A_i) > 0, i = 1, 2 \dots n$ (Fig. 2.7). For an arbitrary event B one has

$$\begin{aligned}
 B &= B \cap S = B \cap (A_1 \cup A_2 \cup \dots \cup A_n) \\
 &= (B \cap A_1) \cup (B \cap A_2) \cup \dots \cup (B \cap A_n)
 \end{aligned}$$

where the events $B \cap A_1, B \cap A_2, \dots, B \cap A_n$ are mutually exclusive.

$$P(B) = \sum_i P(B \cap A_i) = \sum_i P(A_i)P(B|A_i) \tag{2.10}$$

This is called total probability theorem.

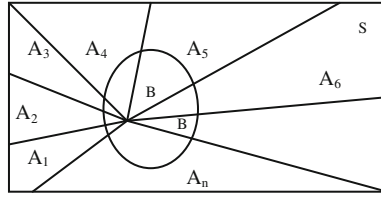


Fig. 2.7 Sample space containing n mutually exclusive events

Bayes Theorem

From the definitions of conditional probability,

$$P(A|B) = \frac{P(A \cap B)}{P(B)}$$

$$P(A \cap B) = P(B) \times P(A|B) - (a)$$

$$P(B|A) = \frac{P(A \cap B)}{P(A)}$$

$$P(A \cap B) = P(A) \times P(B|A) - (b)$$

Equating both (a) and (b) we have: $P(B) \times P(A|B) = P(A) \times P(B|A)$.

We can obtain $P(A|B)$ as follows

$$P(A|B) = \frac{P(A) \times P(B|A)}{P(B)} \quad (2.11)$$

This is a useful result that enables us to solve for $P(A|B)$ in terms of $P(B|A)$.

In general, if $P(B)$ is written using the Total Probability theorem, we obtain the following general result, which is known as *Bayes' Theorem*.

$$P(A_i|B) = \frac{P(A_i)P(B|A_i)}{\sum_i P(A_i)P(B|A_i)} \quad (2.12)$$

Bayes' theorem presents a way to evaluate posterior probabilities $P(A_i|B)$ in terms of prior probabilities $P(A_i)$ and conditional probabilities $P(B|A_i)$. This is very useful in updating failure data as more evidence is available from operating experience.

The basic concepts of probability and statistics are explained in detail in the Refs. [1, 2].

2.2.3 Random Variables and Probability Distributions

It is important to represent the outcome from a random experiment by a simple number. In some cases, descriptions of outcomes are sufficient, but in other cases, it is useful to associate a number with each outcome in the sample space. Because the particular outcome of the experiment is not known in advance, the resulting value of our variable is not known in advance. For this reason, random variables are used to associate a number with the outcome of a random experiment. A random variable is defined as a function that assigns a real number to each outcome in the sample space of a random experiment. A random variable is denoted by a capital letter and numerical value that it can take is represented by a small letter. For example, if X is a random variable representing number of power outages in a plant, then x shows the actual number of outages it can take say 0, 1, 2...n.

Random variable can be classified into two categories, namely, discrete and continuous random variables. A random variable is said to be discrete if its sample space is countable. The number of power outages in a plant in a specified time is discrete random variable. If the elements of the sample space are infinite in number and sample space is continuous, the random variable defined over such a sample space is known as continuous random variable. If the data is countable then it is represented with discrete random variable and if the data is measurable quantity then it is represented with continuous random variable.

Discrete Probability Distribution

The probability distribution of a random variable X is a description of the probabilities associated with the possible values of X . For a discrete random variable, the distribution is often specified by just a list of the possible values along with the probability of each. In some cases, it is convenient to express the probability in terms of a formula.

Let X be a discrete random variable defined over a sample space $S = \{x_1, x_2 \dots x_n\}$. Probability can be assigned to each value of sample space S . It is usually denoted by $f(x)$. For a discrete random variable X , a probability distribution is a function such that

$$(a) \quad f(x_i) \geq 0$$

$$(b) \quad \sum_{i=1}^n f(x_i) = 1$$

$$(c) \quad f(x_i) = P(X = x_i)$$

Probability distribution is also known as probability mass function. Some examples are Binomial, Poisson, Geometric distributions. The graph of a discrete probability distribution looks like a bar chart or histogram. For example, in five flips of a coin, where X represents the number of heads obtained, the probability mass function is shown in Fig. 2.8.

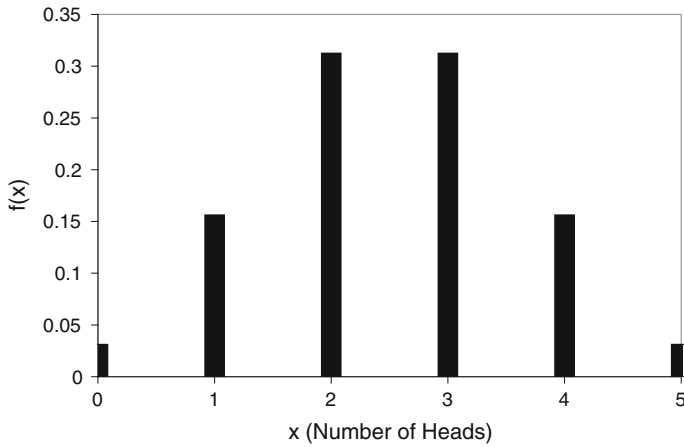


Fig. 2.8 A discrete probability mass function

The cumulative distribution function of a discrete random variable X , denoted as $F(x)$, is

$$F(x) = P(X \leq x) = \sum_{x_i \leq x} f(x_i)$$

$F(x)$ satisfies the following properties for a discrete random variable X .

- (a) $F(x) = P(X \leq x) = \sum_{x_i \leq x} f(x_i)$
- (b) $0 \leq F(x) \leq 1$
- (c) if $x \leq y$ then $F(x) \leq F(y)$

The cumulative distribution for the coin flipping example is given in Fig. 2.9.

Continuous Probability Distributions

As the elements of sample space for a continuous random variable X are infinite in number, probability of assuming exactly any of its possible values is zero. Density functions are commonly used in engineering to describe physical systems. Similarly, a probability density function $f(x)$ can be used to describe the probability distribution of a continuous random variable X . If an interval is likely to contain a value for X , its probability is large and it corresponds to large values for $f(x)$. The probability that X is between a and b is determined as the integral of $f(x)$ from a to b . For a continuous random variable X , a probability density function is a function such that

- (a) $f(x) \geq 0$
- (b) $\int_{-\infty}^{+\infty} f(x) = 1$
- (c) $P(a \leq X \leq b) = \int_a^b f(x) dx$

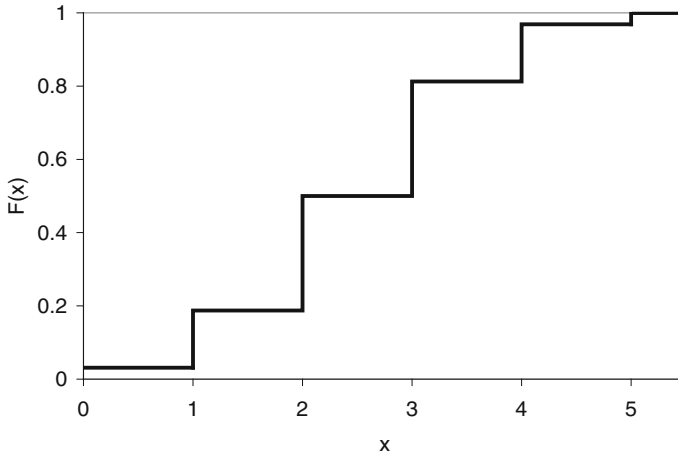


Fig. 2.9 A discrete cumulative distribution function

The cumulative distribution function of a continuous random variable X is

$$F(x) = P(X \leq x) = \int_{-\infty}^x f(\theta) d\theta \quad (2.13)$$

The probability density function of a continuous random variable can be determined from the cumulative distribution function by differentiating. Recall that the fundamental theorem of calculus states that

$$\frac{d}{dx} \int_{-\infty}^x f(\theta) d\theta = f(x)$$

Now differentiating $F(x)$ with respect to x and rearranging for $f(x)$

$$f(x) = \frac{dF(x)}{dx} \quad (2.14)$$

Characteristics of Random Variables

In order to represent probability distribution function of a random variable, some characteristic values such as expectation (mean) and variance are widely used. Expectation or mean value represents the central tendency of a distribution function. It is mathematically expressed as

$$\begin{aligned}
 \text{Mean} = E(x) &= \sum_i x_i f(x_i) \quad \text{for discrete} \\
 &= \int_{-\infty}^{+\infty} x f(x) dx \quad \text{for continuous}
 \end{aligned}$$

A measure of dispersion or variation of probability distribution is represented by variance. It is also known as central moment or second moment about the mean. It is mathematically expressed as

$$\begin{aligned}
 \text{Variance} = E((x - \text{mean})^2) &= \sum_x (x - \text{mean})^2 f(x) \quad \text{for discrete} \\
 &= \int_{-\infty}^{+\infty} (x - \text{mean})^2 f(x) dx \quad \text{for continuous}
 \end{aligned}$$

2.3 Reliability and Hazard Functions

Let 'T' be a random variable representing time to failure of a component or system. Reliability is probability that the system will perform its expected job under specified conditions of environment over a specified period of time. Mathematically, reliability can be expressed as the probability that time to failure of the component or system is greater than or equal to a specified period of time (t).

$$R(t) = P(T \geq t) \quad (2.15)$$

As reliability denotes failure free operation, it can be termed as success probability. Conversely, probability that failure occurs before the time t is called failure probability or unreliability. Failure probability can be mathematically expressed as the probability that time to failure occurs before a specified period of time t.

$$\bar{R}(t) = P(T < t) \quad (2.16)$$

As per the probability terminology, $\bar{R}(t)$ is same as the cumulative distributive function of the random variable T.

$$F(t) = \bar{R}(t) = P(T < t) \quad (2.17)$$

Going by the first axiom of probability, probability of sample space is unity. The sample space for the continuous random variable T is from 0 to ∞ . Mathematically, it is expressed as

$$P(S) = 1$$

$$P(0 < T < \infty) = 1$$

The sample space can be made two mutually exclusive intervals: one is $T < t$ and the second is $T \geq t$. Using third axiom of probability, we can write

$$P(0 < T < \infty) = 1$$

$$P(T < t \cup T \geq t) = 1$$

$$P(T < t) + P(T \geq t) = 1$$

Substituting Eqs. 2.15 and 2.17, we have

$$F(t) + R(t) = 1 \quad (2.18)$$

As the time to failure is a continuous random variable, the probability of T having exactly a precise t will be approximately zero. In this situation, it is appropriate to introduce the probability associated with a small range of values that the random variable can take on.

$$P(t < T < t + \Delta t) = F(t + \Delta t) - F(t)$$

Probability density function $f(t)$ for continuous random variables is defined as

$$\begin{aligned} f(t) &= \lim_{\Delta t \rightarrow 0} \left[\frac{P(t < T < t + \Delta t)}{\Delta t} \right] \\ &= \lim_{\Delta t \rightarrow 0} \left[\frac{F(t + \Delta t) - F(t)}{\Delta t} \right] \\ &= \frac{dF(t)}{dt} \\ &= -\frac{dR(t)}{dt} \quad (\text{from Eq. 2.18}) \end{aligned}$$

From the above derivation we have an important relation between $R(t)$, $F(t)$ and $f(t)$:

$$f(t) = \frac{dF(t)}{dt} = -\frac{dR(t)}{dt} \quad (2.19)$$

Given the Probability Density Function (PDF), $f(t)$ (Fig. 2.10), then

$$F(t) = \int_0^t f(t)dt$$

$$R(t) = \int_t^{\infty} f(t)dt \quad (2.20)$$

The conditional probability of a failure in the time interval from t to $(t + \Delta t)$ given that the system has survived to time t is

$$P(t \leq T \leq t + \Delta t | T \geq t) = \frac{R(t) - R(t + \Delta t)}{R(t)}$$

Then $\frac{R(t) - R(t + \Delta t)}{R(t)\Delta t}$ is the conditional probability of failure per unit of time (failure rate).

$$\begin{aligned} \lambda(t) &= \lim_{\Delta t \rightarrow 0} \frac{R(t) - R(t + \Delta t)}{R(t)\Delta t} = \lim_{\Delta t \rightarrow 0} \frac{-[R(t + \Delta t) - R(t)]}{\Delta t} \frac{1}{R(t)} \\ &= \frac{-dR(t)}{dt} \frac{1}{R(t)} = \frac{f(t)}{R(t)} \end{aligned} \quad (2.21)$$

$\lambda(t)$ is known as the instantaneous hazard rate or failure rate function.

Reliability as a function of hazard rate function can be derived as follows: We have the following relation from the above expression

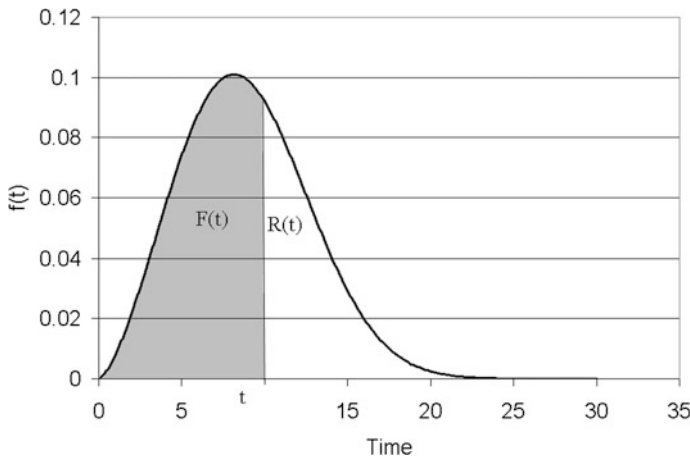


Fig. 2.10 Probability distribution function

$$\lambda(t) = \frac{-dR(t)}{dt} \frac{1}{R(t)}$$

$$\lambda(t)dt = \frac{-dR(t)}{R(t)}$$

Integrating and simplifying, we have

$$R(t) = \exp \left[- \int_0^t \lambda(\theta) d\theta \right] \quad (2.22)$$

2.4 Distributions Used in Reliability and Safety Studies

This section provides the most important probability distributions used in reliability and safety studies. They are grouped into two categories, namely, discrete probability distributions and continuous probability distributions.

2.4.1 Discrete Probability Distributions

2.4.1.1 Binomial Distribution

Consider a trial in which the only outcome is either success or failure. A random variable X with this trial can have success ($X = 1$) or failure ($X = 0$). The random variable X is said to be Bernoulli random variable if the probability mass function of X is given by

$$P(X = 1) = p$$

$$P(X = 0) = 1 - p$$

where p is the probability that the trial is success. Suppose now that n independent trials, each of which results in a ‘success’ with probability ‘ p ’ and in a ‘failure’ with probability $1 - p$, are to be performed. If X represents the number of success that occur in the n trials, then X is said to be a binomial random variable with parameters n, p . The probability mass function of binomial random variable is given by

$$P(X = i) = {}^n C_i p^i (1 - p)^{n-i} \quad i = 0, 1, 2, \dots, n \quad (2.23)$$

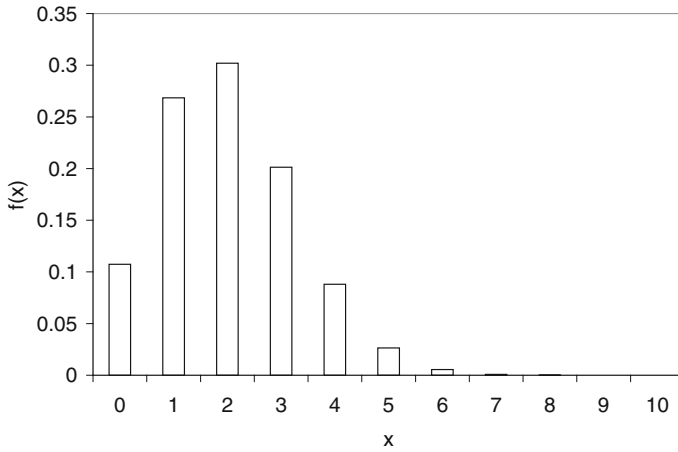


Fig. 2.11 Binomial probability mass function

The probability mass function of a binomial random variable with parameter $(10, 0.2)$ is presented in Fig. 2.11.

The cumulative distributive function is given by

$$P(X \leq i) = \sum_{j=0}^i {}^n c_j p^j (1-p)^{n-j} \quad (2.24)$$

Mean of the binomial distribution is calculated as follows

$$\begin{aligned}
 E(x) &= \sum x f(x) \\
 &= \sum_{i=0}^n i \times {}^n c_i p^i (1-p)^{n-i} \\
 &= np \sum_{i=1}^n {}^{n-1} c_{i-1} p^{i-1} (1-p)^{n-i} \\
 &= np \sum_{j=0}^{n-1} {}^m c_j p^j (1-p)^{m-j} \\
 &= np
 \end{aligned}$$

Similarly variance can also be derived as

$$\text{Variance} = npq$$

Example 1 It has been known from the experience that 4 % of hard disks produced by a computer manufacture are defective. Find the probability that out of 50 disks tested, what is the probability of having (i) Zero Defects and (ii) All are defective.

Solution: $q = 4\%$ of hard disks produced by a computer manufacture are defective.

We know,

$$\begin{aligned} p + q &= 1 \\ p &= 1 - q \\ &= 1 - 0.04 \\ p &= 0.96 \end{aligned}$$

According to Binomial Distribution,

$$P(X = x) = {}^nC_x \cdot p^x \cdot q^{n-x}$$

Now,

(i) In case of 'zero defects', i.e. $p(X = 0)$

$$P(X = 0) = {}^nC_x \cdot p^x \cdot q^{n-x} = {}^{50}C_0 \cdot (0.04)^0 \cdot (0.96)^{(50-0)} = 0.1299$$

(ii) In case of 'all are defective', i.e. $p(X = 50)$

$$P(X = 50) = {}^nC_x \cdot p^x \cdot q^{n-x} = {}^{50}C_{50} (0.04)^{50} (0.96)^{(50-50)} = 0.8701$$

Or in other way,

$$P(X = 50) = 1 - P(X = 0) = 1 - 0.1299 = 0.8701$$

Example 2 To ensure high reliability, triple modular¹ redundancy is adopted in instrumentation systems of Nuclear Power Plant (NPP). It is known that failure probability of each instrumentation channel from operating experience is 0.01. What is the probability of success of the whole instrumentation system?

Solution: $q =$ failure probability from operation experience is 0.01.

We know, $p = 1 - q = 1 - 0.01 = 0.99$

According to Binomial Distribution,

¹Triple modular redundancy denotes at least 2 instruments should be success out of 3 instruments.

Table 2.4 Calculations

	Formula	Numerical solutions	Value
(i)	$P(X = 0) = nC_x \cdot p^x \cdot q^{n-x}$	$P(0) = 3C_0 (0.99)^0 \cdot (0.01)^{(3-0)}$	$P(0) = 1e-6$
(ii)	$P(X = 1) = nC_x \cdot p^x \cdot q^{n-x}$	$P(0) = 3C_1 (0.99)^1 \cdot (0.01)^{(3-1)}$	$P(1) = 2.9e-4$
(iii)	$P(X = 2) = nC_x \cdot p^x \cdot q^{n-x}$	$P(0) = 3C_2 (0.99)^2 \cdot (0.01)^{(3-2)}$	$P(2) = 2.9e-2$
(iv)	$P(X = 3) = nC_x \cdot p^x \cdot q^{n-x}$	$P(0) = 3C_3 (0.99)^3 \cdot (0.01)^{(3-3)}$	$P(3) = 0.97$

$$P(X = x) = nC_x \cdot p^x \cdot q^{n-x}$$

The sample space is then developed as in Table 2.4.

Now the failure probability is sum of (i) and (ii), which is obtained as 2.98e-4 and the success probability is sum of (iii) and (iv), which is obtained as 0.999702.

2.4.1.2 Poisson Distribution

Poisson distribution is useful to model when the event occurrences are discrete and the interval is continuous. For a trial to be a Poisson process, it has to satisfy the following conditions:

1. The probability of occurrence of one event in time Δt is $\lambda \Delta t$ where λ is constant
2. The probability of more than one occurrence is negligible in interval Δt
3. Each occurrence is independent of all other occurrences

A random variable X is said to have Poisson distribution if the probability distribution is given by

$$f(x) = \frac{e^{-\lambda t} (\lambda t)^x}{x!} \quad x = 0, 1, 2, \dots \quad (2.25)$$

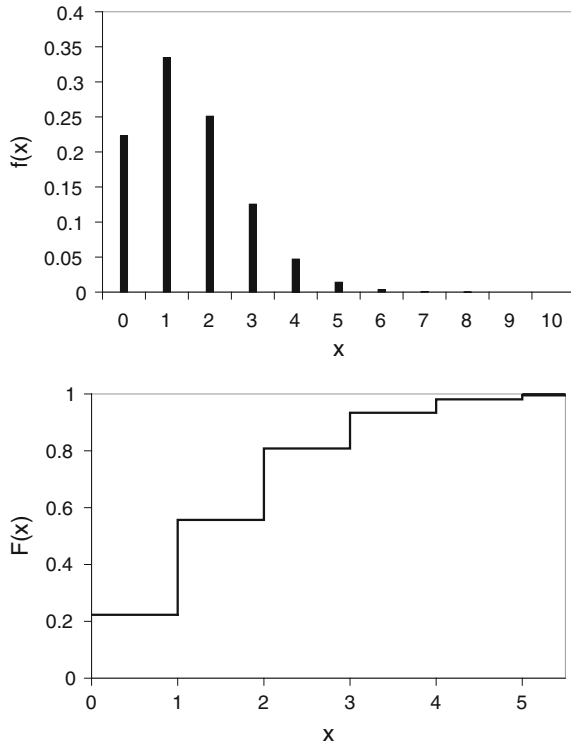
λ is known as average occurrence rate and x is number of occurrences of Poisson events.

The cumulative distribution function is given by

$$F(x) = \sum_{i=0}^x f(X = i) \quad (2.26)$$

The probability mass function and CDF for $\lambda = 1.5/\text{year}$ and $t = 1$ year are shown in Fig. 2.12. Both the mean and variance of Poisson distribution is λt .

Fig. 2.12 Probability functions for Poisson distribution



If the probability of occurrence is near zero and sample size very large, the Poisson distribution may be used to approximate Binomial distribution.

Example 3 If the rate of failure for an item is twice a year, what is the probability that no failure will happen over a period of 2 years?

Solution: Rate of failure, denoted as $\lambda = 2/\text{year}$

Time $t = 2$ years

The Poisson probability mass function is expressed as

$$f(x) = \frac{e^{-\lambda t} (\lambda t)^x}{x!}$$

In a case of no failures, $x = 0$, which leads to

$$f(X = 0) = \frac{e^{-2 \times 2} (2 \times 2)^0}{0!} = 0.0183$$

2.4.1.3 Hyper Geometric Distribution

The hyper geometric distribution is closely related with binomial distribution. In hyper geometric distribution, a random sample of 'n' items is chosen from a finite population of N items. If N is very large with respect to n, the binomial distribution is good approximation of the hyper geometric distribution. The random variable 'X' denote x number of successes in the random sample of size 'n' from population N containing k number of items labeled success. The hyper geometric distribution probability mass function is

$$f(x) = p(x, N, n, k) = \frac{{}^K C_x \cdot {}^{N-k} C_{n-x}}{{}^N C_n}, \quad x = 0, 1, 2, 3, 4, \dots, n. \quad (2.27)$$

The mean of hyper geometric distribution is

$$E(x) = \frac{n \cdot K}{N} \quad (2.28)$$

The variance of hyper geometric distribution is

$$V(x) = \left(\frac{n \cdot K}{N} \right) \left(1 - \frac{K}{N} \right) \left(\frac{N - n}{N - 1} \right) \quad (2.29)$$

2.4.1.4 Geometric Distribution

In case of binomial and hyper geometric distribution, the number of trials 'n' is fixed and number of successes is random variable. Geometric distribution is used if one is interested in number of trials required to obtain the first success. The random variable in geometric distribution is number of trials required to get the first success.

The geometric distribution probability mass function is

$$f(x) = P(x; p) = p(1 - p)^{x-1}, \quad x = 1, 2, 3, \dots, n. \quad (2.30)$$

where 'p' is the probability of success on a style trials.

The mean of geometric distribution is

$$E(x) = \frac{1}{p}$$

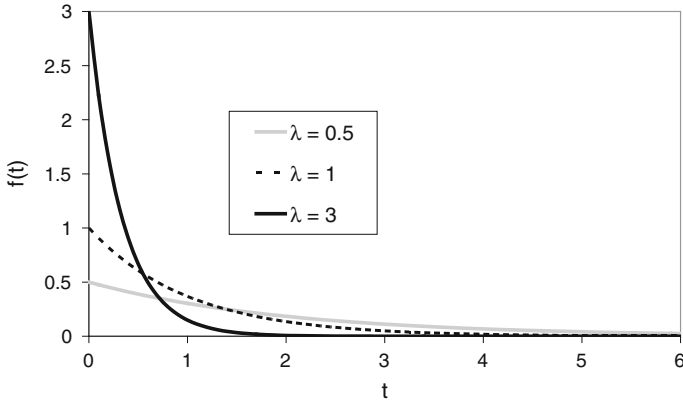


Fig. 2.13 Exponential probability density functions

The variable of geometric distribution is

$$V(x) = \frac{1-p}{p^2}$$

The geometric distribution is the only discrete distribution which exhibits the memory less property, as does the exponential distribution in the continuous case.

2.4.2 Continuous Probability Distributions

2.4.2.1 Exponential Distribution

The exponential distribution is most widely used distribution in reliability and risk assessment. It is the only distribution having constant hazard rate and is used to model ‘useful life’ of many engineering systems. The exponential distribution is closely related with the Poisson distribution which is discrete. If the number of failure per unit time is Poisson distribution then the time between failures follows exponential distribution. The probability density function (PDF) of exponential distribution is

$$\begin{aligned} f(t) &= \lambda e^{-\lambda t} \quad \text{for } 0 \leq t \leq \infty \\ &= 0 \quad \text{for } t < 0 \end{aligned} \quad (2.31)$$

The exponential probability density functions are shown in Fig. 2.13 for different values of λ .

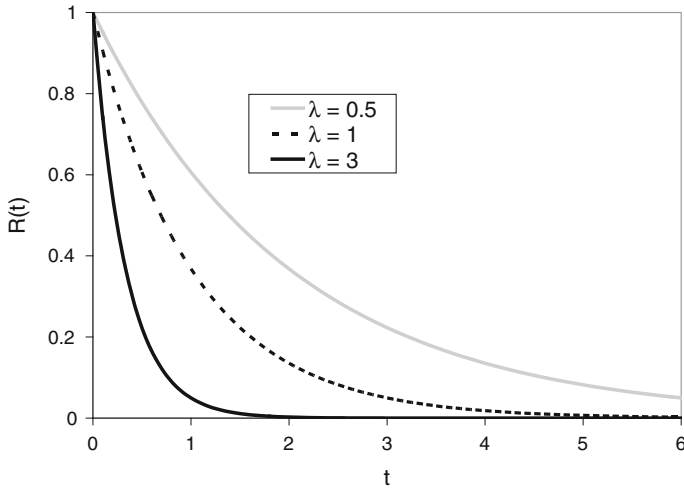


Fig. 2.14 Exponential reliability functions

The exponential cumulative distribution function can be derived from its PDF as follows,

$$F(t) = \int_0^t f(t)dt = \int_0^t \lambda e^{-\lambda t} dt = \lambda \left[\frac{e^{-\lambda t}}{-\lambda} \right]_0^t = \lambda \left[\frac{e^{-\lambda t}}{-\lambda} - \frac{1}{-\lambda} \right] = 1 - e^{-\lambda t} \quad (2.32)$$

Reliability function is complement of cumulative distribution function

$$R(t) = 1 - F(t) = e^{-\lambda t} \quad (2.33)$$

The exponential reliability functions are shown in Fig. 2.14 for different values of λ .

Hazard function is ratio of PDF and its reliability function, for exponential distribution it is

$$h(t) = \frac{f(t)}{R(t)} = \frac{\lambda e^{-\lambda t}}{e^{-\lambda t}} = \lambda \quad (2.34)$$

The exponential hazard function is constant λ . This is reason for memory less property for exponential distribution. Memory less property means the probability of failure in a specific time interval is the same regardless of the starting point of that time interval.

Mean and Variance of Exponential Distribution

$$\begin{aligned} E(t) &= \int_0^{\infty} tf(t) \\ &= \int_0^{\infty} t\lambda e^{-\lambda t} dt \end{aligned}$$

Using integration by parts formula ($\int u dv = uv - \int v du$)

$$\begin{aligned} E(t) &= \lambda \left[t \cdot \frac{e^{-\lambda t}}{-\lambda} \Big|_0^{\infty} - \int_0^{\infty} \frac{e^{-\lambda t}}{-\lambda} dt \right] \\ &= \lambda \left[0 + \frac{1}{\lambda} \left(\frac{e^{-\lambda t}}{-\lambda} \Big|_0^{\infty} \right) \right] = \lambda \left[\frac{1}{\lambda} \left(\frac{1}{\lambda} \right) \right] = \frac{1}{\lambda} \end{aligned}$$

Thus mean time to failure of exponential distribution is reciprocal of failure rate.
Variance(t) = $E(T^2) - (\text{mean})^2$

$$E(T^2) = \int_0^{\infty} t^2 f(t) dt = \int_0^{\infty} t^2 \lambda e^{-\lambda t} dt$$

Using integration by parts formula

$$\begin{aligned} E(T^2) &= \lambda \left[t^2 \cdot \frac{e^{-\lambda t}}{-\lambda} \Big|_0^{\infty} - \int_0^{\infty} \frac{e^{-\lambda t}}{-\lambda} (2t) dt \right] \\ &= \lambda \left[0 + \frac{2}{\lambda^2} \int_0^{\infty} t \lambda e^{-\lambda t} dt \right] \end{aligned}$$

But the integral term in the above expression is $E(T)$ which is equal to $1/\lambda$, substituting the same,

$$E(T^2) = \lambda \left[0 + \frac{2}{\lambda^2} \times \frac{1}{\lambda} \right] = \frac{2}{\lambda^2}$$

Now variance is

$$\text{Variance} = \frac{2}{\lambda^2} - \left(\frac{1}{\lambda} \right)^2 = \frac{1}{\lambda^2} \quad (2.35)$$

Example 4 The failure time (T) of an electronic circuit board follows exponentially distribution with failure rate $\lambda = 10^{-4}/\text{h}$. What is the probability that (i) it will fail before 1000 h (ii) it will survive at least 10,000 h (iii) it will fail between 1000 and 10,000 h. Determine the (iv) mean time to failure and (v) median time failure also.

Solution:

- (i) $P(T < 1000) = F(T = 1000)$

For exponential distribution $F(T) = 1 - e^{-\lambda t}$ and substituting $\lambda = 10^{-4}/\text{h}$

$$P(T < 1000) = 1 - e^{-\lambda t} = 0.09516$$

- (ii) $P(T > 10,000) = R(T = 10,000)$

For exponential distribution $R(T) = e^{-\lambda t}$ and substituting $\lambda = 10^{-4}/\text{h}$

$$P(T > 10,000) = e^{-\lambda t} = 0.3678$$

- (iii) $P(1000 < T < 10,000) = F(10,000) - F(1000) = [1 - R(10,000)] - F(1000)$

From (i), we have $F(1000) = 0.09516$ and from (ii) we have $R(10,000) = 0.3678$,

$$P(1000 < T < 10,000) = [1 - 0.3678] - 0.09516 = 0.537$$

- (iv) Mean time to failure $= 1/\lambda = 1/10^{-4} = 10,000$ h

- (v) *Median time to failure* denote the point where 50 % failures have already occurred, mathematically it is

$$R(T) = 0.5$$

$$e^{-\lambda t} = 0.5$$

Applying logarithm on both sides and solving for t ,

$$t = \frac{-1}{\lambda} \ln(0.5) = 6931.47 \text{ h.}$$

2.4.2.2 Normal Distribution

The normal distribution is the most important and widely used distribution in the entire field of statistics and probability. It is also known as Gaussian distribution and it is the very first distribution introduced in 1733. The normal distribution often occurs in practical applications because the sum of large number of statistically

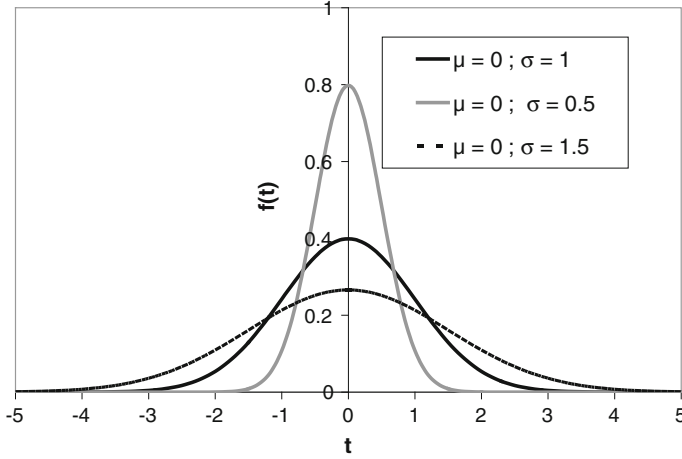


Fig. 2.15 Normal probability density functions

independent random variables converges to a normal distribution (known as central limit theorem). Normal distribution can be used to represent wear-out region of bath-tub curve where fatigue and aging can be modeled. It is also used in stress-strength interference models in reliability studies. The PDF of normal distributions is

$$f(t) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{t-\mu}{\sigma}\right)^2}, \quad -\infty \leq t \leq \infty \quad (2.36)$$

where μ and σ are parameter of the distribution. The distribution is bell shaped and symmetrical about its mean with the spread of distribution determined by σ . It is shown in Fig. 2.15.

The normal distribution is not a true reliability distribution since the random variable ranges from $-\infty$ to $+\infty$. But if the mean μ is positive and is larger than σ by several folds, the probability that random variable T takes negative values can be negligible and the normal can therefore be a reasonable approximation to a failure process.

The normal reliability function and CDF are

$$R(t) = \int_t^{\infty} \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{t-\mu}{\sigma}\right)^2} dt, \quad (2.37)$$

$$F(t) = \int_{-\infty}^t \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{t-\mu}{\sigma}\right)^2} dt \quad (2.38)$$

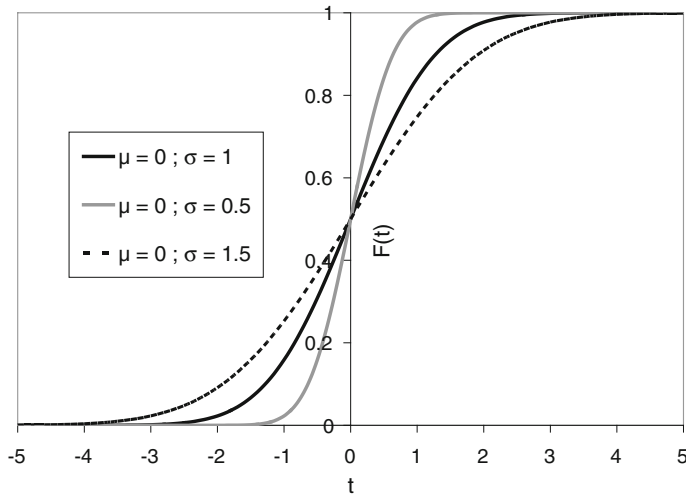


Fig. 2.16 Normal cumulative distribution functions

As there is no closed form solution to these integrals, the reliability and CDF are often expressed as a function of standard normal distribution ($\mu = 0$ and $\sigma = 1$) (Fig. 2.16). Transformation to the standard normal distribution is achieved with the expression

$$z = \frac{t - \mu}{\sigma},$$

The CDF of z is given by

$$\phi(z) = \int_{-\infty}^z \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}} dz \quad (2.39)$$

Table A.1 (see appendix) provides cumulative probability of the standard normal distribution. This can be used to find cumulative probability of any normal distribution. However, these tables are becoming unnecessary, as electronic spread sheets for example Microsoft Excel, have built in statistic functions.

The hazard function can expressed as

$$h(t) = \frac{f(t)}{R(t)} = \frac{f(t)}{1 - \Phi(z)} \quad (2.40)$$

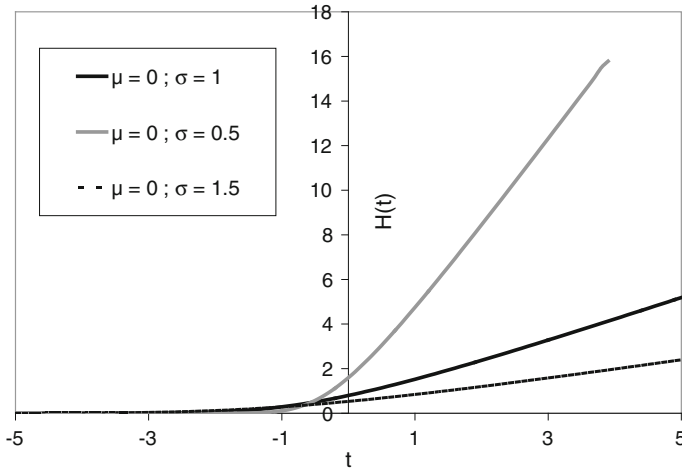


Fig. 2.17 Normal hazard rate functions

Hazard function is an increasing function as shown in Fig. 2.17. This feature makes it suitable to model aging components.

Example 5 Failure times are recorded from the life testing of an engineering component as 850, 890, 921, 955, 980, 1025, 1036, 1047, 1065, and 1120. Assuming a normal distribution, calculate the instantaneous failure rate at 1000 h?

Solution: Given data, $n = 10$, $N = 1000$; using the calculations from Table 2.5,

$$\text{Mean} = \bar{x} = \frac{\sum xi}{n} = \frac{9889}{10} = 988.9$$

Now, the sample S.D. is (σ)

$$\sigma = \sqrt{\frac{n \sum_{i=1}^n xi^2 - (\sum_{i=1}^n xi)^2}{n(n-1)}} = 84.8455$$

The instantaneous failure rate is given by the hazard function, and is established by

$$h(t) = \frac{f(t)}{R(t)} = \frac{f(1000)}{R(1000)} = \frac{\phi(z)}{1 - \Phi(z)} = \frac{0.0046619}{1 - 0.552} = 0.0104$$

Table 2.5 Calculations

xi	xi ²
850	722,500
890	792,100
921	848,241
955	912,025
980	960,400
1025	1,050,625
1036	1,073,296
1047	1,096,209
1065	1,134,225
1120	1,254,400
Σ xi = 9889	Σ xi ² = 9, 844, 021

2.4.2.3 Lognormal Distribution

A continuous positive random variable T is lognormal distribution if its natural logarithm is normally distributed. The lognormal distribution can be used to model the cycles to failure for metals, the life of transistors and bearings and modeling repair times. It appears often in accelerated life testing as well as when a large number of statistically independent random variables are multiplied. The lognormal PDF is

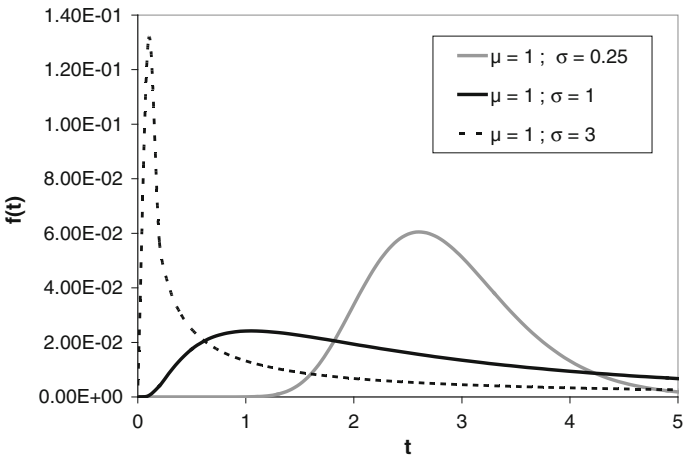


Fig. 2.18 Lognormal probability density functions

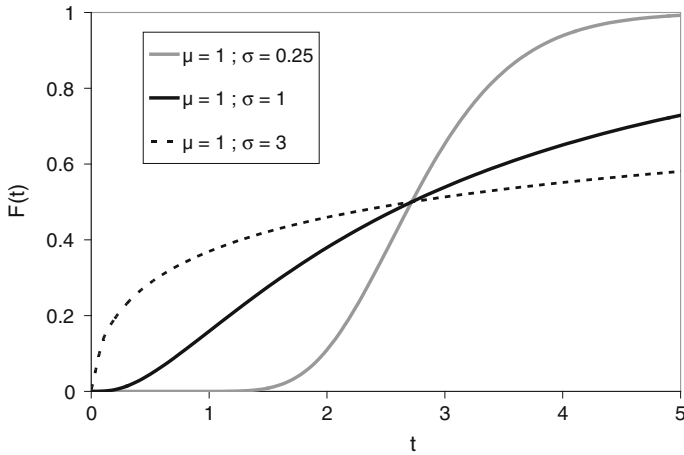


Fig. 2.19 Lognormal cumulative distribution functions

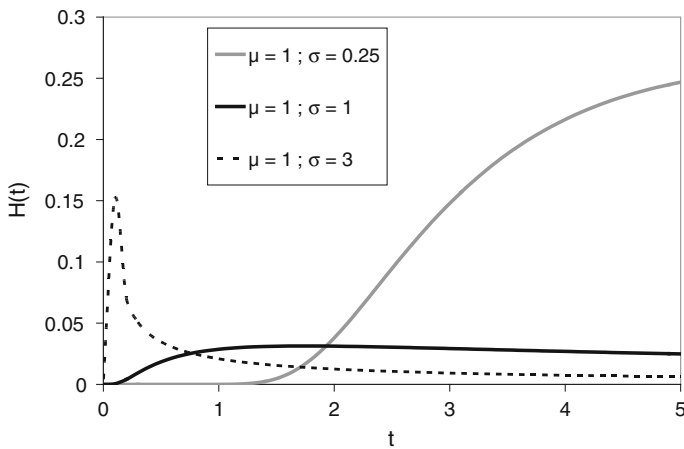


Fig. 2.20 Lognormal hazard functions

$$f(t) = \frac{1}{\sigma t \sqrt{2\pi}} e^{-\frac{1}{2} \left(\frac{\ln t - \mu}{\sigma} \right)^2}, \quad t > 0 \quad (2.41)$$

where μ and σ are known as the location parameter and shape parameters respectively. The shape of distribution changes with different values of σ as shown in Fig. 2.18.

The lognormal reliability function and CDF are

$$R(t) = 1 - \Phi \left[\frac{\ln t - \mu}{\sigma} \right] \quad (2.42)$$

$$F(t) = \Phi \left[\frac{\ln t - \mu}{\sigma} \right] \quad (2.43)$$

Lognormal failure distribution functions and lognormal hazard functions are shown in Figs. 2.19 and 2.20.

The mean of lognormal distribution is

$$E(t) = e^{\mu + \frac{\sigma^2}{2}} \quad (2.44)$$

The variance of lognormal distribution is

$$V(t) = e^{(2\mu + \sigma^2)}(e^{\sigma^2} - 1) \quad (2.45)$$

Example 6 Determine the mean and variance of time to failure for a system having lognormally distributed failure time with $\mu = 5$ years. And $\sigma = 0.8$.

Solution: The mean of lognormal distribution is,

$$E(t) = e^{\left(\mu + \frac{\sigma^2}{2}\right)}$$

$$E(t) = e^{\left(5 + \frac{0.8^2}{2}\right)} = 204.3839$$

The variance of lognormal distribution is,

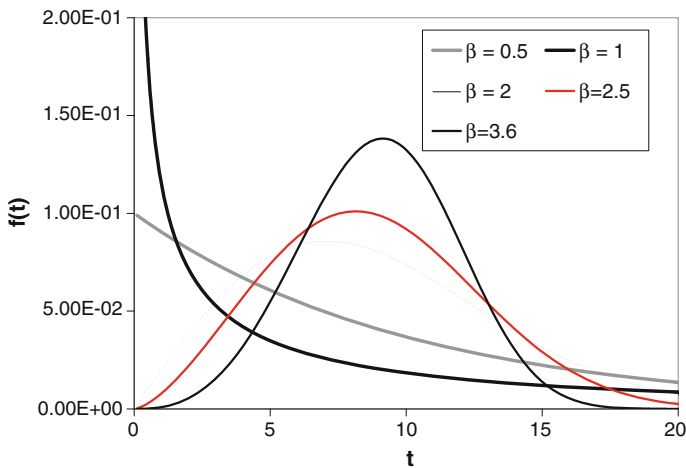


Fig. 2.21 Weibull PDF

Table 2.6 Distributions with different values of β

β	Remarks
1	Identical to exponential
2	Identical to Rayleigh
2.5	Approximates lognormal
3.6	Approximates normal

$$\begin{aligned} V(t) &= e^{(2\mu+\sigma^2)} \times (e^{\sigma^2} - 1) \\ V(t) &= e^{(10+(0.8)^2)} \times (e^{0.8^2} - 1) \\ V(t) &= 37,448.49 \end{aligned}$$

2.4.2.4 Weibull Distribution

Weibull distribution was introduced in 1933 by Rosin and Rammler [3]. Weibull distribution has wide range of applications in reliability calculation due to its flexibility in modeling different distribution shapes. It can be used to model time to failure of lamps, relays, capacitors, germanium transistors, ball bearings, automobile tyres and certain motors. In addition to being the most useful distribution function in reliability analysis, it is also useful in classifying failure types, trouble shooting, scheduling preventive maintenance and inspection activities. The Weibull PDF is

$$f(t) = \frac{\beta}{\alpha} \left(\frac{t}{\alpha}\right)^{\beta-1} e^{-\left(\frac{t}{\alpha}\right)^{\beta}}, \quad t > 0 \tag{2.46}$$

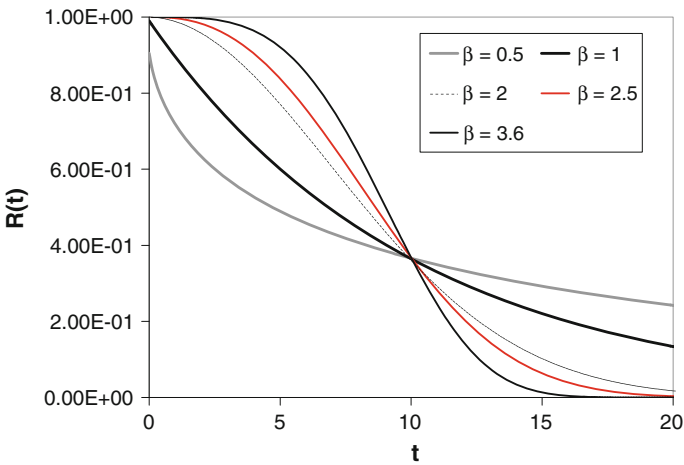


Fig. 2.22 Weibull reliability functions

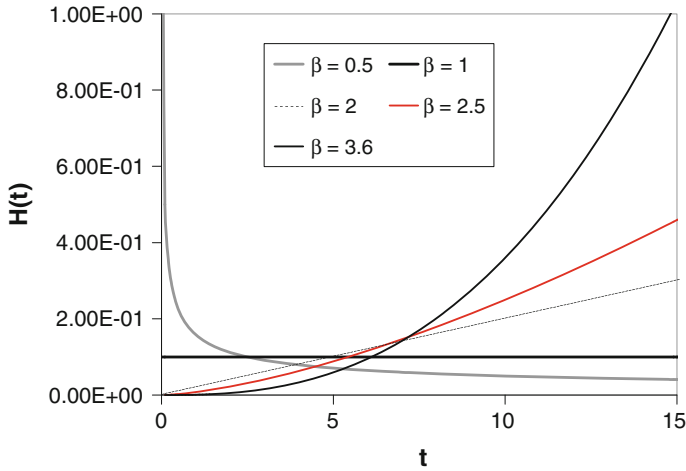


Fig. 2.23 Weibull hazard functions

where α and β are known as scale parameter (or characteristic life) and shape parameter respectively. An important property of Weibull distribution is as β increases, mean of the distribution approaches α and variance approaches zero. Its effect on the shape of distribution can be seen in Fig. 2.21 with different values of β ($\alpha = 10$ is assumed in all the cases).

It is interesting to see from Fig. 2.21, all are equal to or approximately matching with several other distributions. Due to this flexibility, Weibull distribution provides a good model for much of the failure data found in practice. Table 2.6 summarizes this behavior.

Weibull reliability and CDF functions are

$$R(t) = e^{-\left(\frac{t}{\alpha}\right)^\beta} \quad (2.47)$$

$$F(t) = 1.0 - e^{-\left(\frac{t}{\alpha}\right)^\beta} \quad (2.48)$$

Reliability functions with different values of β are shown in Fig. 2.22.

The Weibull hazard function is

$$H(t) = \frac{\beta t^{\beta-1}}{\alpha^\beta} \quad (2.49)$$

The effects of β on the hazard function are demonstrated in Fig. 2.23. All three regions of bath-tub curve can be represented by varying β value.

- $\beta < 1$ results in decreasing failure rate (burn-in period)
- $\beta = 1$ results in constant failure rate (useful life period)
- $\beta > 1$ results in increasing failure rate (Wear-out period)

The mean value of Weibull distribution can be derived as follows:

$$Mean = \int_0^{\infty} t f(t) dt = \int_0^{\infty} t \cdot \left(\frac{\beta}{\alpha}\right) \left(\frac{t}{\alpha}\right)^{\beta-1} e^{-\left(\frac{t}{\alpha}\right)^{\beta}} dt$$

$$\text{Let } x = \left(\frac{t}{\alpha}\right)^{\beta},$$

$$dx = \left(\frac{\beta}{\alpha}\right) \left(\frac{t}{\alpha}\right)^{\beta-1} dt$$

$$\text{Now mean} = \int_0^{\infty} t e^{-y} dy$$

$$\text{Since } t = \alpha x^{\frac{1}{\beta}}$$

$$Mean = \alpha \int_0^{\infty} (x)^{\frac{1}{\beta}} e^{-x} dx = \alpha \Gamma\left(1 + \frac{1}{\beta}\right). \quad (2.50)$$

where $\Gamma(x)$ is known as gamma function.

$$\Gamma(x) = \int_0^{\infty} y^{x-1} \cdot e^{-y} dy$$

Similarly variance can be derived as

$$\sigma^2 = \alpha^2 \left[\Gamma\left(1 + \frac{2}{\beta}\right) - \Gamma^2\left(1 + \frac{1}{\beta}\right) \right] \quad (2.51)$$

Example 7 The failure time of a component follows Weibull distribution with shape parameter $\beta = 1.5$ and scale parameter = 10,000 h. When should the component be replaced if the minimum recurred reliability for the component is 0.95?

Solution: Substituting into the Weibull reliability function gives,

$$R(t) = e^{-\left(\frac{t}{\alpha}\right)^{\beta}}$$

$$0.95 = e^{-\left(\frac{t}{10,000}\right)^{1.5}} \Rightarrow \frac{1}{0.95} = e^{\left(\frac{t}{10,000}\right)^{1.5}}$$

Taking natural logarithm on both sides

$$\ln \frac{1}{0.95} = \left(\frac{t}{10,000} \right)^{1.5}$$

Taking log on both sides,

$$\begin{aligned} \log 0.051293 &= 1.5 \log \frac{t}{10,000} \Rightarrow \frac{-1.2899}{1.5} = \log \frac{t}{10,000} \\ \Rightarrow -0.85996 &= \log t - \log 10,000 \Rightarrow \log 10,000 - 0.85996 = \log t \\ \Rightarrow t &= 1380.38 \text{ h} \end{aligned}$$

2.4.2.5 Gamma Distribution

As the name suggests, gamma distribution derives its name from the well known gamma function. It is similar to Weibull distribution where by varying the parameter of the distribution wide range of other distribution can be derived. The gamma distribution is often used to model life time of systems. If an event takes place after 'n' exponentially distributed events take place sequentially, the resulting random variable follows a gamma distribution. Examples of its application include the time to failure for a system consisting of n independent components, with n – 1 components being stand by comp; time between maintenance actions for a system that requires maintenance after a fixed number of uses; time to failure of system which fails after n shocks. The gamma PDF is

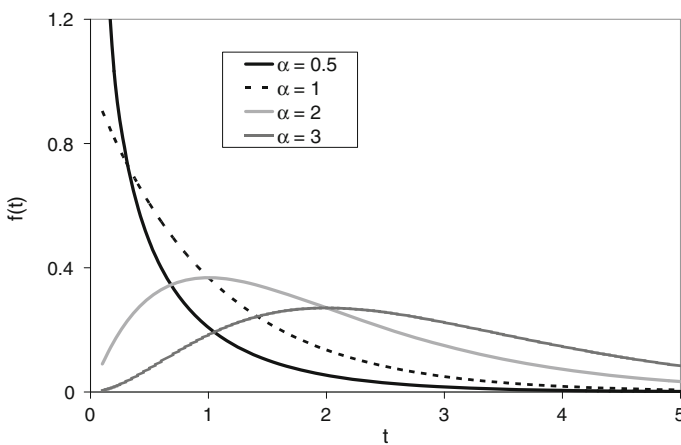


Fig. 2.24 Gamma probability density functions

Table 2.7 Distribution with different values of α

α	Distribution
$\alpha = 1$	Exponential distribution
$\alpha = \text{integer}$	Erlangian distribution
$\alpha = 2$	Chi square distribution
$\alpha > 2$	Normal distribution

$$f(t) = \Gamma(t; \alpha, \beta) = \frac{\beta^\alpha}{\Gamma(\alpha)} t^{\alpha-1} e^{-\beta t}, t \geq 0$$

$$\text{where } \Gamma(\alpha) = \int_0^\infty x^{\alpha-1} e^{-x} dx. \quad (2.52)$$

where α and β are parameters of distribution. The PDF with parameter $\beta = 1$ known as standardized gamma density function. By changing the parameter α , different well known distributions can be generated as shown in Fig. 2.24 and Table 2.7.

The CDF of random variable T having gamma distribution with parameter α and β is given by,

$$F(t) = P(T < t) = \int_0^t \frac{\beta^\alpha}{\Gamma(\alpha)} t^{\alpha-1} e^{-\beta t} dt \quad (2.53)$$

The gamma CDF in general does not have closed form solution. However, tables are available given the values of CDF having standard gamma distribution function.

The mean of gamma distribution is

$$E(T) = \frac{\alpha}{\beta} \quad (2.54)$$

The variable of gamma distribution is

$$V(T) = \frac{\alpha}{\beta^2} \quad (2.55)$$

For integer values of α , the gamma PDF is known as Erlangian probability density function.

2.4.2.6 Erlangian Distribution

Erlangian distribution is special case of gamma distribution where α is an integer. In this case PDF is express as,

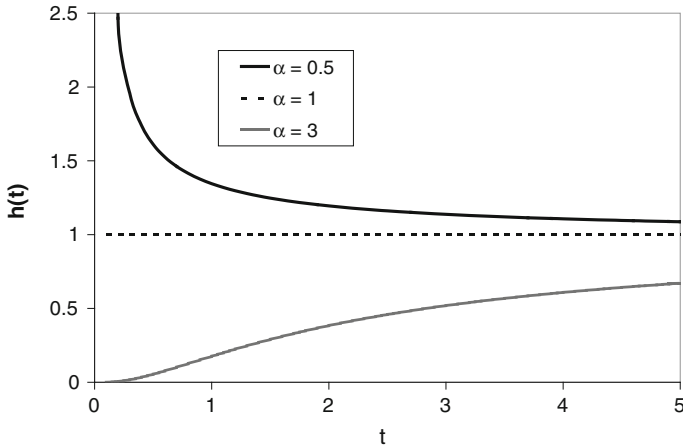


Fig. 2.25 Erlangian hazard functions

$$f(t) = \frac{t^{\alpha-1}}{\beta^\alpha (\alpha-1)!} e^{-\left(\frac{t}{\beta}\right)} \quad (2.56)$$

The Erlangian reliability function is

$$R(t) = \sum_{k=0}^{\alpha-1} \frac{\left(\frac{t}{\beta}\right)^k e^{-\left(\frac{t}{\beta}\right)}}{k!} \quad (2.57)$$

The hazard function is

$$h(t) = \frac{t^{\alpha-1}}{\beta^\alpha \Gamma(\alpha) \sum_{k=0}^{\alpha-1} \frac{\left(\frac{t}{\beta}\right)^k}{k!}} \quad (2.58)$$

By changing the value of α , all three phases of bath-tub curves can be selected (Fig. 2.25). If $\alpha < 1$, failure rate is decreasing, $\alpha = 1$, failure rate is constant and $\alpha > 1$, failure rate is increasing.

2.4.2.7 Chi-Square Distribution

A special case of the gamma distribution with $\alpha = 2$ and $\beta = 2/\nu$, a chi-square (χ^2) distribution is used to determinant of goodness of fit and confidence limits.

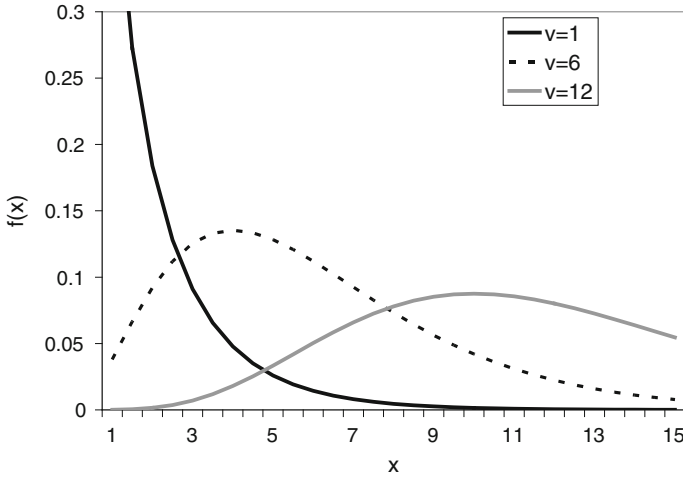


Fig. 2.26 PDF of Chi-Square

The chi-square probability density function is

$$\chi^2(x, v) = f(x) = \frac{1}{2^{v/2}\Gamma(v/2)} x^{(v/2-1)} e^{-x/2}, \quad x > 0 \quad (2.59)$$

The shape of chi-square distribution is shown in Fig. 2.26.

The mean of chi-square distribution is $E(x) = v$.

The variance of chi-square distribution is $V(x) = 2v$.

If $x_1, x_2, \dots, x_n$ are independent, standard normally distributed variables, then the sum of squares of random variable, i.e., $(X_1^2 + X_2^2 + \dots + X_v^2)$ is chi-square distribution with v degree of freedom.

It is interesting to note that the sum of two or more independent chi-square variables is also a chi-square variable with degree-of-freedom equal to the sum of degree-of-freedom for the individual variable. As v become large, the chi-square distribution approaches normal with mean v and variance $2v$.

2.4.2.8 F-Distribution

If χ_1^2 and χ_2^2 are independent chi-square random variable with v_1 and v_2 degrees of freedom, then the random variable F defined by

$$F = \frac{\chi_1^2/v_1}{\chi_2^2/v_2} \quad (2.60)$$

is said to have an F-distribution with v_1 and v_2 degrees of freedom.

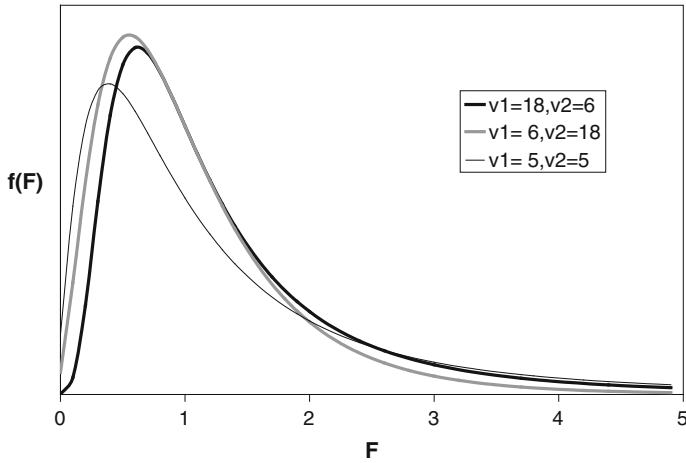


Fig. 2.27 F PDFs with different v_1 and v_2

The PDF of random variable F is given by

$$f(F) = \left[\frac{\Gamma\left(\frac{v_1+v_2}{2}\right) \left(\frac{v_1}{v_2}\right)^{\frac{v_1}{2}}}{\Gamma\left(\frac{v_1}{2}\right) \Gamma\left(\frac{v_2}{2}\right)} \right] \left[\frac{F^{\frac{v_1}{2}-1}}{\left(1 + v_1 \frac{F}{v_2}\right)^{\left(\frac{v_1+v_2}{2}\right)}} \right], \quad F > 0 \quad (2.61)$$

Figure 2.27 shows F PDF with different v_1 and v_2 .

The values of F-distribution are available from tables. If $f_\alpha(v_1, v_2)$ represent area under the F pdf, with degree of freedom v_1 and v_2 , to the right of α , then

$$F_{1-\alpha}(v_1, v_2) = \frac{1}{F_\alpha(v_2, v_1)} \quad (2.62)$$

It is interesting to observe that if s_1^2 and s_2^2 are the variance of independent random samples of size n_1 and n_2 drawn from normal population with variance of σ_1^2 and σ_2^2 respectively then the statistic

$$F = \frac{s_1/\sigma_1^2}{s_2/\sigma_2^2} = \frac{\sigma_2^2 \cdot s_1^2}{\sigma_1^2 \cdot s_2^2} \quad (2.63)$$

has an F distribution with $v_1 = n_1 - 1$ and $v_2 = n_2 - 1$ degree of freedom.

2.4.2.9 t-Distribution

If Z is normally distributed random variable and the independent random variable χ^2 follows a chi square distribution with v degree of freedom then the random variable t defined by

$$t = \frac{z}{\sqrt{\chi^2/v}} \quad (2.64)$$

is said to be have t-distribution with v degree of freedom.

PDF of t is given by

$$f(t) = \frac{\Gamma\left(\frac{v+1}{2}\right)}{\Gamma(v/2)\sqrt{\Pi v}} \left[1 + \frac{t^2}{v}\right]^{-\frac{(v+1)}{2}}, \quad -\infty < t < \infty. \quad (2.65)$$

Table 2.8 Summary of application areas

Distribution	Areas of application in reliability studies
Poisson distribution	To model occurrence rates such as failures per hour or defects per item (defects per computer chip or defects per automobile)
Binomial distribution	To model K out of M or voting redundancy such as triple modular redundancies in control and instrumentation
Exponential distribution	To model useful life of many items Life distribution of complex non-repairable systems
Weibull distribution	$\beta > 1$ often occurs in applications as failure time of components subjected to wear out and/or fatigue (lamps, relays, mechanical components) Scheduling inspection and preventive maintenance activities
Lognormal distribution	To model the cycles to failure for metals, the life of transistors, the life of bearings. Size distribution of pipe breaks To model repair time Prior parameter distribution in Bayesian analysis
Normal distribution	Modeling buildup of tolerances Load-resistance analysis (stress-strength interference) Life distribution of high stress components
Gamma distribution	To model time to failure of system with standby units To model time between maintenance actions Prior parameter distribution in Bayesian analysis
Chi-square distribution	Count the number of failures in an interval Applications involving goodness of fit and confidence limits
F distribution	To make inferences about variances and to construct confidence limits
t distribution	To draw inferences concerning means and to construct confidence intervals for means when the variances is unknown

Like the standard normal density, the t-density is symmetrical about zero. In addition, as v become larger, it becomes more and more like standard normal density.

Further,

$$E(t) = 0$$

$$\text{and } v(t) = v/(v - 2) \quad \text{for } v > 2.$$

2.4.3 Summary

The summary of applications of the various distributions is described in the Table 2.8.

2.5 Failure Data Analysis

The credibility of any reliability/safety studies depend upon the quality of the data used. This section deals with the treatment of failure data and subsequent usage in reliability/safety studies. The derivation of system reliability models and various reliability measures is an application of probability theory, where as the analysis of failure data is primarily an application of statistics.

The objective of failure data analysis is to obtain reliability and hazard rate functions. This is achieved by two general approaches. The first is deriving empirical reliability and hazard functions directly from failure data. These methods are known as non parametric methods or empirical methods. The second approach is to identify an approximate theoretical distribution, estimate the parameter(s) of distribution, and perform a goodness of fit test. This approach is known as parametric method. Both the methods are explained in this section.

2.5.1 Nonparametric Methods

In this method empirical reliability distributions are directly derived from the failure data. The sources of failure data are generally from (1) Operational or field experience and/or (2) Failures generated from reliability testing. Nonparametric method is useful for preliminary data analysis to select appropriate theoretical distribution. This method is also finds application when no parametric distribution adequately fits the failure data.

Consider life tests on a certain unit under exactly same environment conditions with N number of units ensuring that failures of the individual units are independent and do not affect each other. At some predetermined intervals of time, the number of failed units is observed. It is assumed that test is carried out till all the units have failed. Now let us analyze the information collected through this test.

From the classical definition of probability, the probability of occurrence of an event A can be expressed as follows

$$P(A) = \frac{n_s}{N} = \frac{n_s}{n_s + n_f} \quad (2.66)$$

where

n_s is the number of favorable outcomes

n_f is number of unfavorable outcomes

N is total number of trials = $n_s + n_f$

When N number of units are tested, let us assume that $n_s(t)$ units survive the life test after time t and that $n_f(t)$ units have failed over the time t . Using the above equation, the reliability of such a unit can be expressed as:

$$R(t) = \frac{n_s(t)}{N} = \frac{n_s(t)}{n_s(t) + n_f(t)} \quad (2.67)$$

This definition of reliability assumes that the test is conducted over a large number of identical units.

The unreliability $Q(t)$ of unit is the probability of failure over time t , equivalent to Cumulative Distribution Function (CDF) and is given by $F(t)$,

$$Q(t) \equiv F(t) = \frac{n_f(t)}{N} \quad (2.68)$$

We know that the derivative of the CDF of a continuous random variable gives the PDF. In reliability studies, failure density function $f(t)$ associated with failure time of a unit can be defined as follows:

$$f(t) \equiv \frac{dF(t)}{dt} = \frac{dQ(t)}{dt} = \frac{1}{N} \frac{dn_f}{dt} = \frac{1}{N} \lim_{\Delta t \rightarrow 0} \left\{ \frac{n_f(t + \Delta t) - n_f(t)}{\Delta t} \right\} \quad (2.69)$$

Hazard rate can be derived from Eq. 2.21 by substituting $f(t)$ and $R(t)$ as expressed below

$$h(t) = \frac{1}{n_s(t)} \lim_{\Delta t \rightarrow 0} \left\{ \frac{n_f(t + \Delta t) - n_f(t)}{\Delta t} \right\} \quad (2.70)$$

Equations 2.67, 2.69 and 2.70 can be used for computing reliability, failure density and hazard functions from the given failure data.

The preliminary information on the underlying failure model can be obtained if we plot the failure density, hazard rate and reliability functions against time. We can define piece wise continuous functions for these three characteristics by selecting some small time interval Δt . This discretization eventually in the limiting conditions i.e., $\Delta t \rightarrow 0$ or when the data is large would approach to the continuous function analysis. The number of interval can be decided based on the range of data and accuracy desired. But higher is the number of intervals, better would be the accuracy of results. However the computational effort increases considerably if we choose a large number of intervals. However, there exist an optimum number of intervals given by Sturges [4], which can be used to analyze the data. If n is the optimum number of intervals and N is the total number of failures, then

$$n = 1 + 3.3 \log_{10}(N) \quad (2.71)$$

Example 8 To ensure proper illumination in control rooms, higher reliability of electric-lamps is necessary. Let us consider that the failure times (in hours) of a population of 30 electric-lamps from a control room are given in the following Table 2.9. Calculate failure density, reliability and hazard functions?

Solution:

The optimum number of intervals as per Sturge's formula (Eq. 2.71) with $N = 30$ is

$$n = 1 + 3.3 \log(30) = 5.87$$

Table 2.9 Failure data

Lamp	Failure time	Lamp	Failure time	Lamp	Failure time
1	5495.05	11	3511.42	21	4037.11
2	8817.71	12	6893.81	22	933.79
3	539.66	13	1853.83	23	1485.66
4	2253.02	14	3458.4	24	4158.11
5	18,887	15	7710.78	25	6513.43
6	2435.62	16	324.61	26	8367.92
7	99.33	17	866.69	27	1912.24
8	3716.24	18	6311.47	28	13,576.97
9	12,155.56	19	3095.62	29	1843.38
10	552.75	20	927.41	30	4653.99

Table 2.10 Data in ascending order

Bulb	Failure time	Bulb	Failure time	Bulb	Failure time
1	99.33	11	1912.24	21	5495.05
2	324.61	12	2253.02	22	6311.47
3	539.66	13	2435.62	23	6513.43
4	552.75	14	3095.62	24	6893.81
5	866.69	15	3458.4	25	7710.78
6	927.41	16	3511.42	26	8367.92
7	933.79	17	3716.24	27	8817.71
8	1485.66	18	4037.11	28	12,155.56
9	1843.38	19	4158.11	29	13,576.97
10	1853.83	20	4653.99	30	18,887

In order to group the failure times under various intervals, the data is arranged in increasing order. Table 2.10 is the data with ascending order of failure times. The minimum and maximum of failure time is 99.33 and 18,887 respectively.

$$Interval\ size = \Delta t_i = \frac{18,887 - 99.33}{6} = 3131.27 \approx 3150$$

We can now develop a table showing the intervals and corresponding values of $R(t)$, $F(t)$, $f(t)$ and $h(t)$ respectively. The following notation is used. The summary of calculations is shown in Table 2.11.

$n_s(t_i)$ number of survivors at the beginning of the interval

$n_f(t_i)$ number of failures during i th interval

The plots of $f(t)$ and $h(t)$ are shown in Figs. 2.28 and 2.29 where as the plots of $R(t)$ and $F(t)$ are given in Fig. 2.30.

Table 2.11 Calculations

Interval	$n_s(t_i)$	$n_f(t_i)$	$R(t_i)$	$F(t_i)$	$f(t_i) = \frac{n_f(t_i)}{N\Delta t_i}$	$h(t_i) = \frac{n_f(t_i)}{n_s(t_i)\Delta t_i}$
0–3150	30	14	1	0	1.48e-4	1.48e-4
3151–6300	16	7	0.53	0.47	7.4e-5	1.38e-4
6301–9450	9	6	0.3	0.7	6.35e-5	2.11e-4
9451–12,600	3	1	0.1	0.9	1.06e-5	1.05e-4
12,601–15,750	2	1	0.066	0.934	1.06e-5	1.58e-4
15,751–18,900	1	1	0.033	0.967	1.06e-5	3.17e-4

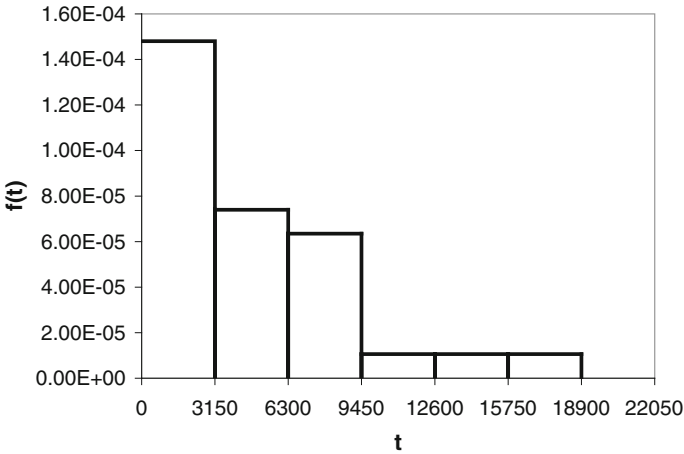


Fig. 2.28 Failure density function

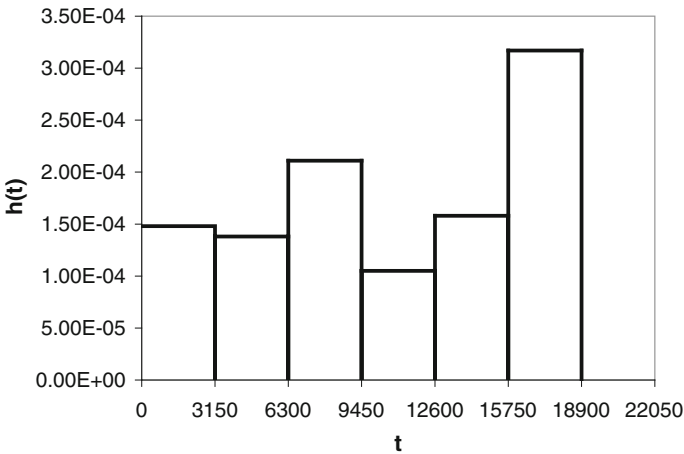


Fig. 2.29 Hazard rate function

2.5.2 Parametric Methods

Preceding section discussed methods for deriving empirical distributions directly from failure data. The second, and usually preferred, method is to fit a theoretical distribution, such as the exponential, Weibull, or normal distributions. As theoretical distributions are characterized with parameters, these methods are known as parametric method. Nonparametric methods have certain practical limitations compared with parametric methods.

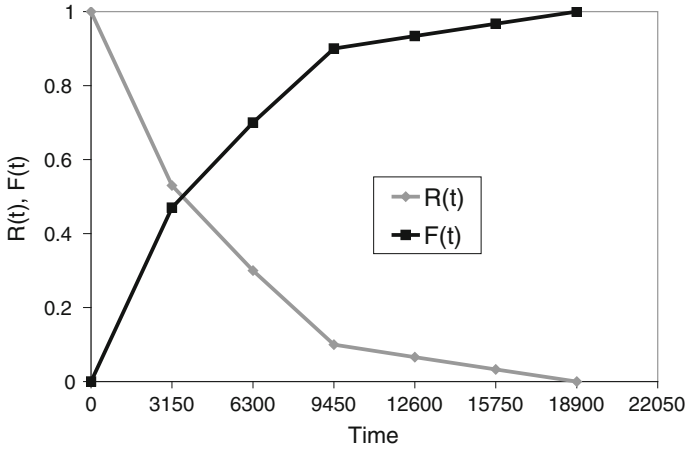


Fig. 2.30 Reliability function/CDF

1. As nonparametric methods are based on sample data, information beyond the range of data cannot be provided. Extrapolation beyond the censored data is possible with a theoretical distribution. This is significant in reliability/safety studies as the tails of the distribution attract more attention.
2. The main concern is determining the probabilistic nature of the underlying failure process. The available failure data may be simple a subset of the population of failure times. Establishing the distribution the sample came from and not sample itself is the focus.
3. The failure process is often a result of some physical phenomena that can be associated with a particular distribution.
4. Handling a theoretical model is easy in performing complex analysis.

In parametric approach, fitting of a theoretical distribution, consists of the following three steps:

1. Identifying candidate distribution
2. Estimating the parameters of distributions
3. Performing goodness-of-fit test

All these steps are explained in the following sections.

2.5.2.1 Identifying Candidate Distributions

In the earlier section on nonparametric methods, we have seen how one can obtain empirical distributions or histograms from the basic failure data. This exercise helps one to guess a failure distribution that can be possibly employed to model the failure data. But nothing has been said about an appropriate choice of the distribution. Probability plots provide a method of evaluating the fit of a set of data to a distribution.

A probability plot is a graph in which the scales have been changed in such a manner that the CDF associated with a given family of distributions, when represented graphically on that plot, becomes a straight line. Since straight lines are easily identifiable, a probability plot provided a better visual test of a distribution than comparison of a histogram with a PDF. Probability plots provide a quick method to analyze, interpret and estimate the parameters associated with a model. Probability plots may also be used when the sample size is too small to construct histograms and may be used with incomplete data.

The approach to probability plots is to fit a linear regression line of the form mentioned below to a set of transformed data:

$$y = mx + c \quad (2.72)$$

The nature of transform will depend on the distribution under consideration. If the data of failure times fit the assumed distribution, the transformed data will graph as a straight line.

Consider exponential distribution whose CDF is $F(t) = 1 - e^{-\lambda t}$, rearranging $1 - F(t) = e^{-\lambda t}$, taking the natural logarithm of both sides,

$$\begin{aligned} \ln(1 - F(t)) &= \ln(e^{-\lambda t}) \\ -\ln(1 - F(t)) &= \lambda t \\ \ln\left(\frac{1}{1 - F(t)}\right) &= \lambda t \end{aligned}$$

Comparing it with Eq. 2.72: $y = mx + c$, we have

$$\begin{aligned} y &= \ln\left(\frac{1}{1 - F(t)}\right) \\ m &= \lambda; x = t; c = 0; \end{aligned}$$

Now if y is plotted on the ordinate, the plot would be a straight line with a slope of λ .

The failure data is generally available in terms of the failure times of n items that have failed during a test conducted on the original population of N items. Since $F(t)$ is not available, we can make use of $E[F(t_i)]$

$$E[F(t_i)] = \sum_{i=1}^n \frac{i}{N+1} \quad (2.73)$$

Example 9 Table 2.12 gives chronological sequence of the grid supply outages at a process plant. Using probability plotting method, identify the possible distributions.

Table 2.12 Class IV power failure occurrence time since 01.01.1998

Failure number	Date/time	Time to failure (in days)	Time between failure (in days)
1	11.04.1998/14:35	101	101
2	17.06.1998/12:30	168	67
3	24.07.1998/09:19	205	37
4	13.08.1999/10:30	590	385
5	27.08.1999	604	14
6	21.11.1999	721	117
7	02.01.2000	763	42
8	01.05.2000/15:38	882	119
9	27.10.2000/05:56	1061	179
10	14.05.2001	1251	190
11	03.07.2001/09:45	1301	50
12	12.07.2002/18:50	1674	374
13	09.05.2003/08:43	1976	301
14	28.12.2005	2940	964
15	02.05.2006/11:02	3065	125
16	17.05.2007/11:10	3445	380
17	02.06.2007/16:30	3461	16

Table 2.13 Time between failure (TBF) values for outage of Class IV (for Weibull plotting)

I	Failure number	TBF (in days) (t)	$F(t) = (i - 0.3)/(n + 0.4)$	$y = \ln(\ln(1/R(t)))$	$x = \ln(t)$
1	5	14	0.04023	-3.19268	2.639057
2	17	16	0.097701	-2.27488	2.772589
3	3	37	0.155172	-1.78009	3.610918
4	7	42	0.212644	-1.43098	3.73767
5	11	50	0.270115	-1.1556	3.912023
6	2	67	0.327586	-0.92412	4.204693
7	1	101	0.385057	-0.72108	4.615121
8	6	117	0.442529	-0.53726	4.762174
9	8	119	0.5	-0.36651	4.779123
10	15	125	0.557471	-0.20426	4.828314
11	9	179	0.614943	-0.04671	5.187386
12	10	190	0.672414	0.109754	5.247024
13	13	301	0.729885	0.269193	5.70711
14	12	374	0.787356	0.437053	5.924256
15	16	380	0.844828	0.622305	5.940171
16	4	385	0.902299	0.844082	5.953243
17	14	964	0.95977	1.16725	6.871091

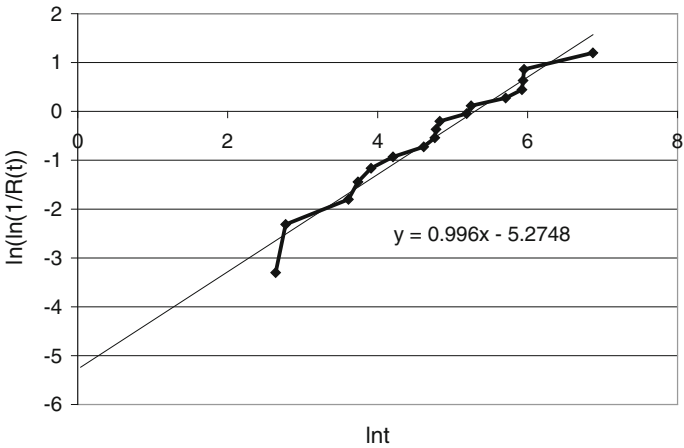


Fig. 2.31 Weibull plotting for the data

Table 2.14 Coordinates of distributions for probability plotting

Distribution	(x, y)	$y = mx + c$
Exponential $F(t) = 1 - e^{-\lambda t}$	$\left(t, \ln \left[\frac{1}{1-F(t)} \right] \right)$	$m = \lambda$ $c = 0$
Weibull $F(t) = 1 - e^{-(\frac{t}{\beta})^\alpha}$	$\left(\ln t, \ln \ln \left[\frac{1}{1-F(t)} \right] \right)$	$m = \alpha$ $c = \ln(1/\beta)$
Normal $F(t) = \Phi \left[\frac{t-\mu}{\sigma} \right]$	$\left(t, \Phi^{-1}[F(t)] \right)$	$m = \frac{1}{\sigma}$ $c = \frac{-\mu}{\sigma}$

Solution:

Table 2.13 gives the summary of calculations for x and y coordinates. The same are plotted in Fig. 2.31.

The plot is approximated to a straight line as mentioned below

$$y = 0.996x - 5.2748$$

The shape parameter $\alpha = 0.996$

Scale parameter, $\beta = e^{5.2748} = 194.4$ days

As shape parameter is close to unity, the data fits exponential distribution.

Table 2.14 summarizes (x, y) coordinates of various distributions used in probability plotting.

2.5.2.2 Estimating the Parameters of Distribution

The preceding section on probability plotting focused on the identification of distribution for a set of data. Specification of parameters for the identified distribution is the next step. The estimation of parameters of the distribution by probability plotting is not considered to be best estimates. This is especially true in certain goodness of fit tests that are based on Maximum Likelihood Estimator (MLE) for the distribution parameters. There are many criteria based on which an estimator can be computed, viz., least square estimation and MLE. MLE provides maximum flexibility and is widely used.

Maximum Likelihood Estimates

Let the failure times, $t_1, t_2, \dots, t_n$ represent observed data from a population distribution, whose PDF is $f(t|\theta_1, \dots, \theta_k)$ where θ_i is the parameter of the distribution. Then the problem is to find likelihood function given by

$$L(\theta_1 \dots \theta_k) = \prod_{i=1}^n f(t_i|\theta_1 \dots \theta_k) \quad (2.74)$$

The objective is to find the values of the estimators of $\theta_1, \dots, \theta_k$ that render the likelihood function as large as possible for given values of $t_1, t_2, \dots, t_n$. As the likelihood function is in the multiplicative form, it is to maximize $\log(L)$ instead of L but these two are identical since maximizing L is equivalent to maximizing $\log(L)$.

By taking partial derivatives of the equation with respect to $\theta_1, \dots, \theta_k$ and setting these partial equal to zero, the necessary conditions for finding MLEs can be obtained.

$$\frac{\partial \ln L(\theta_1 \dots \theta_k)}{\partial \theta_i} = 0 \quad i = 1, 2, \dots, k \quad (2.75)$$

Exponential MLE

The likelihood function for a single parameter exponential distribution whose PDF is $f(t) = \lambda e^{-\lambda t}$ is given by

$$L(t_1 \dots t_n|\lambda) = (\lambda e^{-\lambda t_1})(\lambda e^{-\lambda t_2}) \dots (\lambda e^{-\lambda t_n}) = \lambda^n e^{-\lambda \sum_{j=1}^n t_j} \quad (2.76)$$

Taking logarithm, we have

$$\ln L(t_1, t_2, \dots, t_n|\lambda) = n \ln \lambda - \lambda \sum_{j=1}^n t_j \quad (2.77)$$

Partially differentiating the Eq. 2.77 with respect to λ and equating to zero, we have

$$\hat{\lambda} = \frac{n}{\sum_{j=1}^n t_j} \quad (2.78)$$

where $\hat{\lambda}$ is the MLE of λ .

Interval Estimation

The point estimates would provide the best estimate of the parameter where as the interval estimation would offer the bounds with in which the parameter would lie. In other words, it provides the confidence interval for the parameter. A confidence interval gives a range of values among which we have a high degree of confidence that the distribution parameter is included.

Since there is always an uncertainty associated in this parameter estimation, it is essential to find upper confidence and lower confidence limit of these two parameters.

Upper and Lower Confidence of the Failure Rate

The Chi square distribution is used to find out upper and lower confidence limits of Mean Time To Failure. The Chi square equation is given as follow

$$\theta_{LC} \equiv \frac{2T}{\chi_{2r, \alpha/2}^2} \quad (2.79)$$

$$\theta_{UC} \equiv \frac{2T}{\chi_{2r, 1-\alpha/2}^2} \quad (2.80)$$

where

θ_{LC} and θ_{UC}	Lower and Upper Confidence limits of mean time to failure
r	Observed number of failures
T	Operating Time
α	Level of significance

The mean time represents the Mean Time Between Failure (MTBF) or Mean Time To Failure (MTTF). When failure model follows an exponential distribution, the failure rate can be expressed as

$$\lambda = \frac{1}{\theta}$$

Thus, the inverse of θ_{LC} and θ_{UC} will be the maximum and minimum possible value of the failure rate, i.e. the upper and lower confidence limit of the failure rate.

Upper and Lower Confidence Limit of the Demand Failure Probability:

In case of demand failure probability, F-Distribution is used to derive the upper and the lower confidence limit.

$$P_{LC} = \frac{r}{r + (D - r + 1)F_{0.95}(2D - 2r + 2, 2r)} \quad (2.81)$$

$$P_{UC} = \frac{(r + 1)F_{0.95}(2r + 2, 2D - 2r)}{D - r + (r + 1)F_{0.95}(2r + 2, 2D - 2r)} \quad (2.82)$$

where,

P_{LC} and P_{UC} Lower and Upper Confidence limits for demand failure probabilities

r number of failures

D number of demands

$F_{0.95}$ 95 % confidence limit for variables from F-distribution Table A.4.

Example 10 Estimate the point and 90 % confidence interval for the data given in the previous example on grid outage in a process plant.

Solution: Total Number of Outages: 17

Total Period: 10 year.

Mean failure rate = $17/10 = 1.7/\text{year} = 1.94 \times 10^{-4}/\text{h}$.

The representation of Lower (5 %) and Upper (95 %) limits of (Chi-square) χ^2 distribution is as follows for failure terminated tests is as follows;

$$\frac{\chi^2_{\alpha/2; 2\gamma}}{2T} \leq \lambda \leq \frac{\chi^2_{1-\alpha/2; 2\gamma}}{2T} \quad (2.83)$$

For the case under consideration

α 100 - 90 = 10 %;

n 17;

Degree of freedom $\gamma = n = 17$;

T 10 year.

$$\frac{\chi^2_{0.05; 2 \cdot 17}}{2 \cdot 10} \leq \lambda \leq \frac{\chi^2_{0.95; 2 \cdot 17}}{2 \cdot 10}$$

Obtaining the respective values from the χ^2 Table A.3, $1.077 \leq \lambda \leq 2.55$.

The mean value of grid outage frequency is 1.7/year ($1.94 \times 10^{-4}/h$) with lower and upper limit of 1.077/year ($1.23 \times 10^{-4}/h$) and 2.55/year ($2.91 \times 10^{-4}/h$) respectively.

2.5.2.3 Goodness-of-Fit Tests

The last step in the selection of a parametric distribution is to perform a statistical test for goodness of fit. Goodness-of-fit tests have the purpose to verify agreement of observed data with a postulated model. A typical example is as follows:

Given $t_1, t_2, \dots, t_n$ as n independent observations of a random variable (failure time) t , a rule is asked to test the null hypothesis

H_0 The distribution function of t is the specified distribution

H_1 The distribution function of t is not the specified distribution

The test consists of calculating a statistic based on the sample of failure times. This statistic is then compared with a critical value obtained from a table of such values. Generally, if the test statistic is less than the critical value, the null hypothesis (H_0) is accepted, otherwise the alternative hypothesis (H_1) is accepted. The critical value depends on the level of significance of the test and the sample size. The level of significance is the probability of erroneously rejecting the null hypothesis in favor of the alternative hypothesis.

A number of methods are available to test how closely a set of data fits an assumed distribution. For some distribution functions used in reliability theory, particular procedures have been developed, often with different alternative hypotheses H_1 and investigation of the corresponding test power. Among the distribution free procedures, chi-square (χ^2) is frequently used in practical applications to solve the goodness-of-fit problems.

The chi-square (χ^2) goodness-of-fit test

The χ^2 test is applicable to any assumed distribution provided that a reasonably large number of data points are available. The assumption for the χ^2 goodness-of-fit tests is that, if a sample is divided into n cells (i.e. we have v degrees of freedom where $v = n-1$), then the values within each cell would be normally distributed about the expected value, if the assumed distribution is correct, i.e., if x_i and E_i are the observed and expected values for cell i :

$$\chi^2 = \sum_{i=1}^n \frac{(x_i - E_i)^2}{E_i} \quad (2.84)$$

If we obtain a very low χ^2 (e.g. less than the 10th percentile), it suggests that the data corresponds more closely to the proposed distribution. Higher values of χ^2 cast doubt on the null hypothesis. The null hypothesis is usually rejected when the value of

χ^2 falls outside the 90th percentile. If χ^2 is below this value, there is insufficient information to reject the hypothesis that the data come from the supposed distribution.

For further reading on treatment of statistical data for reliability analysis, interested readers may refer Ebeling [5] and Misra [6].

Exercise Problems

1. A continuous random variable T is said to have an exponential distribution with parameter λ , if PDF is given by $f(t) = \lambda e^{-\lambda t}$, calculate the mean and variance of T?
2. Given the following PDF for the random variable time to failure of a circuit breaker, what is the reliability for a 1500 h operating life?

$$f(t) = \frac{\beta}{\alpha} \left(\frac{t}{\alpha}\right)^{\beta-1} e^{-\left(\frac{t}{\alpha}\right)^\beta} \text{ with } \alpha = 1200 \text{ h and } \beta = 1.5.$$

3. Given the hazard rate function $\lambda(t) = 2 \times 10^{-5}t$, determine R(t) and f(t) at $t = 500$ h?
4. The diameter of bearing manufactured by a company under the specified supply conditions has a normal distribution with a mean of 10 mm and standard deviation of 0.2 mm
 - (i) Find the probability that a bearing has a diameter between 10.2 and 9.8 mm?
 - (ii) Find the diameters, such that 10 % of the bearings have diameters below the value?
5. While testing ICs manufactured by a company, it was found that 5 % are defective. (i) What is the probability that out of 50 ICs tested more than 10 are defective? (ii) what is the probability that exactly 10 are defective?
6. If the rate of failure for a power supply occurs at a rate of once a year, what is the probability that 5 failures will happen over a period of 1 year?
7. Given the following 20 failure times, estimate R(t), F(t), f(t), and $\lambda(t)$: 100.84, 580.24, 1210.14, 1630.24, 2410.89, 6310.56, 3832.12, 3340.34, 1420.76, 830.24, 680.35, 195.68, 130.72, 298.76, 756.86, 270.39, 130.0, 30.12, 270.38, 720.12.
8. Using the data given in problem 7, identify possible distribution with the help of probability plotting method?

References

1. Ross SM (1987) Introduction to probability and statistics for engineers and scientists. Wiley, New York
2. Montgomery DC, Runger GC (1999) Applied statistics and probability for engineers. Wiley, New York
3. Weibull W (1951) A statistical distribution of wide applicability. J Appl Mech 18:293–297

4. Sturges HA (1976) The choice of class interval. *J Am Stat Assoc* 21:65–66
5. Ebeling CE (1997) An introduction to reliability and maintainability engineering. Tata McGraw-Hill Publishing Company Ltd, New Delhi
6. Misra KB (1992) Reliability analysis and prediction. Elsevier, Amsterdam

Reliability and Safety Engineering

Verma, A.K.; Ajit, S.; Karanki, D.R.

2016, XX, 571 p. 249 illus., Hardcover

ISBN: 978-1-4471-6268-1