

Some Common Misconceptions about the Modeling of Repairable Components

C. J. Zapata, *Member IEEE*, A. Torres, *Senior Member IEEE*, D. S. Kirschen, *Fellow IEEE*,
M. A. Ríos, *Member IEEE*

Abstract—Although stochastic point process theory has been successfully applied in many fields of knowledge, in power systems reliability it has not received so much attention what is reflected in the low number of reported applications. This may be due to some common misconceptions about the modeling of repairable components which falsely show this method is the same than other popular ones. All these misconceptions originate in the incorrect practice of analyzing the reliability of repairable components using concepts that were developed only for non repairable ones and, specifically, in the misleading idea that a stationary random process model can represent a non stationary random process. This paper discusses these misconceptions expecting clarity of concepts can foster the development of more applications of this theory in power systems reliability.

Index Terms— Point processes, Poisson processes, power system reliability, reliability, stochastic processes.

I. INTRODUCTION

SINCE long ago, stochastic point process (SPP) theory has been successfully applied in many fields of knowledge such as biology, physics, queuing analysis, engineering reliability, etc. [1]–[3]; statistical procedures for applying this type of modeling to real problems have been developed and several SPP models have gained wide acceptance. On the other hand, SPP has not received as much attention in power system reliability as in other fields and only a small number of applications have been reported, e. g. [4]–[10]. This may be due to some common misconceptions about the reliability modeling of repairable components. In particular, it is often believed that SPP is identical to others widely used such as the analyses based on the Weibull distribution. The aim of this paper is to bring some clarity about SPP theory by discussing the origin of these misconceptions.

II. REVIEW OF BASIC CONCEPTS [11]–[13]

Before discussing the misconceptions, it is necessary to review some fundamental concepts about random processes.

A. Definitions

The term random process denotes a random phenomenon that is observed in the real world. The term stochastic process is reserved for a kind of modeling for random processes. The period of interest for studying a random process is denoted t . A random variable x represents the random process.

B. Stationary and Homogeneous Random Process

A random process is stationary if their statistical properties, the expectation $E[x]$ and the variance $V[x]$, are constant during t . The opposite is true for a non stationary random process.

A random process is time homogeneous if its probability density function $f(x)$ does not change during t . The opposite is true for a non homogeneous random process.

Homogeneous and stationary are interchangeable terms because: *i.* If $f(x)$ does not change during t then $E[x]$ and $V[x]$ are constant during this period. *ii.* If $E[x]$ and $V[x]$ are constant during t it is necessary that $f(x)$ does not change during this period. Non homogeneous and non stationary are also interchangeable terms.

C. Distribution Model

A distribution is a mathematical model for a stationary random process in which t does not explicitly appear. A distribution is defined by means of a probability density function $f(x)$ which do not change during t . All mathematical functions used as distributions produce $E[x]$ and $V[x]$ because this kind of model always refers to a stationary random process; hence $E[x]$ and $V[x]$ are only functions of the distribution parameter which are also constant.

D. Stochastic Process Model

A stochastic process is a mathematical model for a stationary or non-stationary random process in which t appears explicitly. The random variable that represents the process can then be written x_t and t is called the process index. Thus, a stochastic process is a collection of random variables $x_{t_1}, x_{t_2}, \dots, x_{t_N}$, one for each value of the index t . There is thus, a collection of probability density functions $f_{t_1}(x), f_{t_2}(x), \dots, f_{t_N}(x)$ one for each random

C. J. Zapata is a professor at Universidad Tecnológica de Pereira, Pereira, Colombia, and a PhD student at Universidad de los Andes, Bogotá, Colombia. (e-mail: cjzapata@utp.edu.co).

A. Torres is a professor at Universidad de los Andes, Bogotá, Colombia. (e-mail: atorres@uniandes.edu.co).

D. S. Kirschen is a professor at The University of Manchester, Manchester, U. K. (e-mail: daniel.kirschen@manchester.ac.uk).

M. A. Ríos is a professor at Universidad de los Andes, Bogotá, Colombia. (e-mail: mrrios@uniandes.edu.co).

variable. If for a given t_i the statistics of the random process are constant, it is stationary and time homogeneous because $f_{t_i}(x)$ do not change during this period. The opposite is true for a non-stationary non-homogeneous random process.

E. How to Select a Model for a Random Process [14]-[16]

Fig. 1 shows the basic procedure for selecting a model that is a proper representation of a random process. Omitting any of the three steps of this procedure can lead to an unsuitable model. A sample $x_1, x_2, \dots, x_n$ is the input data for this procedure.

The first step is to determine whether the random process is stationary or non stationary. Several statistical methods are available for this. However, only trend tests [17] are discussed because only sequences of times to failure (t_{tf}) and times to repair (t_{tr}) are considered in this paper.

Fig. 2 shows a simple trend test where the bar graph shows the chronologically ordered inter arrival time magnitudes. If this graph shows a pattern of increasing or decreasing inter arrival time magnitudes, then the random process is deemed to have a tendency or that it is non stationary. If this test does not show that the random process has a tendency, it is deemed to be stationary. The basic condition to guarantee the validity of a trend test is to keep the chronological order in which the inter arrival times occurred.

If the sample data for the random process shows that it is non stationary, a non stationary stochastic process model has to be selected. This can be done by applying the procedures for parameter estimation and a goodness of fit test, which are particular for each model in this class and should not be confused with the ones used for distributions. Two important families of non-stationary stochastic processes are the non-homogeneous Markov chains and the non-homogeneous Poisson processes.

If the sample data for the random process under study shows it is stationary, is necessary to apply a test for independency such as the scatter diagram or the correlation plot [18]. Two cases arise here:

1. If the sample data is not independent, a model for dependent events has to be selected. An example of these kinds of models is the branching point process or time series.
2. If the sample data is independent, a distribution must be selected if t is not necessary to explain the random process. If that is not the case, a stationary stochastic process must be selected. In both cases it is necessary to apply the procedures for parameter estimation and a goodness of fit test to select the distribution or the stationary stochastic process model that can represent the random process under study.

The importance of performing trend and independency tests is discussed by Ascher and Hansen [15] who point out that:

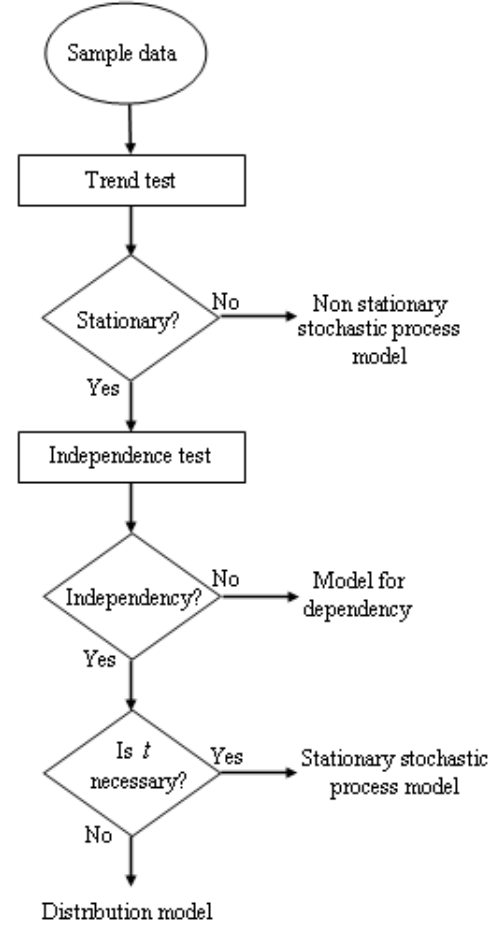


Fig. 1. Procedure to select a model for a random process

1. It is incorrect to fit a sample of inter arrival times to a distribution model without performing first a trend test to check that the random process from which the sample was taken is stationary. Goodness of fit tests sorts out sample values by magnitude therefore losing the chronological order in which they occurred.
2. It is incorrect to fit a sample of inter arrival times to a distribution model without performing first an independency test because the goodness of fit tests, such as chi square and Kolmogorov-Smirnov, were developed assuming sample independency. This also applies to the maximum likelihood method for parameter estimation.

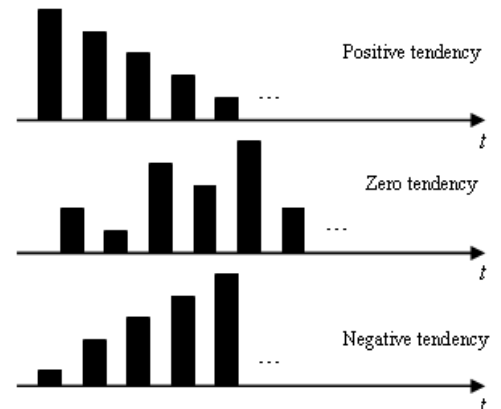


Fig. 2. Bar graphs of inter arrival times magnitudes for trend test

III. RELIABILITY ANALYSIS OF NON REPAIRABLE COMPONENTS

A non-repairable component is one that dies when the first failure f_1 occur. The classical model for this kind of component is shown in Fig. 3. It only considers two operating states and ttf is used to represent the failure process.

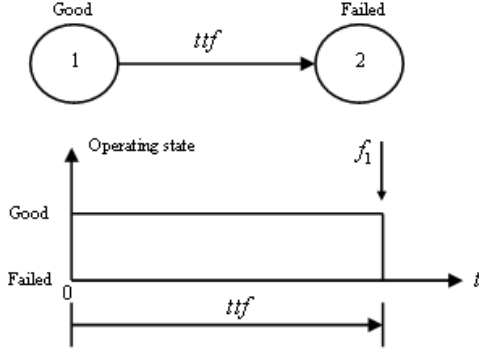


Fig. 3. Operating states of a non-repairable component

Because a non-repairable component can fail only once, a sample $ttf_1, tt f_2, \dots, tt f_n$ obtained from a group of identical components that have failed is necessary to build its reliability model. Fig 4 shows such a sample. These values are not ordered in a chronological sequence and each ttf_i has no connection with the other sample values. Furthermore, the instant when the observation of the operating time was taken does not matter.

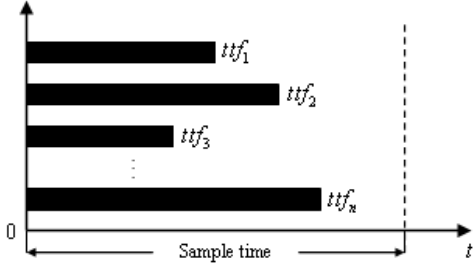


Fig. 4. Sample of ttf of a group of identical non-repairable components

The ttf sample is fitted to a distribution $f_{ttf}(t)$ that is called life model. $F_{ttf}(t)$ gives the probability of failure and its complement $R_{ttf}(t) = 1 - F_{ttf}(t)$ is the reliability.

One important aspect to study for non-repairable components is the risk that a component that has not failed until a given time t fails after it. This is a conditional probability that leads to the famous equation for $\lambda(t)$ called “failure rate” or “hazard rate” [19], [20]:

$$\lambda(t) = \lim_{\Delta t \rightarrow 0} \frac{R(t) - R(t + \Delta t)}{\Delta t \cdot R(t)} = \frac{f_{ttf}(t)}{[1 - F_{ttf}(t)]} \quad (1)$$

Depending on the kind of distribution used for the life model or the values of its parameters, $\lambda(t)$ can be constant or a function of time; only for the exponential distribution $\lambda(t)$ is a constant, for a Gaussian distribution it is an increasing function of time, etc.

For a Weibull distribution with scale parameter λ and shape parameter β , $\lambda(t)$ is defined by (2). As shown in Fig. 5 the form of $\lambda(t)$ depends on the value of β .

$$\lambda(t) = \lambda \beta t^{\beta-1} \quad (2)$$

Equation (2) has a ubiquitous place in reliability. Unfortunately, as will be discussed later, this has led some authors to forget its real meaning and origin.

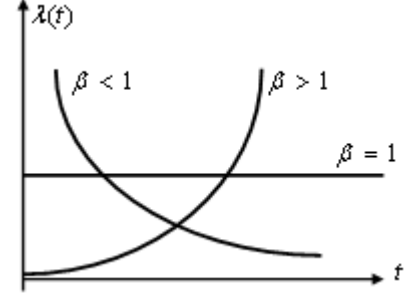


Fig. 5. Failure rate for a non-repairable component with Weibull life model

IV. RELIABILITY ANALYSIS OF REPAIRABLE COMPONENTS

A repairable component is one that can withstand a sequence of failures $f_1, f_2, \dots, f_n$. Its simplest representation in terms of reliability is the two-state diagram shown in Fig. 6.

The independent processes of failures and repairs can be illustrated by the operating sequence shown below the two state diagram in Fig. 6. Unlike the case of a non-repairable component, in this case, the sample values $ttf_1, tt f_2, \dots, tt f_n$ and $ttr_1, ttr_2, \dots, ttr_n$ must be chronologically ordered sequences to keep the tendency of the failure and repair processes.

The two main families of models that have been applied to the reliability analysis of repairable components are discussed below.

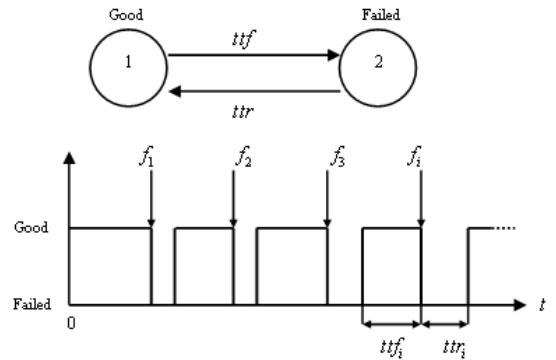


Fig. 6. Two state diagram and operating sequence of a repairable component

A. Markov Chain Models

The term Markov chain refers here to a family of models which couples the processes of failures and repairs in a two state diagram representation such as the one shown in Fig. 6. This definition is adopted because there is no agreement about the names for the different extensions to the basic continuous-time exponential Markov chain model.

1) Homogeneous Exponential Markov Chain

If the samples of tff and ttr show no tendency, are independent and meet a goodness of fit test for exponential distributions with parameters $\lambda = 1/\overline{tff}$ and $\mu = 1/\overline{ttr}$, respectively, the coupled process of failures and repairs is described by [19]-[20]:

$$\begin{pmatrix} dP_1(t)/dt \\ dP_2(t)/dt \end{pmatrix} = \begin{pmatrix} -\lambda & \mu \\ \lambda & -\mu \end{pmatrix} \begin{pmatrix} P_1(t) \\ P_2(t) \end{pmatrix} \quad (3)$$

$P_1(t)$ and $P_2(t)$ are the probabilities of finding the component in states 1 (good) and 2 (failed), respectively. λ and μ are called “failure rate” and “repair rate”, respectively, or more generally “transition rates”. Overlines symbols denote a statistical mean. The most appealing characteristic of this model is that it has an analytical solution.

This model is memoryless or Markovian, i.e. the transition to another state depends only on the current state and the trajectory before reaching the present state does not matter. This model is commonly called homogeneous Markov process or homogeneous Markov chain.

2) General Homogeneous Markov Chain

In this case, samples of tff and ttr show no tendency, are independent and one or both of them meet the goodness of fit test with a non exponential distribution. When both distributions are not exponential, this model is called “non Markovian process” and for the case where one is exponential but the other not it is called “semi-Markov process”. We adopt the name general homogeneous Markov chain because “general” indicates that any kind of distributions can be used and “homogeneous” specifies that these distributions do not change with time. This model does not have the memoryless property i. e. it is non-Markovian And cannot be solved using (3). Solutions methods include Monte Carlo simulation, the device of stages and the technique of adding variables [19], [20].

This model is very important because it is unusual for both the failure and repair distributions to be exponential. While the failure process for non-aged components generally fits an exponential distribution, repair times are generally lognormally distributed [21], [22].

3) Non Homogeneous Markov Chain

In this case, λ and μ in (3) are not constant but functions of time. The failure and repair processes are thus not homogeneous because as time evolves the expected number of failures and the expected number of repairs are not constant. Therefore, the failure and repair processes cannot be represented by means of distributions. This model also does not have the memoryless property i. e. it is non-Markovian.

Popular solutions to this process are numerical methods of differential equations and sequential Monte Carlo simulation. However, it has problems for adjusting the operating times and of tractability for some types of time varying rates [23].

B. SPP Models

As shown in Fig. 7, this kind of modeling decouples the processes of failures and repairs of the component. Failures and repairs are represented by sequences of events that arrive independently.

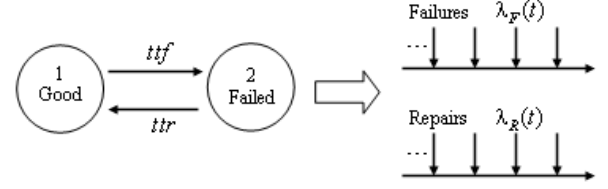


Fig. 7. In SPP modeling the process of failures and repairs are uncoupled

In many applications the repair process is neglected because the repair times are much shorter than the typical interval of time separating failures. For example, the repair time may be of the order of hours, compared to times to failure in the order of years.

1) Definition

A SPP is a model that counts the number of events N that occur in t . This model has as basic condition that one and only one event can occur at every instant. Fig. 8 shows a pictorial representation of a SPP; x_i denotes an inter-arrival interval and t_i an arrival time.

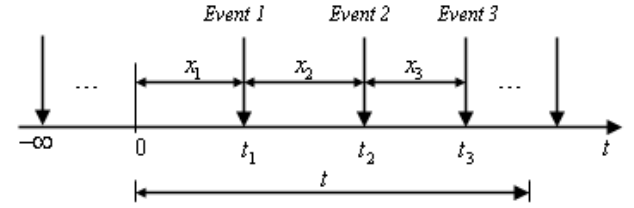


Fig. 8. The concept of SPP

A SPP model is defined by means of the parameter $\lambda(t)$ called the intensity function, which is defined mathematically, as follows:

$$\lambda(t) = dE[N(t)]/dt \quad (4)$$

Depending on the application of the SPP, $\lambda(t)$ can be a failure rate, a repair rate, a flooding rate, etc.

2) Classification

The tendency of the inter-arrival times allows the classification of SPP models shown in Fig. 9.

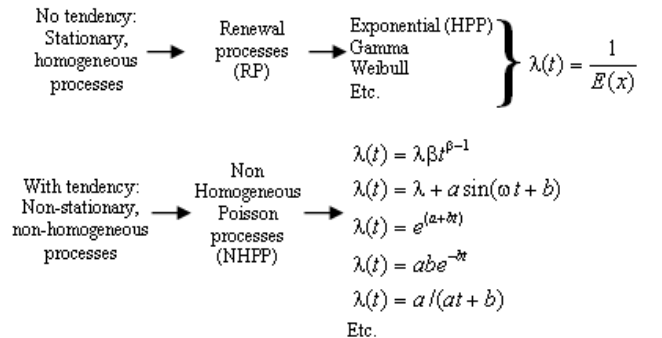


Fig. 9. A basic classification of SPP

The name for a RP is given after the x 's distribution. The most famous RP is the exponential one, commonly called Homogeneous Poisson process (HPP). For $t \rightarrow \infty$, the intensity function of every RP is a constant defined as [24]:

$$\lambda(t) = 1 / E(x) \quad (5)$$

3) The Power Law NHPP

The Power Law Process (PLP) developed by Crow in 1974 [25] is a NHPP model that is widely used in the field of electrotechnics to represent the failure process of repairable components [26]. Its intensity function is (2). The shape parameter determines the tendency of the model; for $\beta > 1$ the tendency is positive, for $\beta < 1$ it is negative and for $\beta = 1$ the tendency is zero and the PLP is equal to the HPP.

V. THE MISCONCEPTIONS

A. The meaning of the term "failure rate"

The first problem that arises is the failure rate given by (4) is confused with the one in (1) when a SPP is used to model the failure process of a repairable component. The two concepts are different:

1. The failure rate (1) refers to failures that affect a population of identical non-repairable components and kill them. For a single non-repairable component, it can neither be calculated nor measured.
2. The failure rate (4) refers to failures that affect a single repairable component if the sample was taken from a particular component. Also, it can refer to failures that affect a population of identical or non-identical repairable components if component failure data were pooled.

In order to distinguish the two concepts, Ascher and Feingold [14] proposed the term *ROCOF* (Rate of Occurrence Of Failures) for (4). While this lexical distinction is useful, it is essential to understand what definition applies for repairable and non-repairable components; it is incorrect to use the definition (1) for repairable components or the definition (4) for non-repairable ones. However, in many papers (1) is presented as the failure rate of components that are repairable such as power transformers, generators, etc. In [27] Thompson discusses the uses and abuses on the application of (1).

B. The use of a life model for a repairable component

The life model of a non-repairable component $f_{iff}(t)$ refers to the arrival of one and only one failure that kills it. Thus, is incorrect to apply this concept to a repairable component as it can withstand several failures. But what happen if an analyst takes a sample of tff from a repairable component and, after applying required tests, shows that a given distribution is a valid representation of this failure process and calls it the component's life model with failure rate defined by (1)? Although the procedure is correct, the way the analyst conceives the model is flawed:

1. As explained before, the failure rate (1) does not apply.

2. The distribution represents the inter-arrival times of failures. It can be used to calculate the probability that tff is less or equal than a given value, for generating a sequence of time to failures or for defining a RP failure model with failure rate given by (5).
3. The distribution is not a life model because it does not defines the death of the repairable component. Such an event is defined mainly by economic consideration: a failed component is deemed to have died and is replaced if its repair cost is equal or higher than its replacement cost or if the expected cost of its unavailability during a planning period is higher than its replacement cost.

C. A distribution can represent a non stationary random process

This is the most misleading idea in reliability! A distribution can only be used to model stationary random processes. All mathematical functions used as distributions produce constant statistics. This fact can be easily proven using a bar diagram of a sequence of values generated from any distribution. Fig. 10 shows this for a realization of a Weibull distribution with $\lambda = 5$ [years] and different values of β . As can be seen, in all cases there is not tendency.

Similarly, RP are always stationary because they are defined on the basis of the distribution of inter arrival times. Thus, in [27] Thompson points out a RP cannot model component aging and discusses this misconception.

This misconception originates from (1); as it can produce increasing or decreasing failure rates depending on the kind of distribution or in accordance to the value of its shape parameter it is believed (or more precisely, misbelieved) that this is a natural property of some distributions. Thus, in some papers a time varying failure rate is defined for a repairable component and without a theoretical support the tff are generated using an exponential or Weibull distribution.

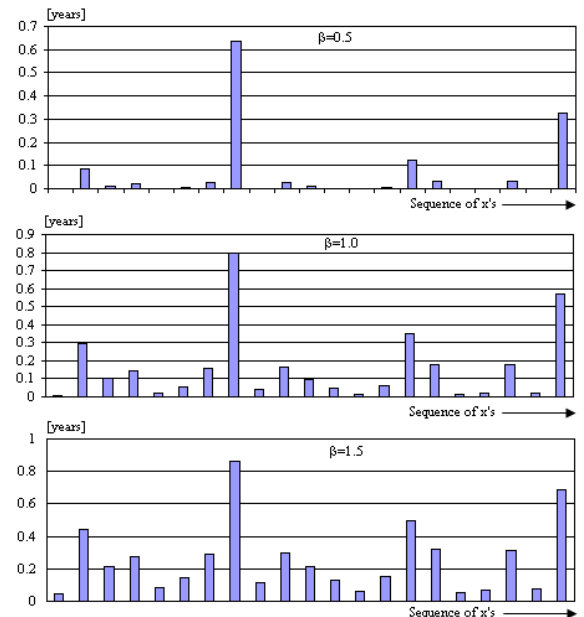


Fig. 10. Bar graphs of the values generated from a Weibull distribution.

D. Equation (2) generates a random process whose model is the Weibull distribution

This misconception is a consequence of the previous one. The truth is that if (2) is used as intensity function for a SPP or as transition rate for a Markov chain, an HPP is obtained when $\beta=1$ and a non stationary one when $\beta \neq 1$. This can be proven using the algorithm given in Appendix A. More importantly, this is valid for any random process and not only for those which pertain to failures. The relationship between the Weibull distribution and (2) is restricted only to the case where the reliability of a non-repairable component is studied.

This misconception originates in the fact that the concept expressed by (2) has been applied extensively in the reliability field forgetting in many cases its origin and meaning. For example:

1. Many books and papers show it as a natural property of the Weibull distribution. Results obtained by means of (2) are only valid when referring to the reliability of a non-repairable component, a particular result of an application where the Weibull distribution is applied.
2. Many papers define the failure rate for a repairable component using (2) and tell it belongs to the Weibull distribution although they are applying a proper method for a non stationary analysis. This is, the analysis is correct but they are bringing a concept that does not apply.

E. A general homogeneous Markov chain can represent a non stationary process

This misconception is also a consequence of misconception C. It is not true because a distribution always refers to a stationary process.

The bar diagram shown in Fig 5 proves this for a Weibull distribution. In addition, let us consider now the method called the device of stages [19], [20]; for some pairs of distributions (exponential-lognormal, exponential-Weibull, etc.), it transforms the two state general homogeneous Markov chain in an exponential one that has more than two states. Fig. 11 shows an example: the exponential-lognormal chain is transformed into an exponential one where state 2 is replaced by k stages in series ($2S_1$ to $2S_k$) and two stages in parallel ($2P_1$ and $2P_2$). Transition rates ρ , $\rho\omega_1$, $\rho\omega_2$, ρ_1 , and ρ_2 are constants obtained from the fourth first moments of the lognormal distribution. If the equivalent exponential Markov chain obtained using the device of stages, which is stationary, solves the two state general homogeneous Markov chain, how can the latter be not-stationary? However some papers apply the device of stages and say that for the case Weibull – lognormal it represents a non stationary process!

This misconception originates from misbelieving the transition rates of a general homogeneous Markov chain are defined by means of (1). This is wrong because there is no connection between the concepts of transition rate of a Markov chain and hazard rate of a non repairable component. Concept (1) can not be extended to failures of a repairable component neither to other events such as repairs!

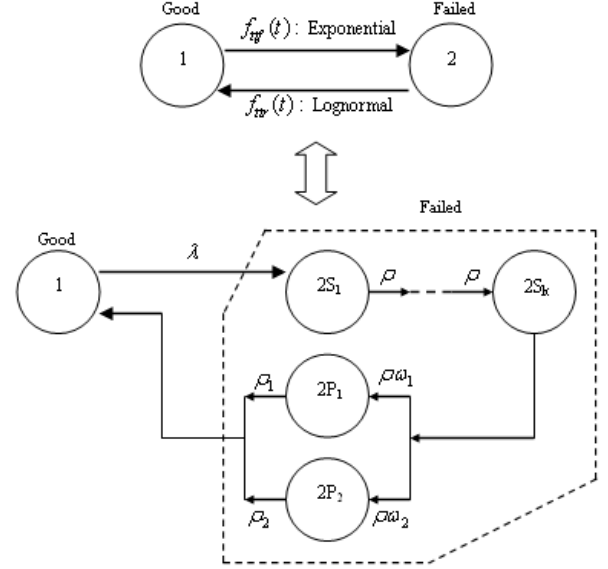


Fig. 11. The device of stages for solving a given homogeneous Markov chain

F. The PLP is the same thing as a Weibull distribution.

The arguments presented in section V-D show that this is false. The PLP has no connection with the Weibull distribution. The origin of this misconception is the fact that PLP intensity function is the mathematical function (2). However, when applying (2) it should be remembered the context of application:

1. For a non repairable component, it refers to a sequence of failures that affect a population of identical non-repairable components, not to the process of failure arrivals to a single non-repairable component neither to the arrival of other, non-failure, events.
2. For a repairable component, it refers to a sequence of events that arrive. It is not confined to the case of failures. And in the case of failures, it can represent the process of failure arrival to a repairable component or to a population of repairable components.

G. The PLP is the same thing as a Weibull RP

The arguments presented in section V-D can be used to show that this is a misconception. In a PLP an exponential stationary process is obtained when $\beta=1$ and a non-stationary one when $\beta \neq 1$. When $\beta=1$ it generates a HPP not a Weibull RP. This misconception has the same origin that the one discussed in section V-F. Another factor that reinforces this misconception is that PLP has received other names with the word Weibull such as Weibull process, Weibull-Poisson process, Rasch-Weibull process [28].

H. The only model for a stationary failure process is the HPP

This is probably the most common of all misconceptions, but it is not as misleading as the one discussed in V-C. This statement is only valid when a sample of t_{tf} taken from repairable component shows no tendency, is independent and complies with the goodness of fit test for an exponential distribution. But, what happens if the sample shows no

tendency, is independent, but does not comply with a goodness of fit test for the exponential distribution? In this case, it is incorrect to assume an exponential distribution; the failure process of the repairable component has to be represented by means of the RP of a distribution that satisfies a goodness of fit test.

This misconception originates again from the concept of failure rate for a non-repairable component (1); it produces a constant failure rate only for the case of an exponential distribution. Thus, “constant failure rate = HPP model” has been applied as a rule of thumb for any type of components, forgetting this results was obtained only for non repairable ones. For the case of a repairable component with stationary failure process, all RP are possible failure models.

VI. RELATIONSHIP BETWEEN SPP AND MARKOV CHAINS

A two state Markov chain is generated by two SPP process as shown in Fig. 12; every time a failure arrives to the component it is sent from the good state to the failed one and every time a repair is performed, the component comes back to the good state; the sources of this motion are the SPP.

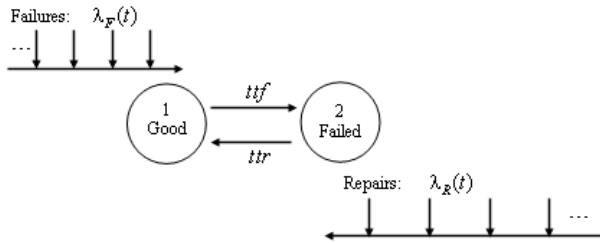


Fig. 12. Relationship between a two state Markov chain and SPP

Intensity functions $\lambda_F(t)$ and $\lambda_R(t)$ in the SPP models are equal to transition rates $\lambda_{12}(t)$ and $\lambda_{21}(t)$ in the Markov chain, respectively, regardless of whether the models are defined using distributions or non stationary stochastic processes.

One could therefore argue that, since both types of models are equivalent, there is no reason to use a SPP when Markov chains are a more popular method. While this would be true at the component level, analysts usually deal with systems of repairable components. When dealing with large repairable systems, the repair process should not be included in the component level because [4]:

1. It is equivalent to assume repair resources are unlimited because every time the component fails a crew is available to repair it, or in other words, there is a repair team dedicated to each component. Hence, an implicit assumption is made that repair times depend only on the particular actions taken to fix each type of component.
2. As shown in Fig 13, for maintenance activities, a power system is usually split into several zones or service territories, and repair teams are assigned to each area. The repair process performed in each service territory is really a queuing system like the one shown in Fig 14.

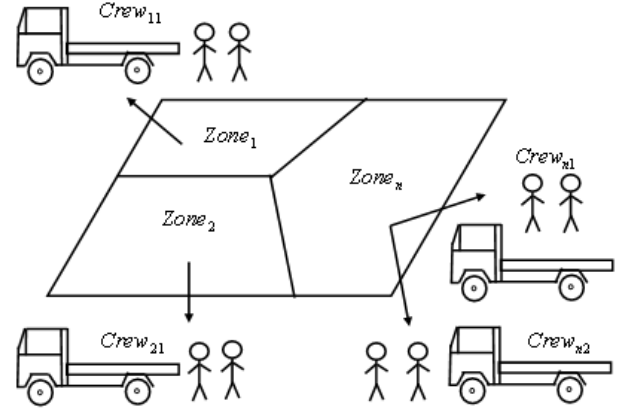


Fig. 13. Zones for maintenance in a power system

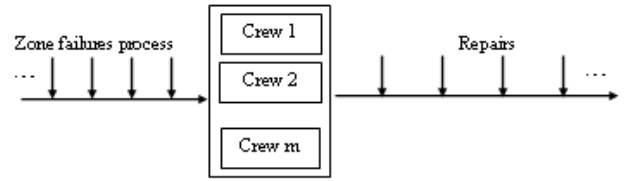


Fig. 14. The repair process in a maintenance zone of a power system

For this queuing system the following must be defined:

- Input process: The superimposed failure process generated by all components located in the service territory. These failures are related to service interruptions.
- Service process: The SPP that represents the crews' capacity and generates the repair times.
- Output process: The SPP of the repairs performed by crews. These repairs are related to service restorations.

Since repair resources are limited some failures will have to wait while others that failed before are repaired.

SPP modeling thus makes it possible to represent the repair process performed in each area of a large repairable system as it really happens. This is something that Markov chain modeling is unable to do.

VII. CONCLUSIONS

There are several common misconceptions about the modeling of repairable components for reliability studies. In particular, it is often assumed that SPP are identical to other methods currently in widespread use, for example, the popular analyses based on the Weibull distribution.

All these misconceptions originate in the incorrect practice of analyzing the reliability of repairable components using concepts that were developed only for non-repairable ones and, specifically, in the misleading idea that a stationary random process model can represent a non-stationary random process.

Reliability engineers must consider carefully the concepts of homogeneity and stationarity of random processes, the procedure for selecting a type of model for a random process and the differences among the main types of models that are available.

VIII. APPENDIX A – ALGORITHM TO GENERATE SAMPLES FROM NHPP [18]

1. Generate a sequence of n inter arrival times $x_1', x_2', \dots, x_n'$ of a HPP with intensity function $\lambda = 1.0$ which covers a sample period T using:

$$x_i' = -1/\lambda * LN(U_i) \quad (A1)$$

Where U is a uniformly distributed random number.

2. Convert the sequence $x_1', x_2', \dots, x_n'$ to a sequence of arrival times called $t_1', t_2', \dots, t_n'$ using:

$$t_i' = \sum_{k=1}^i x_k' \quad (A2)$$

3. Find the expected number of events

$$\Lambda(t) = \int_0^t \lambda(t) dt \quad (A3)$$

4. Find (Λ^{-1}) .

5. Calculate the arrival times of the NHPP using:

$$t_i = \Lambda^{-1}(t_i') \quad (A4)$$

6. Calculate the sequence of arrival times of the NHPP $x_1, x_2, \dots, x_n$ using:

$$x_i = t_i - t_{i-1} \quad (A5)$$

Where $t_0 = 0$.

As pointed out by Law and Kelton [18] the application of this algorithm depends on how easy the inversion of Λ is.

In the case of the PLP the recursive equation is:

$$t = (t' \setminus \lambda)^{1/\beta} = \Lambda^{-1}(t') \quad (A6)$$

IX. REFERENCES

- [1] Cox D. R, Isham V, *Point Processes*, Chapman and Hall, 1980.
- [2] Snyder D. L, *Random Point Processes*, Wiley, 1975.
- [3] Srinivasan S. K, *Stochastic Point Processes*, Griffin, 1974.
- [4] Zapata C. J, Silva S. C, Gonzales H. I, Burbano O. L, Hernández J. A, "Modeling the repair process of a power distribution system", *IEEE Transmission & Distribution Latin America Conference & Exhibition*, 2008.
- [5] Schilling M. T, Praca J. C. G, Queiroz J. F, Singh C, Ascher H, "Detection of ageing in the reliability analysis of thermal generators", *IEEE Transactions on Power Systems*, Vol. 3, No.2, 1988.
- [6] Kogan V. I, Jones T. L, "An explanation for the decline in URD cable failures and associated nonhomogeneous Poisson Process", *IEEE Transactions on Power Delivery*, Vol. 9, No. 1, January 1994.
- [7] Kogan V. I, Gursky R. J, "Transmission towers inventory", *IEEE Transactions on Power Delivery*, Vol. 11, No. 4, October 1996.
- [8] Stillman R. H, "Modeling failure data of overhead distribution systems", *IEEE Transactions on Power Delivery*, Vol. 15, No. 4, October 2000.
- [9] Stillman R. H, "Power line maintenance with minimal repair and replacement", *IEEE Annual Reliability and Maintainability Symposium*, 2003.
- [10] Balijepalli N, Subrahmanyam S. Venkata, Christie R. D, "Modeling and analysis of distribution reliability indices", *IEEE Transactions on Power Delivery*, Vol. 19, No. 4, October 2004.
- [11] Papoulis Athanasios, *Probability, Random Variables and Stochastic Processes*, Mc-Graw Hill, 1991.
- [12] Lawyer G. F, *Introduction to Stochastic Processes*, Chapman and Hall, 1995
- [13] Cox D. R, Miller H. D, *The Theory of Stochastic Processes*, Methuen, 1965.
- [14] Ascher H, Feingold H, *Repairable systems reliability: Modeling, inference, misconceptions and their causes*, Marcel Dekker, 1984.

- [15] Ascher H. E, Hansen C. K, "Spurious exponentially observed when incorrectly fitting a distribution to nonstationary data", *IEEE Transactions on Reliability*, Vol. 47, No. 4, December 1998.
- [16] Choy S, English J. R, Landers T, Yan L, "Collective approach for modeling complex system failures", *Proceedings of Annual Reliability and Maintainability Symposium*, IEEE, 1996.
- [17] Wang P, Coit D. W, "Repairable systems reliability trend test and evaluation", *IEEE Annual Reliability and Maintainability Symposium*, 2005.
- [18] Law A. M, Kelton D. W, *Simulation Modeling and Analysis*, Mc-Graw Hill, 2000.
- [19] Billinton R, Allan R. N, *Reliability Evaluation of Engineering Systems: Concepts and Techniques*, Plenum Press, 1992.
- [20] Singh C, Billinton R, *System reliability modeling and evaluation*, Hutchinson Educational Publishers, 1977.
- [21] Zapata C. J, Silva S. C, Burbano O. L, Hernández J. A, "Repair models for power distribution components", *IEEE Transmission & Distribution Latin America Conference & Exhibition*, 2008.
- [22] Jacobson D. W, Arorn S. R, "A non exponential approach to availability modeling", *IEEE Annual Reliability and Maintainability Symposium*, 1995.
- [23] Hassett T. F, Dietrich D. L, Szidarovszky F, "Time-varying failure rates in the availability & reliability analysis of repairable systems", *IEEE Transactions on Reliability*, Vol. 44, No. 1, March, 1995.
- [24] Cox D. R, *Renewal Theory*, Methuen, 1962.
- [25] Crow L. H, "Evaluating the reliability of repairable systems", *IEEE Annual Reliability and Maintainability Symposium*, 1990.
- [26] IEC *Power law model – Goodness-of-fit test and estimation methods*, IEC Standard 61710, 2000.
- [27] Thomson W. A, *Point Processes Models with Applications to Safety and Reliability*, Chapman and Hall, 1988.
- [28] Klefsjo B, Kumar U, "Goodness-of-fit tests for the power law process based on the TTT plot", *IEEE Transactions on Reliability*, Vol. 41, No. 4, December 1992.
- [29] Rigdon S. E, Basu A. P, *Statistical methods for the reliability of repairable systems*, Wiley, 2000.

X. BIOGRAPHIES

Carlos J. Zapata (S'1993, AM'1997, M'2004) obtained his BScEE from the Universidad Tecnológica de Pereira, Pereira, Colombia, in 1991 and his MScEE from the Universidad de los Andes, Bogotá, Colombia, in 1996. He worked for eleven years for Concol S. A, Bogotá, Colombia and since 2001 he has worked as a professor at the Universidad Tecnológica de Pereira, Pereira, Colombia. Currently, he is working towards his PhD at the Universidad de los Andes, Bogotá, Colombia.

Alvaro Torres (M'1972, SM'1989) obtained his BScEE from the Universidad Industrial de Santander, Bucaramanga, Colombia, in 1975. In 1977 he received his MEng in electrical engineering from Rensselaer Polytechnic Institute, Troy, USA. In 1978 he received his MSc in computer science and his PhD in electrical engineering from Rensselaer Polytechnic Institute, Troy, USA. Since 1982 he has worked as a professor at the Universidad de los Andes, Bogotá, Colombia, and as Technical Manager for Concol S. A, Bogotá, Colombia.

Daniel S. Kirschen (S'1980, M'1984, SM'1993) obtained his Electrical and Mechanical Engineer's degree from the Université Libre de Bruxelles, Brussels, Belgium, in 1979 and his MSc (1980) and PhD (1985) degrees in Electrical Engineering from the University of Wisconsin-Madison, USA. He worked for ten years for Control Data Corporation and Siemens Energy and Automation, Minneapolis, USA, and since 1994 he has worked as a professor at UMIST and The University of Manchester, Manchester, United Kingdom.

Mario A. Ríos (M'1989) obtained his BScEE (1991), MScEE (1992) and PhD (1998) from the Universidad de los Andes, Bogotá, Colombia. He also received in 1998 the PhD from the Institut National Polytechnique du Grenoble, Grenoble, France. He worked for 12 years for Concol S. A, Bogotá, Colombia, and since 2005 he has worked as a professor at the Universidad de los Andes, Bogotá, Colombia.