



UNIVERSIDADE ESTADUAL PAULISTA
“JÚLIO DE MESQUITA FILHO”
Faculdade de Ciências e Tecnologia
Câmpus de Presidente Prudente

Objective Bayesian Analysis of Process Capability Indexes for Some Lifetime Models

Marcello Henrique de Almeida

Academic Supervisor: Prof. Dr. Fernando Antonio Moala

Post-Graduate Program in Applied and Computational Mathematics

Presidente Prudente, February 2021

UNIVERSIDADE ESTADUAL PAULISTA

Faculdade de Ciências e Tecnologia de Presidente Prudente

Post-Graduate Program in Applied and Computational Mathematics

Objective Bayesian Analysis of Process Capability Indexes for Some Lifetime Models

Marcello Henrique de Almeida

Academic Supervisor: Prof. Dr. Fernando Antonio Moala

Dissertation presented to the Graduate Program in Applied and Computational Mathematics of the Faculty of Sciences and Technology of the Universidade Estadual Paulista "Júlio de Mesquita Filho", Presidente Prudente campus to obtain a Master's degree in Applied and Computational Mathematics.

Presidente Prudente, February 2021

A447o	Almeida, Marcello Henrique Objective Bayesian analysis of process capability indexes for some lifetime models / Marcello Henrique Almeida. -- Presidente Prudente, 2021 65 p. : il., tabs., fotos Dissertação (mestrado) - Universidade Estadual Paulista (Unesp), Faculdade de Ciências e Tecnologia, Presidente Prudente Orientador: Fernando Antonio Moala 1. Objective Bayesian Inference. 2. Matching prior, Reference priors. 3. Gamma Distribution. 4. Weibull Distribution. 5. Process Capability Index. I. Título.
-------	---

Sistema de geração automática de fichas catalográficas da Unesp. Biblioteca da Faculdade de Ciências e Tecnologia, Presidente Prudente. Dados fornecidos pelo autor(a).

Essa ficha não pode ser modificada.



UNIVERSIDADE ESTADUAL PAULISTA

Câmpus de Presidente Prudente

CERTIFICADO DE APROVAÇÃO

TÍTULO DA DISSERTAÇÃO: Objective Bayesian Analysis of Process Capability Indexes for Some Lifetime Models

AUTOR: MARCELLO HENRIQUE DE ALMEIDA

ORIENTADOR: FERNANDO ANTONIO MOALA

Aprovado como parte das exigências para obtenção do Título de Mestre em MATEMÁTICA APLICADA E COMPUTACIONAL, pela Comissão Examinadora:

Prof. Dr. FERNANDO ANTONIO MOALA (Participação Virtual)
Departamento de Estatística / Faculdade de Ciências e Tecnologia de Presidente Prudente

VIDEOCONFERÊNCIA

Prof. Dr. EDILSON FERREIRA FLORES (Participação Virtual)
Departamento de Estatística / Faculdade de Ciências e Tecnologia de Presidente Prudente

VIDEOCONFERÊNCIA

Prof. Dr. PEDRO LUIZ RAMOS (Participação Virtual)
Pós-doutorado / Instituto de Ciências Matemáticas e de Computação / Universidade de São Paulo

Presidente Prudente, 26 de fevereiro de 2021

*I dedicate this work to my mother Margarida
and my father Antonio.*

Acknowledgements

First, I thank God for sustaining and blessing me throughout my journey to this day.

I thank my parents Antonio and Margarida for supporting me throughout my academic path, always encouraging me.

I am very grateful to my girlfriend Bruna Daiane Freitas Souza for always being by my side even from far away, supporting and encouraging me during the whole master's period.

I cannot forget to thank my sister Marcia for her support and incentive and for always being by my side.

I would like to thank my supervisor Professor Dr. Fernando Antonio Moala for guiding me during the process of this work and for all the times I needed help.

I must give my immense thanks to Pedro Luiz Ramos, for his great partnership since the beginning of the dissertation writing, for his patience, for guiding me through my difficulties, because I believe that without his help I would not have been able to develop all the work in such a short time. I still hope to be able to produce countless articles through this partnership.

I thank my classmates who were by my side during these two years, especially the student Liner Samer Albornoz Sandoval for his friendship during all this time.

I thank the other people from UNESP and also from the city of Presidente Prudente, especially Anderson José Arias dos Santos, because they were people who somehow were part of my trajectory throughout the master's period.

Finally, I would like to thank the Virtual University of the State of São Paulo (UNIVESP) for the opportunity to be part of the facilitators of this University and also for the financial support, making it possible for me to stay in the program.

*“Maybe I couldn’t do the best, but I fought for the best to be done.
I’m not what I should be, but thank God, I am not what I was before”.*
Marthin Luther King

Resumo

A distribuição Gama tem sido aplicada em pesquisas em diversas áreas do conhecimento, devido a sua boa flexibilidade e adaptabilidade. A distribuição Weibull desempenha um papel importante no controle de confiabilidade e monitoramento da qualidade. Estes modelos têm sido amplamente utilizados para descrever o índice de capacidade de processo (PCI) quando os dados não seguem uma distribuição normal. Sob este cenário, os estudos atuais se concentram na estimativa dos parâmetros usando a inferência clássica. Neste trabalho, consideramos métodos Bayesianos para estimar o PCI denominado C_{pk} a partir de uma perspectiva objetiva, utilizando matching priori e priores de referência. A inferência proposta é ampliada para uma versão generalizada da distribuição Weibull que fornece um bom ajuste para dados mais complexos com comportamento de risco não-monotônico. As distribuições posteriores são construídas e são propostos estimadores Bayesianos com base na mediana para a distribuição Weibull e estimadores Bayesianos com base na média para a distribuição Gama. Neste caso, os métodos Markov Chain Monte Carlo são usados para obter as estimativas e, a partir de um extenso estudo de simulação, observamos que bons resultados são observados em termos de erros médios relativos e quadráticos. A abordagem proposta também é utilizada para construir intervalos de credibilidade adequados com baixo custo computacional e probabilidades de cobertura precisas. É apresentada uma aplicação de dados reais que confirma que nossa abordagem proposta tem um desempenho superior aos métodos atuais.

Palavras chave: *inferência bayesiana objetiva; matching priori; prioris de referência; distribuição Gama; distribuição Weibull; índice de capacidade de processo C_{pk} .*

Abstract

The Gamma distribution has been applied in research in several areas of knowledge, due to its good flexibility and adaptability nature. The Weibull distribution plays an important role in reliability control and quality monitoring. These models have been widely used to describe the process capability index (PCI) when the data do not follow a normal distribution. Under this scenario, current studies focus on parameter estimation using classical inference. In this work, we consider Bayesian methods to estimate the PCI denominated C_{pk} from an objective perspective using matching prior and reference priors. The proposed inference is extended to a generalized version of the Weibull distribution that provides a good fit for more complex data with non-monotonic risk behavior. The posterior distributions are constructed and median based Bayesian estimators are proposed for the Weibull distribution and mean-based Bayesian estimators for the Gamma distribution. In this case, Markov Chain Monte Carlo methods are used to obtain the estimates and from an extensive simulation study, we observe that good results are observed in terms of mean relative and quadratic errors. The proposed approach is also used to construct adequate credibility intervals with low computational cost and accurate coverage probabilities. An application of real data is presented that confirms that our proposed approach outperforms current methods.

Keywords: *objective Bayesian inference; matching prior; reference priors; Gamma distribution; Weibull distribution; Process capability index C_{pk} .*

List of Figures

2.1	Graph of the probability density functions for different values of α and β and their respective cumulative density functions.	26
2.2	Risk function and survival function of the Gamma distribution with different values for the parameters.	27
2.3	Graph of the probability density functions for different values of λ and α and their respective cumulative density functions	28
2.4	Risk function and survival function of the Weibull distribution with different values for the parameters.	28
2.5	Graph of the probability density functions for different values of ϕ , λ and α and their respective cumulative density functions.	30
2.6	Risk function and survival function of the distribution GW with different values for the parameters.	31
4.1	Simulation study for the Weibull distribution assuming different values for the C_{pk} , $N = 10.000$ and $n = 50, 60, \dots, 200$	47
5.1	Simulation study for the Generalized Weibull distribution assuming different values for the C_{pk} , $N = 10.000$ and $n = 50, 60, \dots, 300$	52
6.1	Control Chart for the Juice data set.	54

List of Tables

3.1	MRE, MSE and CP from the estimates of Cpk considering different values of n with $N = 10.000$ simulated samples using Classical and Bayesian inference.	39
3.2	MRE, MSE and Coverage Probability from the estimates of Cpk considering different values of n with $N = 10.000$ simulated samples using Classical and Bayesian inference.	40
6.1	Data set related to 30 packages of strawberry and grape juice	53
6.2	Estimates and 95% credibility intervals for the C_{pk} with its respective AIC and BIC.	54

List of Acronyms

MLE: Maximum Likelihood Estimation.

MRE: Mean Relative Error.

MSE: Mean Square Error.

MCMC: Markov Chain Monte Carlo.

ASQC: American Society for Quality Control.

ABNT: Brazilian Association of Technical Standards.

ISO: International Standardization Organization.

INMETRO: National Institute of Metrology Quality and Technology.

PCI: Process Capability Index.

PPM: Part per Million.

CDF: Cumulative Density Function.

GW: Generalized Weibull.

PDF: Probability Density Function.

LSL: Lower specification Limit.

USL: Upper Specification Limit.

SPC: Statistical Process Control.

IPEM: Weights and Measures Institute.

Contents

Resumo	7
Abstract	9
List of Figures	10
List of Table	11
List of Acronyms	15
Chapters	
1 Introduction	19
2 Background	21
2.1 Process Capability Indexes	21
2.1.1 C_p Index	22
2.1.2 C_{pk} Index	22
2.1.3 C_{pm} index	23
2.1.4 C_p index for non normal distribution	23
2.1.5 C_{pk} index for non normal distribution	24
2.1.6 C_{pm} index for non normal distribution	24
2.2 Probability Distributions	25
2.2.1 Gamma Distribution	25
2.2.2 Weibull Distribution	27
2.2.3 Generalized Weibull Distribution	29
2.3 Maximum likelihood inference	31
2.4 Bayesian Inference	32
3 Gamma Distribution	35
3.1 Introduction	35
3.2 Maximum Likelihood Estimation	36
3.3 Bayesian Inference for Gamma Distribution	37
3.4 Markov Chain Monte Carlo method	37
3.5 Simulation Study for Gamma Distribution	39
4 Weibull distribution	43
4.1 Introduction	43
4.2 Bayesian Inference for the Weibull distribution	44
4.2.1 MCMC methods for Weibull Distribution	45
4.3 Simulation Study for Weibull distribution	46

5	Generalized Weibull distribution	49
5.1	Bayesian Inference for the Generalized Weibull distribution	49
5.1.1	Markov Chain Monte Carlo Method for GW	50
5.2	Simulation Study Generalized Weibull distribution	52
6	Application	53
7	Final Considerations and Future Proposals	57
	Bibliography	58

Introduction

The concept of statistical quality control was first introduced in the mid-1920s by mathematician Walter Shewhart. From then on several organizations related to quality in general emerged, such as the American Society for Quality Control (ASQC), the Brazilian Association of Technical Standards (ABNT) and the International Standardization Organization (ISO), National Institute of Metrology Quality and Technology (INMETRO).

Later, principally with the advance of technology, industries began to fight more and more for space in the labor market. Thus, it is clear that in all industrial sectors it is essential to have quality in their production in order to be able to guarantee themselves in the face of such a disputed market.

Still in the production scope, in the confection of products in general can occur failures in their specification due to lack of monitoring. In this sense, there are two possibilities: the first of them is related to the dissatisfaction of the clients by the lack of quality of the product; on the other hand it can happen that the company is wasting its input during the production to try to avoid that the client is deceived.

From this perspective we have two tools: the control chart and the process capability indexes. These two mechanisms can contribute to the verification of the measurement system used by the company so that we have good results in terms of the products manufactured and whether the process is being able to generate items within the desired specifications.

Studies through probabilistic models is one of the main challenges currently highlighted in the field of statistical analysis. It is important to realize that with the advent and the new discoveries connected with the advance of new technologies has made possible the analysis through more robust models in order to be able to apply them in real life problems.

In the literature we can find several probability distributions that are generally applied to problems related to survival study, reliability, quality control, among others. Among them we can mention Gamma, Weibull, Exponential, Generalized Weibull (or Gamma-Weibull), among others.

In this work we use three probabilistic models to develop the study of interest. First, we will talk about the Gamma distribution, which is a model that has two parameters (α, β) where α represents the shape and β denotes the scale. This model is one of the best known in statistical analysis and commonly used in several application areas due to its great flexibility in adjusting data in areas such as environmental analysis, clinical trials among other real situations.

The Weibull distribution is another distribution widely used in studies related to practical applications because this model presents a monotonous risk function, that is, increasing or decreasing, but it can also be constant. It is worth remembering that

this model has great importance because this distribution is the generalization of the exponential model that unlike the Exponential Distribution, the Weibull Distribution has an importance property of keeping memory, in other words, here is considered the time already used of an equipment in relation to a new unit and without use. See some applications of this Distribution in Teimouri et al. [56] and Ramos et al. [38].

The Generalized Weibull Distribution is also a model that has proven very useful in the field of analysis of survival data due to its various forms presented in its risk function such as: growing, decreasing, constant, bathtub shape and inverse bathtub. This model also has several particular cases of distributions that can be reached from its parameters. Some examples are: Gamma, Weibull, Exponential, Chi-square, Rayleigh. Thus, this model can be very effective in explaining data in different fields of study see Ahsanullah et al. [1], Raju and Srinivasan [37], Marani et al. [26] and Santos [44].

In the literature, it is easy to find results related to the process capability analysis assuming the data come from a Normal Distribution. In the opposite direction, it does not happen the same thing for data with Non-Normal Distribution, with that there are very few references that approach this subject. Thus, this work aims to treat the analysis of process capability through data from Non-Normal Distributions. For more details on the approach to non-normal data and see González and Herner [14].

We showed some of the capability indexes that are available in the literature, but in the present work we have used the index called C_{pk} to obtain results about the capacity of a certain process to produce items with the least possible variation. C_{pk} has the objective of saying if a specific process is being viable in relation to loss, time, resources and costs and, mainly, if it is being able to offer good manufacturing and that it meets the needs and desires of the clientele.

In order to exemplify the present work, all the content presented is exposed as follows: In the first chapter, we present an introduction to the subject of interest where we briefly talk about statistical quality control, the emergence of industries and the need to produce with quality and respecting all the established standards, we also emphasize the importance of computer technology for the development of new research applied to real data. We also mention some important distributions that are very well known and also widely used in studies related to the most varied types of data of interest. Finally, in the introduction we also talk about the C_{pk} capability index that we use in the study under the perspective of data originated from Non-Normal Distribution. In the second chapter, we dedicate to the different capability indexes arranged in the literature under the normal and non-normal approach. We defined each index and explained in a way that stimulates the calculation of each one of them and the C_{pk} classification. Still in chapter 2, we present in a general way the maximum likelihood estimator, we show each of the distributions that we will use throughout the work and also their properties through the moments of each one and we also graphically illustrate their functions PDF, CDF, risk functions and survival functions. Finally, in the second chapter, we briefly review Bayesian inference. Chapter 3, composes the first paper we developed with the objective of estimating the C_{pk} index through the Gamma distribution. Chapter 4 and 5 is based on the second paper that we developed for estimating the C_{pk} index using the Weibull and the GW distributions, respectively. Finally, in chapter 6 illustrates our proposed approach in a real data set. Finally, in chapter 7 we conclude by summarizing the present study and also some future proposals.

Background

2.1 Process Capability Indexes

Process capability indexes are employed to appraise management aspiration to raise its earnings potential and at the same time to compete with existing company's by focus on plummeting the variability in the manufacturing process. Also, at the same time produced items should meet the company's specifications and the customer's anticipation. To estimate the amount of process capability of the manufacturing process, it is required to characterize a quantitative measure of the product that justifies the performance of the company. Sometimes the product quality can be assessed by sampling inspection whereas it may not interpret quality specifications laid down by the company or customer's anticipation. To analyze quality of the product by using manufacturing technology, the process capability analysis can be measured through various stages of the manufacturing process including process, product design, manufacturing, and manufacturing planning; for more details, see Statistic and Tehnike [53]. A number of process capability indices have been proposed to determine whether the manufactured product is capable or not and also product meets the specifications given by the company or satisfies the customer. More literature review on process capability indexes can be checked in Kotz and Johnson [18], Spiring et al. [51], Sappakitkamjorn and Niwitpong [45] and Yum and Kim [61] and Piña-Monarrez et al [34]. Zhang rightly pointed out that the two most frequently process capability indices are the C_p and C_{pk} .

To estimate these process capability indexes the basic assumption on quality characteristics under study should follow normal distribution. Whereas for industrial production which is concerned with most of quality characteristics may not follow the normal distribution or quality characteristic distribution, in these situations the process capability indices are not accurate; for more details, see Sennaroglu and Senvar [47], Piña-Monarrez et al. [34], Senvar and Kahraman [49, 48]. Hosseinifard et al. [16] studied the effect of non-normal data to estimate the process capability indexes and assessed that data lead to inaccurate results. For non-normally distributed data, a quantile approach to measure the parametric estimation of process capability indexes were developed by Clements [9]. Hence, more researchers focused on the study of the process capability analysis for the non-normal data, for more details refer Kashif et al [17], Panichkitkosolkul [32], Leiva et al. [20], Senvar and Sennaroglue [50], Saha et al. [43, 42], Meng et al. [27], Mahmoud et al [24], and Wu [60].

The PCI when used for asymmetric data is a measure that is achieved through the percentiles of the target distribution and the appropriate specifications for the process.

The specifications are given by the lower specification limit (LSL) and upper specification limit (USL) and can be established according to the interest of the researcher or through regulatory quality control agencies, as we will see in the application section. There are several indexes available in the literature, among them, we can mention some of the commonly used indexes in statistical quality control which are referred to as C_p , C_{pk} , C_{pm} , C_{pmk} .

The demands on the quality of products produced to meet consumer demand have increased in the past decades. Quality inspections began with the purpose of verifying the production and studying actions to improve the stages of the process aiming at a final product with acceptable characteristics for the customers' satisfaction. In this sense, an important tool used in reliability analysis is linked to quantitative measurements and aligned with manufacturing processes that are generally applied in the measurement of reliability and are called process capability indexes.

González and Werner [14] show in detail that problems arise when companies perform capacity analysis assuming normality of data when in fact the correct would be to consider data coming from non-normal models. In this way, the responsible professional may be making wrong conclusions about the process.

Therefore, we will describe some capability indexes utilizing data from a normal distribution and also data not normally distributed.

2.1.1 C_p Index

The C_p is denominated index of potential process capability and was the first index to be introduced in the literature. The calculation of the C_p is based on a ratio that compares two values where the numerator is defined as the dispersion of the specification and the denominator is established by the variation of the process. It is worth remembering that the C_p does not take into account the location of the data, but only the width of the process, because here it is considered that the process is always accommodated around the average. Thus, the C_p for normally distributed data can be defined as follows

$$C_p = \frac{USL - LSL}{6\sigma} \quad (2.1)$$

Where LSL and USL are the lower and upper specification limits, respectively, and σ is defined as the process standard deviation.

In general, this index considers that for higher C_p values it means that the process is considered capable, in contrast, for lower C_p values indicate that the process needs improvement. However, as we know that the C_p does not consider the location of the process, then it is not possible to assess whether this process is within or near the region defined by the specification limits. Therefore, to obtain more reliable information it is necessary to combine the results obtained here with another capability index, for example use C_{pk} to obtain desired conclusions for the data set of interest.

2.1.2 C_{pk} Index

As we saw before the C_p index has some deficiencies in its accuracy, so in order to overcome this impasse we have the C_{pk} which is a better index than the C_p index. The present index is also based on the minimum between C_{pl} and C_{pu} . The minimum is also defined from a ratio that compares the values of the distance from the process mean to the closest specification limit (LSL or USL) in relation to the one-sided variation of the process. Unlike the C_p this index considers the location and variation of the data. In this

way, C_{pk} for data from a normal distribution is defined as

$$C_{pk} = \min \left(\frac{USL - \mu}{3\sigma}, \frac{\mu - LSL}{3\sigma} \right) \quad (2.2)$$

where LSL and USL are the lower and upper specification limits, respectively. μ and σ are defined as the mean and standard deviation coming from the data. As we have mentioned before, it is worth noting that C_{pk} evaluates potential capability based on the location of the data and also on the dispersion of the process. In general for higher C_{pk} values indicate that the process is being capable, already for lower C_{pk} values it means that the process is flawed and needs improvements.

An important piece of information is that if $C_p = C_{pk}$ is the same it means that the process is centered around the average of the specification, but if $C_p \neq C_{pk}$ then the process is not centered around the average, i.e. the process average is not the same as the average of the specifications given by (LSL and USL).

2.1.3 C_{pm} index

As an alternative to the indexes shown above, we have the C_p index, which is a capability index that besides considering the process variation also takes into account the distance of the process average of interest compared to the average of the nominal specification value. Analogously to the items mentioned above, the C_{pm} is also given through a ratio where the numerator is defined as the distance between the lower and upper specification limits and the denominator is given through a radical involving the process mean, the nominal specification mean and the variance of the target process. Thus, the C_p can be achieved as follows

$$C_{pm} = \left(\frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}} \right) \quad (2.3)$$

Where USL and LSL are the lower and upper specification limits, respectively; σ^2 is the variance given by the process; μ and T are the process mean and the nominal specification mean, in that order. According to its definition, if the process has a large variability it causes the denominator to increase, consequently the C_{pm} value will decrease dramatically. Likewise, if μ and T have a considerable distance it can also influence the C_{pm} to decrease. However, C_{pm} has advantages when compared to C_p for example, because this index produces a good representation about the capacity of the process both for data that have close to the nominal value and for data that are further from the nominal average. It is worth noting that if $\mu = T$ we have $C_{pm} = C_p$ for data originating from a normal distribution.

2.1.4 C_p index for non normal distribution

When it is not possible to assume that the process data can be explained through a normal distribution it is not viable to apply the C_p and C_{pk} indexes shown above. Therefore, it is necessary to study alternative methods to be able to perform the capacity analysis through a convenient distribution for the available data. Under this perspective we will define two new indexes that have been developed and are available in the literature for data that do not come from the normal distribution. These new indexes are quite similar to the traditional indexes seen previously when assuming normality of the data. The first index is the C_p which is calculated using a ratio where the numerator is defined

as the width between the upper and lower specification limits and the denominator is given by varying the 0.135 and 99.865 percentiles of the distribution of interest. Here as in the case of normality the C_p does not take into account the location of the process with respect to the specification limits, only its dispersion, due to the fact that in this case it is considered that the process is always located around the nominal mean of specification. Of this fact, the C_p is given by

$$C_p = \frac{USL - LSL}{F_{99.865} - F_{0.135}} \quad (2.4)$$

Where LSL and USL are the lower and upper specification limits; $F_{99.865}$ e $F_{0.135}$ are the 99.865 and 0.135 position percentiles related to the process data, respectively.

2.1.5 C_{pk} index for non normal distribution

According to the above writing, the C_p index may not be effective to analyze the process, because according to its definition we do not have the idea of location. So from this assumption, we will present another more robust index that besides considering the variation of the process also takes into account its respective location. This index is denominated C_{pk} and can be found from the minimum between two ratios that are called C_{pl} and C_{pu} and can be calculated by:

$$C_{pk} = \min \left(\frac{USL - M}{U_p - M}, \frac{M - LSL}{M - L_p} \right), \quad (2.5)$$

where U_p , L_p and M are the percentiles 99.865^{th} , 0.135^{th} and 50^{th} from the Gamma distribution, respectively, USL and LSL denote the upper and lower specification limits. It is important to point out that the percentiles U_p , L_p and M mentioned above, can be easily computed from a function already implemented in the statistical software R.

The interpretation for the C_{pk} index is given by the items in non-conformity assuming one million items produced and describe by

- If $C_{pk} < 1$, then the process is producing over 2700 items in nonconformities and the process is considered incapable.
- If $1 \leq C_{pk} \leq 1.33$, then the process is generating between 64 and 2700 non-conforming manufacturing, and the procedure is considered acceptable.
- If $C_{pk} \geq 1.33$, then the process reproduces a quantity below 64 nonconforming items, hence, the process is classified as capable.

2.1.6 C_{pm} index for non normal distribution

As the C_{pm} index seen above only considers data that originate from a normal distribution, so the literature provides us with another index that is also called C_{pm} , but more robust, and can be used for any non-normal distribution. This index is very similar to the C_{pm} index when assuming normal data. Similarly, the index here is defined by a division where the numerator is also given by the width of the lower and upper specification limits, while the denominator is calculated by a radical involving the 99.865 and 0.135 percentiles of the target distribution; the process mean and the nominal mean of the specification limits. In view of this, the C_{pm} can be measured as follows

$$C_{pm} = \frac{USL - LSL}{6 \sqrt{\left[\frac{F_{99.865} - F_{0.135}}{6} \right]^2 + (M - T)^2}} \quad (2.6)$$

Where USL and LSL are the upper and lower specification limits, $F_{99.865}$ and $F_{0.135}$ are the 99.865 and 0.135 percentiles of the target distribution, respectively; M and T are defined as the process median and the nominal value of the specifications in that order.

In this chapter we present three important groups of indexes that are commonly used in the analysis of process capability using normal and non-normal data. We show the characteristics of each of these indexes and also the way used in the calculation of each model of index.

From the indexes presented in this chapter, it is possible to notice that the only index that mentions both the variability given through the lower and upper specification limits and the distance between the process data average and the specification limits in its numerator is the C_{pk} index, that is, besides considering the permitted variation of the process, the C_{pk} also takes into account the location of the data. In this way, the C_{pk} has an advantage compared to other indexes, because in practice it is not always possible to establish the process centered on the nominal specification mean.

In chapter 3, 4 and 5, we will use the C_{pk} index based on non-normal data to perform the study that aims to estimate the C_{pk} for non-normal data with the support of the Gamma, Weibull and WG distributions. In addition to C_{pk} we also present the indexes C_p and C_{pm} , which remains as a suggestion for future studies.

It is worth pointing out that we present the interpretation for the C_{pk} index only for convenience, because the same clarification is valid for the other capability indexes presented in this chapter.

2.2 Probability Distributions

Here we will talk about the probability distributions we use in the course of the work. We will show their respective probability density functions (PDF), cumulative density function (CDF), survival functions and risk functions. The distributions chosen are well known in the field of statistical analysis due to their wide applicability in different real problems.

2.2.1 Gamma Distribution

Let X be a continuous random variable whose probability distribution is given by the model Gamma. Its probability density function and cumulative density function are defined, respectively as

$$f(x|\alpha, \beta) = \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\beta x} \quad \text{and} \quad F(x|\alpha, \beta) = 1 - \int_x^{+\infty} \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\beta x} \quad (2.7)$$

with $x > 0$, $\alpha > 0$, $\beta > 0$ e $\Gamma(\alpha)$ is given by

$$\Gamma(\alpha) = \int_0^{\infty} e^{-x} x^{\alpha-1} dx \quad (2.8)$$

and is known in mathematics as a gamma function. The mean and variance of the gamma distribution is, respectively given by

$$E(X) = \int_0^{\infty} \frac{\beta^\alpha x^\alpha e^{-\beta x}}{\Gamma(\alpha)} dx = \frac{\alpha}{\beta} \quad (2.9)$$

$$Var(X) = \int_0^{\infty} \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha+1} e^{-\beta x} dx - E(X)^2 = \frac{\alpha}{\beta^2} \quad (2.10)$$

This distribution belongs to the exponential family of continuous probability distributions with two parameters. In the literature, there are reparametrized version of this model, hereafter, we will use the form given in (2.7). The Gamma distribution is an important generalization of the exponential distribution ($\alpha = 1$), assuming that $\alpha = \frac{n}{2}$, $n \in \mathbb{Z}$ and $\beta = \frac{1}{2}$ one can show that the model reduces to the chi-square distribution with n degrees of freedom.

Since the Gamma distribution is used to describe nonnegative data with positive skewness, the model is widely used in survival and reliability analysis (see, e.g., Lawless 1982), as it is a model that is adaptable and returns, in most cases, a good adjustment for the event of interest, for example, in the areas of engineering, meteorology, climatology, and among other situations.

In the Figure 2.1 we can see the different forms of the PDF and CDF of the Gamma distribution for different α and β values. Note that the PDF chart shows different shapes according to the choice of parameters, consequently, the CDF also varies according to this choice.

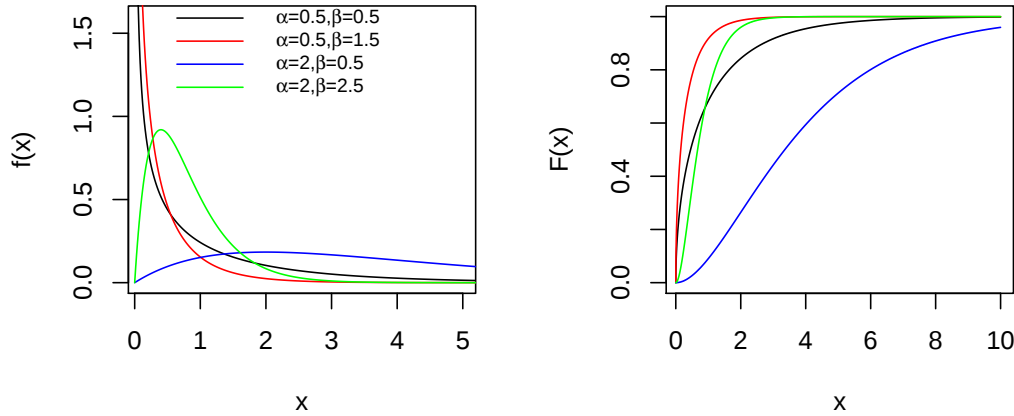


Figure 2.1: Graph of the probability density functions for different values of α and β and their respective cumulative density functions.

The risk functions (failure rate function) and survival function (reliability function) for the Gamma distribution are defined, respectively by

$$h(x|\alpha, \beta) = \frac{\beta^\alpha x^{\alpha-1} \exp(-\beta x)}{\Gamma[\beta x, \alpha]} \quad (2.11)$$

$$S(x|\alpha, \beta) = 1 - P(X \leq x) = \frac{\Gamma[\beta x, \alpha]}{\Gamma(\alpha)} \quad (2.12)$$

where $P(X \leq x)$ is the cumulative density function of (2.7) and $\Gamma[\beta x, \alpha] = \int_\alpha^\infty m^{\beta x-1} e^{-m} dm$ it is known as incomplete Gamma function.

We can see in Figure 2.2 the graphs of the risk and survival functions for some α and β values, respectively.

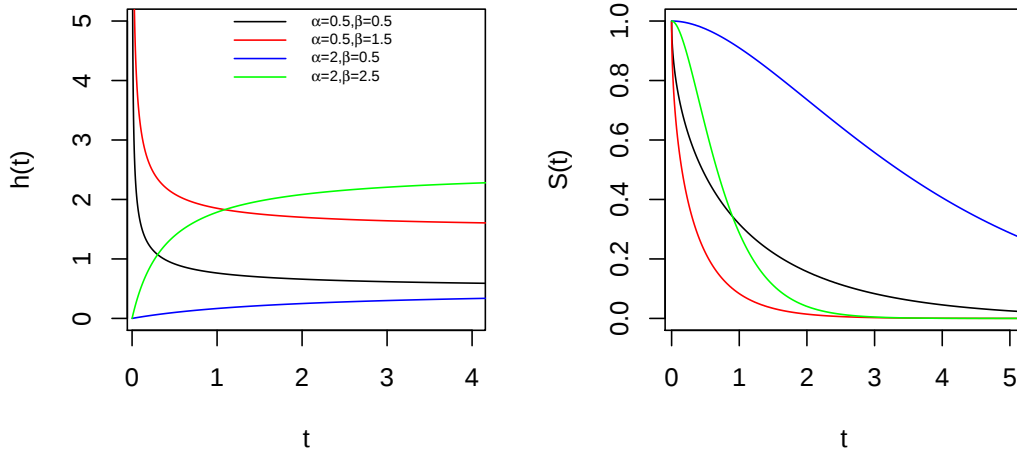


Figure 2.2: Risk function and survival function of the Gamma distribution with different values for the parameters.

2.2.2 Weibull Distribution

Although many asymmetric models have been proposed in recent years. The most used model is the Weibull [59] distribution, such a model arises in a reliability context and it is used as the base for the development of new methodologies and applications in industry. Let T be a continuous random variable whose probability distribution follows Weibull. Its probability density function and cumulative density function are expressed, respectively by

$$f(t|\lambda, \alpha) = \lambda \alpha t^{\alpha-1} e^{-\lambda t^\alpha} \quad \text{and} \quad F(t|\lambda, \alpha) = 1 - e^{-\lambda t^\alpha} \quad (2.13)$$

for all $t > 0$ and the quantities $\alpha > 0$ and $\lambda > 0$ are the shape and the scale parameters respectively. Here, we focus on the use of the parametrization in the form of (2.13), however, other reparametrizations can be considered. The mean and variance of the random variable that follows the Weibull distribution can be expressed, respectively by

$$E(T) = \alpha \Gamma \left[1 + \frac{1}{\lambda} \right] \quad \text{and} \quad Var(T) = \alpha^2 \left[\Gamma \left[1 + \frac{1}{\lambda} \right] - \Gamma \left[1 + \frac{1}{\lambda} \right]^2 \right] \quad (2.14)$$

For $\forall t, \lambda$ and α positive.

The Weibull distribution is an important generalization of the exponential distribution ($\alpha = 1$) as well as the Rayleigh distribution ($\lambda = 2$), among others. The model arises in different areas with interesting physical interpretations [25] and its parameters can be used to understand its hazard behavior. Due to the popularity of the Weibull distribution, many studies have considered this model as a baseline distribution to estimate PCIs as an alternative to the Gaussian model. The parameter estimates are usually derived under the maximum likelihood estimators [34] and bootstrap techniques are applied to construct confidence intervals [12, 58]. A major drawback is that the MLEs are usually biased for small samples which may compromise the analysis, in addition, the computational cost of bootstrap approaches are very significant due to the high number of replications as well as the application of numerical techniques to estimates the parameters of each sub-sample.

Note that Figure 2.3 illustrates the graphs of the probability density function and cumulative density function for the Weibull distribution for different pairs of parameters λ and α .

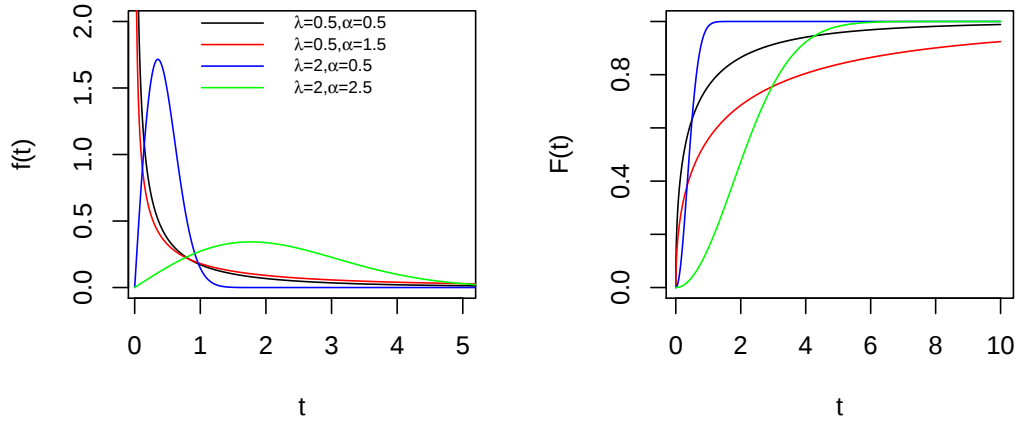


Figure 2.3: Graph of the probability density functions for different values of λ and α and their respective cumulative density functions .

The risk functions (failure rate function) and survival function (reliability function) for the Weibull distribution can be expressed respectively by

$$h(t|\lambda, \alpha) = \frac{\lambda}{\alpha} \left(\frac{t}{\alpha}\right)^{\alpha-1} \quad (2.15)$$

$$S(t|\lambda, \alpha) = 1 - P(T \leq t) = \exp\left(-\left(\frac{t}{\alpha}\right)^\lambda\right) \quad (2.16)$$

where $P(T \leq t)$ is the cumulative density function of the Weibull distribution.

The Figure 2.4 represents the risk and survival functions of the Weibull distribution with different pairs of parameters λ and α .

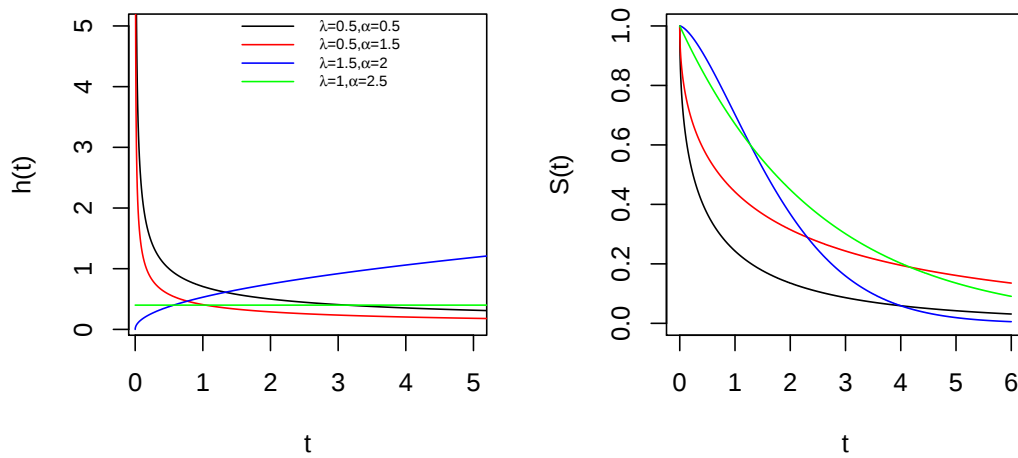


Figure 2.4: Risk function and survival function of the Weibull distribution with different values for the parameters.

2.2.3 Generalized Weibull Distribution

Although the Weibull distribution can be used in many situations to describe asymmetric data, it cannot be used to describe datasets with non-monotone behavior in the hazard rate. In this case, a flexible generalization of the Weibull model can be considered instead. Stacy [52] introduced a unified model with three-parameters that contain the Weibull, Gamma, Lognormal, Nakagami-m, half-normal, Maxwell-Boltzmann, chi distribution, among others. This model is referred to by different names such as generalized Gamma, Gamma-Weibull, or generalized Weibull (GW) distribution. Hereafter, we will refer to such a model as GW distribution due to its relation with the Weibull model. The parameter estimation under the MLE was previously discussed by [15]. However, Prentice [35] showed that the MLEs may not return good estimates, as well as the asymptotic property of the confidence intervals may not be achieved even for sample sizes larger than 400. Therefore, we have generalized the Bayes estimators for the GW distribution, the Bayesian reference posterior is used to achieve the estimates of the PCI returning precise estimates with good coverage properties. Finally, our proposed approach is applied in a real dataset aiming to improve the quality process and reduce losses.

Let T be a continuous random variable whose probability distribution can be explained by the probability density function given by

$$f(x|\phi, \lambda, \alpha) = \frac{\alpha}{\Gamma(\phi)} \lambda^{\alpha\phi} x^{\alpha\phi-1} \exp(-(\lambda x)^{\alpha}), \quad (2.17)$$

where t, ϕ, λ, α positives and $\Gamma(\phi)$ is the mathematical function called Gamma.

The cumulative density function of the distribution GW is given by

$$\int_0^{(\lambda t)^{\alpha}} \frac{1}{\Gamma(\phi)} m^{\phi-1} e^{-m} dm = \frac{\gamma[\phi, (\lambda t)^{\alpha}]}{\Gamma(\phi)} \quad (2.18)$$

where $\gamma[y, x] = \int_0^x m^{y-1} e^{-m} dm$.

The mean and variance of the Generalized Weibull distribution can be found through the Moment Generator function and are expressed as follows:

$$E(T) = \frac{\phi + \frac{1}{\alpha}}{\lambda \Gamma(\phi)} \quad \text{and} \quad Var(T) = \frac{1}{\lambda^2} \left\{ \frac{\Gamma(\phi + \frac{2}{\alpha})}{\Gamma(\phi)} - \left[\frac{\Gamma(\phi) + \frac{1}{\alpha}}{\Gamma(\phi)} \right]^2 \right\} \quad (2.19)$$

If a continuous random variable T follows a distribution GW, then its risk and survival function are defined as:

$$h(t|\phi, \lambda, \alpha) = \frac{\alpha \lambda^{\alpha\phi} t^{\alpha\phi-1} \exp(-(\lambda t)^{\alpha})}{\gamma[\phi, (\lambda t)^{\alpha}]}, \quad (2.20)$$

$$S(t|\phi, \lambda, \alpha) = 1 - P(T \leq t) = 1 - \int_0^{(\lambda t)^{\alpha}} \frac{1}{\Gamma(\phi)} m^{\phi-1} e^{-m} dm \quad (2.21)$$

$$= 1 - \frac{\gamma[\phi, (\lambda t)^{\alpha}]}{\Gamma(\phi)} \quad (2.22)$$

where $P(T \leq t)$ is the cumulative density function from the GW distribution.

The Figure 2.5 illustrates some forms of the probability density function and the cumulative density function of the GW distribution for different values of the parameters (ϕ, λ, α) .

The distribution GW has different forms in its risk function and can be: constant, increasing, decreasing, bathtub (or inverse bathtub). Due to this flexibility the GW has been commonly used to model problems in the field of reliability. There is a method to

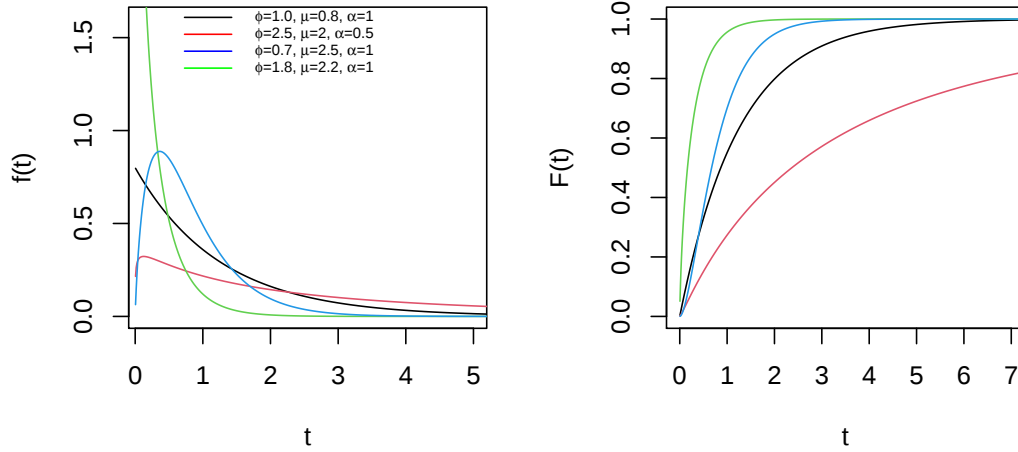


Figure 2.5: Graph of the probability density functions for different values of ϕ , λ and α and their respective cumulative density functions.

analyze the behavior of the risk function. According to Gi and Almhana [33] we will analyze such behaviors through the values of α and ϕ . If $\alpha = 1$, we have the distribution (2.7). With that, the risk function can have the following characteristics:

- increasing: if $0 < \phi < 1$
- Constant: if $\phi = 1$
- decreasing: if $\phi > 1$

If $\alpha \neq 1$ and if $\frac{(1 - \alpha\phi)}{\alpha(\alpha - 1)}$ assume values greater than zero, then so the risk function will take the following forms:

- bathtub: if $\phi > 1$
- inverse bathtub: if $0 < \phi < 1$

In the Figure 2.6 we can observe these characteristics of the risk function and survival function of the distribution GW with different values for the parameters (ϕ, λ, α) .

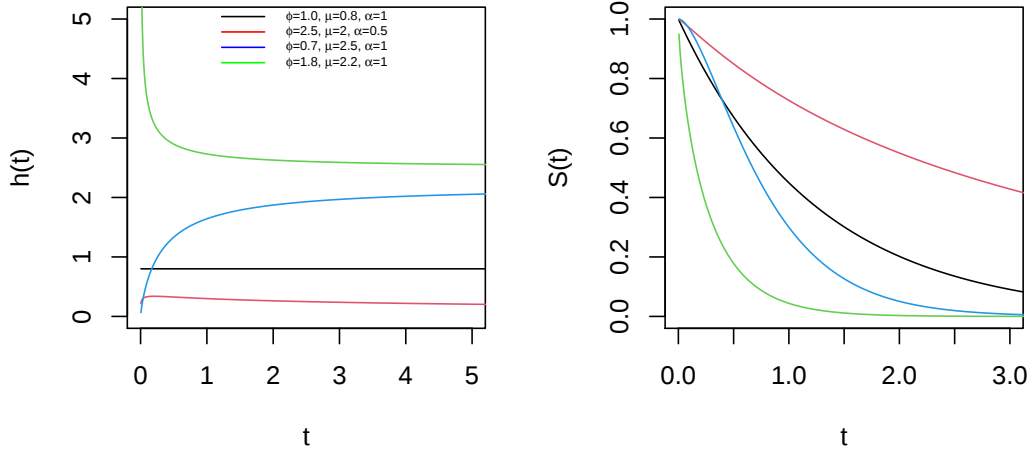


Figure 2.6: Risk function and survival function of the distribution GW with different values for the parameters.

2.3 Maximum likelihood inference

Commonly in statistical analysis, when we perform a study with the objective of making an inference about a certain phenomenon taking into consideration some degree of confidence, we usually employ two alternative methods: frequentist or Bayesian.

In the classical method, we consider the vector of $\boldsymbol{\theta}$ parameters as being an unknown constant. We can find different procedures for the inference of parameters in the literature such as the maximum likelihood estimator (MLE), methods of moments (MM), least squares (LQ), L-moments, percentile, maximum product of spacing among others. See in detail [41], [11], [4], [21]. Here we will use the concepts of maximum likelihood to obtain the estimates from the point of view of frequency due to its good asymptotic properties. This method is based on maximizing the likelihood function to reach the estimators for the parameters see Casela and Berger [8].

Thus, we will show the method of maximum likelihood in a generic way for the vector of parameters $\boldsymbol{\theta}$ under a sample $\mathbf{x}$ observed is given by

$$L(\boldsymbol{\theta}, \mathbf{x}) = \prod_{i=1}^n f(x_i | \boldsymbol{\theta}) \quad (2.23)$$

where $\boldsymbol{\theta} = (\theta_1, \theta_2, \dots, \theta_k)$ and that by hypothesis (2.23) represents all the information about the parameter $\boldsymbol{\theta}$ through the data of the experiment. Generally, from obtaining the likelihood equation, it can often be complicated to make algebraic manipulations through this equation, so one usually works with the logarithm of the function $L(\boldsymbol{\theta}, \mathbf{x})$ which is equivalent to the function given in (2.23). With this, assuming that the likelihood function is differentiable in each of its k parameters, then we can obtain the scoring functions that are given by the calculations

$$\frac{\partial \log(L(\boldsymbol{\theta}, \mathbf{x}))}{\partial \theta_i} = 0, \quad i=1, \dots, k. \quad (2.24)$$

Solving the system involving these equations, we find the maximum likelihood estimators for each of the θ_i . However, most of the time it can become an algebraically complex task, in these cases some numerical method is commonly adopted to obtain the solution.

It is important to note that for a small number of samples these estimators are addicted, but while we have a considerably large sample size the estimators obtained are consistent, efficient and converge to an asymptotically normal joint distribution given by the expression

$$\hat{\boldsymbol{\theta}} \sim N_k(\boldsymbol{\theta}, I^{-1}(\boldsymbol{\theta})) \text{ when } n \rightarrow \infty \quad (2.25)$$

where $I(\boldsymbol{\theta})$ is defined with the Fisher Information Matrix $k \times k$, whose elements $I_{ij}(\boldsymbol{\theta})$ of each entry are given by

$$I_{ij}(\boldsymbol{\theta}) = -E \left[\frac{\partial^2}{\partial \theta_i \partial \theta_j} L(\boldsymbol{\theta}, \mathbf{x}) \right], \text{ when } i, j=1, \dots, k. \quad (2.26)$$

Still under the hypothesis of large samples, it is possible to build the asymptotic confidence intervals individually, i.e., for each $\boldsymbol{\theta}_i$ with a $100(1-\gamma)\%$ confidence coefficient from the marginal distributions defined as follows

$$\hat{\boldsymbol{\theta}}_i \sim N(\boldsymbol{\theta}, I_{ii}^{-1}(\boldsymbol{\theta})) \text{ when } n \rightarrow \infty \quad (2.27)$$

With respect to the estimation via the maximum likelihood method, a property of extreme importance is the invariance principle of this model, i.e., given a random sample $X_1, \dots, X_n$ taken via the variable X with PDF $f(\mathbf{x}|\boldsymbol{\theta})$. If $\hat{\boldsymbol{\theta}}$ is a maximum likelihood estimator for $\boldsymbol{\theta}$, then $g(\hat{\boldsymbol{\theta}})$ will also be a maximum likelihood estimator for $g(\boldsymbol{\theta})$.

2.4 Bayesian Inference

The information related to the parameter of interest $\boldsymbol{\theta}$ plays a key role in statistics. Since the true value of $\boldsymbol{\theta}$ is unknown, our objective is to estimate such a parameter. The distribution that describes the information related to $\boldsymbol{\theta}$ before the data is observed is commonly known in Bayesian inference as a prior distribution $\pi(\boldsymbol{\theta})$. In the literature there are different ways to construct prior distributions, our focus here will be on priors that can be constructed with formal rules without and can be used in an objective inference [10]. Here, we will consider the use of reference priors that were proposed by Bernardo [6] with further development see Bernardo [7]. The class of priors is achieved by maximizing the expected Kullback-Leibler distance between the prior distribution and the posterior distribution, i.e., and the information achieved by the posterior distribution is not overshadowed by the prior. Additionally, these class of prior enjoys important properties such as consistency and invariance under one-to-one transformations see [7]. Similarly, let we also consider the matching prior suggested by Ramos et al. [39]. We will see its construction in the following.

This important objective prior was introduced by Tibshirani [57] assuming that for a prior distribution $\pi(\theta_1, \theta_2)$ where the parameter of interest is θ_1 and θ_2 is a nuisance parameter, the credibility interval for that parameter in has a convergence error $O(n^{-1})$ in the frequentist sense, that is,

$$P[\theta_1 < \theta_1^{1-\alpha}(\pi, X) | (\theta_1, \theta_2)] = 1 - \alpha - O(n^{-1}). \quad (2.28)$$

The class of priors that satisfies the condition above is usually referred to as matching priors with a convergence order of $O(n^{-1})$. Mukerjee and Dey [30] further discussed necessary and sufficient conditions for this class matching priors also be a matching prior but with a convergence of order $o(n^{-1})$.

The following proposition will be used to obtain the prior distributions for the Weibull and Generalized Weibull.

Proposição 1 [Bernardo [7], pg 40, Theorem 14]. Let $\boldsymbol{\theta} = (\theta_1, \dots, \theta_d)'$ be a vector of the ordered parameters of interest and consider $\mathbf{X} = (x_1, \dots, x_n)'$ a random sample of size n taken from a probability density function (PDF) $f(\cdot|\boldsymbol{\theta})$. Now be $\mathcal{P}$ the class that represents all a priors continuous with support in $\mathcal{A}$. In addition, be $I(\boldsymbol{\theta})$ Fisher Information Matrix. Finally, if the parameter space θ_j is independent of the parameter space θ_{-j} and the element $h_{jj}(\boldsymbol{\theta})$ ($j = 1, \dots, d$) of the fisher matrix can be factored through the form

$$h_{jj}^{\frac{1}{2}}(\boldsymbol{\theta}) = f_j(\theta_j)g_j(\theta_{-j})$$

where $h_{jj}(\boldsymbol{\theta})$ are the inputs (j, j) of the Fisher information matrix and the vector of parameters θ_{-j} is defined as $\theta_{-j} = (\theta_1, \dots, \theta_{j-1}, \theta_{j+1}, \dots, \theta_d)'$. Thus, a prior reference for the parameter $\boldsymbol{\theta}$ is given by

$$\pi_{\mathcal{R}} = \prod_{j=1}^d f_j(\theta_j)$$

Now, consider a random sample $x_1, x_2, \dots, x_n$ that were observed from the PDF $f(x|\boldsymbol{\theta})$, in this case, the likelihood function $L(\boldsymbol{\theta}|x) = \prod_{i=1}^n f(x_i|\boldsymbol{\theta})$ summarizes all the information of the parameter $\boldsymbol{\theta}$ based on observed sample data. Be it $\pi(\boldsymbol{\theta})$ a prior defined from $\boldsymbol{\theta}$, with the help of Bayes Theorem we combine the information about the a prior Distribution and the likelihood function we find the a posterior Distribution of $\boldsymbol{\theta}$ given x , known as $p(x|\boldsymbol{\theta})$. This result can be explained from the equality given below

$$p(\boldsymbol{\theta}|x) = \frac{p(\boldsymbol{\theta}, x)}{p(x)} = \frac{p(\boldsymbol{\theta}|x)p(\boldsymbol{\theta})}{p(x)} = \frac{p(\boldsymbol{\theta}|x)p(\boldsymbol{\theta})}{\int p(\boldsymbol{\theta}, x)d\boldsymbol{\theta}} = \frac{L(\boldsymbol{\theta}|x)\pi(\boldsymbol{\theta})}{\int L(\boldsymbol{\theta}|x)\pi(\boldsymbol{\theta})d\boldsymbol{\theta}}$$

It can be notice that $1/p(x)$ does not depend on the parameter $\boldsymbol{\theta}$ and it is defined as a normalizing constant of posterior distribution $p(\boldsymbol{\theta}|x)$. Since we are using objective priors that can be improper, i.e., $\int \pi(\boldsymbol{\theta}) = \infty$, the posterior distributions can also be improper, therefore, we have to prove such posteriors are proper functions to be used in Bayesian analysis.

In the case of proper posterior, our objective is to perform the estimation of the parameters, then, a very useful and simplified way to represent $p(\boldsymbol{\theta}|x)$ is using the following expression

$$p(\boldsymbol{\theta}|x) \propto L(\boldsymbol{\theta}|x)\pi(\boldsymbol{\theta})$$

There are numerous ways to summarize the information from the posterior distribution, usually, this step requires the calculation of probabilities and expectation on marginal distributions when $\boldsymbol{\theta}$ is multidimensional. Note that, usually we are not able to calculate such integrals in an analytical way due to their complexity, in these cases, we resort to Markov Chain Monte Carlo methods that will be presented after the construction of the posterior distributions.

Gamma Distribution

3.1 Introduction

An important non-normal distribution that has been widely considered in reliability is the Gamma distribution. The inference for the parameter of the model are standard in most statistical books. Under the classical approach the maximum likelihood estimator (MLE) is usually considered due to its good properties such as consistency, efficiency and invariant under one-to-one transformations. On the other hand, this method has two drawbacks, the first is that the method is usually biased for small samples, and secondly due to the mathematical form of the capability index it is very difficult to apply the delta method to obtain the confidence intervals of the estimates. In order to overcome these obstacles we can apply the objective Bayesian methods. Previously, several researchers discussed procedures to conduct Bayesian inference for the parameters of the Gamma distribution (see Louzada and Ramos [23], Miller [28], Sun and Ye. [55] and Berger et al. [5]). Ramos et al. [39] compared the effect of the different objective priors have in the posterior distribution and concluded that in terms of bias, mean square error and coverage probabilities the matching prior was the most appropriate.

Here, we consider a fully objective Bayesian analysis to estimate the C_{pk} measure. Due to the complex form of the index, we obtain the posterior distribution using the invariant property of the Markov Chain Monte Carlo (MCMC) methods. The Metropolis-Hasting algorithm is considered to obtain the chains of the parameters as well as the C_{pk} . Further, we show that the results obtained from objective inference are superior to the ones obtained through maximum likelihood. From the obtained results, we can achieve Bayes estimates with smaller bias and construct precise credibility intervals from the posterior density. The proposed approach is applied to check if the control process of powdered juice company is under control or if requires improvements to decrease losses.

The remainder of this chapter is organized as follows. Section (3.2) provides the maximum likelihood estimates. Section (3.3) and (3.4) presents the steps to achieve the Bayesian inference assuming an objective prior and the MCMC algorithm to sample from the posterior distribution. In Section (3.5), we present a simulation study to illustrate the usefulness of our proposed approach.

3.2 Maximum Likelihood Estimation

The maximum likelihood function is one of the most commonly used estimation procedures in frequentist inference, the method is usually considered due to its flexibility in construct the likelihood function, that contains all the information of θ obtained from the data. The MLEs have important properties such as invariance, consistency, and efficiency. Moreover, due to its asymptotic properties, we may be able to construct confidence intervals for the parameters of interest. Let $X_1, \dots, X_n$ a random sample of size n of the random variable give in (2.7) then the α and β likelihood function is give by

$$L(\alpha, \beta | \mathbf{x}) = \frac{\beta^{n\alpha}}{[\Gamma(\alpha)]^n} \left\{ \prod_{i=1}^n x_i^{\alpha-1} \right\} \exp \left\{ -\beta \sum_{i=1}^n x_i \right\}. \quad (3.1)$$

Given that the maximum likelihood of θ , $\theta = \{(\alpha, \beta); \alpha > 0, \beta > 0\}$, which is represented by $\hat{\theta}$ that maximize $L(\alpha, \beta | \mathbf{x})$, it may be difficult to direct maximize $L(\cdot)$ and as the logarithm of the likelihood function $l(\alpha, \beta | \mathbf{x})$ is a growing function, we have that maximize $L(\alpha, \beta | \mathbf{x})$ is equivalent to maximizing its logarithm. Hence, in this case the logarithm of the likelihood function is given by

$$l(\alpha, \beta | \mathbf{x}) = n\alpha \log \beta - n \log \Gamma(\alpha) + (\alpha - 1) \sum_{i=1}^n \log x_i - \beta \sum_{i=1}^n x_i. \quad (3.2)$$

From

$$\frac{\partial \log L(\alpha, \beta | \mathbf{x})}{\partial \alpha} \quad \text{and} \quad \frac{\partial \log L(\alpha, \beta | \mathbf{x})}{\partial \beta},$$

the score functions of α and β are given by for

$$\frac{\partial \log L(\alpha, \beta | \mathbf{x})}{\partial \alpha} = n \log(\beta) - n \frac{\Gamma'(\hat{\alpha})}{\Gamma(\hat{\alpha})} + \sum_{i=1}^n \log(x_i) = 0 \quad (3.3)$$

$$\frac{\partial \log L(\alpha, \beta | \mathbf{x})}{\partial \beta} = \frac{n\alpha}{\hat{\beta}} - \sum_{i=1}^n x_i = 0. \quad (3.4)$$

Performing some changes algebraically in (3.3) and (3.4), we find the maximum likelihood estimators of α and β given, respectively, by

$$\log(\hat{\alpha}) - \psi(\hat{\alpha}) = \log(\bar{X}) - \frac{1}{n} \sum_{i=1}^n \log(x_i) \quad (3.5)$$

$$\hat{\beta} = \frac{\hat{\alpha}}{\bar{X}} \quad (3.6)$$

where $\psi(\hat{\alpha}) = \frac{\partial \log \Gamma(\alpha)}{\partial \alpha} = \frac{\Gamma'(\hat{\alpha})}{\Gamma(\hat{\alpha})}$ and $\bar{X} = \frac{\sum_{i=1}^n x_i}{n}$.

There are situations, mainly, in the case of more complex models, the maximum likelihood estimators does not have closed-form expressions. In this case, the maximum likelihood estimators must be solved using numerical methods. There is also the possibility of making some modifications to the MLE that does not have closed-form expression that leads to estimators with the closed form [22]. Through an important result that guarantees that if the derivatives up to the second order of the likelihood function are defined and the Fisher information matrix is not null, then the distribution of $\hat{\alpha}$ and $\hat{\beta}$ tends to a

normal distribution with mean α , β and equal variance to the elements I_{11} and I_{22} of the matrix $I(\alpha, \beta)^{-1}$, respectively. The Fisher information matrix is given by

$$I(\alpha, \beta) = \begin{bmatrix} \psi'(\alpha) & -\frac{1}{\beta} \\ -\frac{1}{\beta} & \frac{\alpha}{\beta^2} \end{bmatrix},$$

where $\psi'(\hat{\alpha}) = \frac{\partial^2 \log \Gamma(\alpha)}{\partial \alpha^2}$ is known as the trigamma function.

3.3 Bayesian Inference for Gamma Distribution

For the Gamma distribution, Sun and Ye [55] showed that the matching prior of order $o(n^{-1})$ when β is the parameter of interest is given by

$$\pi(\alpha, \beta) \propto \frac{\alpha \psi'(\alpha) - 1}{\beta \sqrt{\alpha}}.$$

As discussed earlier this prior is one of the many possible objective priors for the parameters of the Gamma distribution. On the other hand, Ramos et al. [39] showed that the cited prior returned a posterior distribution that provided excellent posterior estimates in terms of smaller bias and MSE. Moreover, the obtained posterior has precise coverage probabilities for both parameters and important properties such as invariance under one-to-one transformation, hence, we will consider such prior to performing the inference for the C_{pk} index.

Through the matching prior we obtain the posterior joint distribution for α and β .

$$\pi(\alpha, \beta | \mathbf{x}) \propto \frac{(\alpha \psi'(\alpha) - 1)}{\sqrt{\alpha}} \frac{\beta^{n\alpha-1}}{[\Gamma(\alpha)]^n} \left\{ \prod_{i=1}^n x_i^{\alpha-1} \right\} \exp \left\{ -\beta \sum_{i=1}^n x_i \right\}. \quad (3.7)$$

The matching prior is improper and may lead to improper posterior. However, Ramos et al. [39] proved that the posterior distribution (3.7) is proper if and only if $n \geq 2$ with $n \in \mathbb{N}$ as well as its higher moments (posterior mean, variance among others). Therefore we can use such posterior to conduct the Bayesian inference.

By integrating (3.7) with respect to the parameter β we find the later marginal distribution for α which is given by

$$\pi(\alpha | \mathbf{x}) \propto \frac{(\alpha \psi'(\alpha) - 1)}{\sqrt{\alpha}} \frac{\Gamma(n\alpha)}{\Gamma(\alpha)^n} \left(\frac{\sqrt[n]{\prod_{i=1}^n x_i}}{\sum_{i=1}^n x_i} \right)^{n\alpha}. \quad (3.8)$$

The conditional posterior distribution for β is

$$\pi(\beta | \alpha, \mathbf{x}) \sim \text{Gamma} \left(n\alpha, \sum_{i=1}^n x_i \right). \quad (3.9)$$

3.4 Markov Chain Monte Carlo method

Extracting information from the posterior distributions is usually not an easy task. The joint posterior distribution follows an unknown model and the normalized constant is difficult to be computed. In order to extract the information of interested we can considered the Metropolis-Hastings (MH) algorithm to sample from the posterior distribution. The algorithm is used as follows:

1. Select an initial value for $\alpha^{(1)}$ and take $j = 1$ as the loop counter iteration.
2. Generate a candidate $\alpha^{(*)}$ from the proposed distribution $Gamma(\alpha^{(j)}, b)$;
3. Calculate the probability of acceptance given by

$$\lambda(\alpha^{(j)}, \alpha^{(*)}) = \min \left(1, \frac{p(\alpha^{(*)}|\mathbf{x})}{p(\alpha^{(j)}|\mathbf{x})} \frac{q(\alpha^{(j)}, \alpha^{(*)}, b)}{q(\alpha^{(*)}, \alpha^{(j)}, b)} \right),$$

where $p(\cdot)$ is the marginal posterior distribution given in (3.8). Select a random value u from the uniform distribution in the interval of $(0, 1)$;

4. If $\lambda(\alpha^{(j)}, \alpha^{(*)}) \geq u(0, 1)$ then accept the new value $\alpha^{(j+1)} = \alpha^{(*)}$, otherwise reject it and consider $\alpha^{(j+1)} = \alpha^{(j)}$;
5. Generate a random value $\beta^{(j+1)}$ from the conditional posterior distribution for β given by

$$\beta^{(j+1)} \sim Gamma \left(n\alpha^{(j+1)}, \sum_{i=1}^n x_i \right).$$

6. Let

$$C_{pu} = \frac{USL - M(\alpha^{(j+1)}, \beta^{(j+1)})}{U_p(\alpha^{(j+1)}, \beta^{(j+1)}) - M(\alpha^{(j+1)}, \beta^{(j+1)})}$$

$$C_{pl} = \frac{M(\alpha^{(j+1)}, \beta^{(j+1)}) - LSL}{M(\alpha^{(j+1)}, \beta^{(j+1)}) - L_p(\alpha^{(j+1)}, \beta^{(j+1)})},$$

where $U_p(\alpha^{(j+1)}, \beta^{(j+1)})$, $L_p(\alpha^{(j+1)}, \beta^{(j+1)})$ e $M(\alpha^{(j+1)}, \beta^{(j+1)})$ are the percentiles given respectively by 99.865, 0.135 e 50 obtained from the gamma distribution.

7. Compute

$$C_{pk}^{(j+1)} = \min(C_{pu}, C_{pl}).$$

8. Increase the counter from j to $j + 1$. The process must be repeated until a specified number of iterations is achieved.

It is important to mention that when using marginal posterior distribution of α , there is significant difficulty in calculating the ratio

$$\frac{p(\alpha^{(*)}|\alpha^{(t)}, \mathbf{x})}{p(\alpha^{(t)}|\alpha^{(*)}, \mathbf{x})},$$

this impasse can be explained by the fact that $p(\alpha^{(*)}|\alpha^{(t)}, \mathbf{x})$ has the terms $\Gamma(n\alpha)$ and $(\Gamma(\alpha)^n)^{-1}$ and depending on the values of ϕ and n the software may not compute such terms. For instance, consider $\alpha = 3$ and assuming that $n \geq 100$, the R will not be able to calculate such terms, returning the symbol ∞ in the first case and 0 for the second case. This difficulty can be resolved by rewriting the way in which we calculate the probability of acceptance, given by

$$\lambda(\alpha^{(t)}, \alpha^{(*)}) = \min(1, e^{\Omega}),$$

where $\Omega = \log(p(\alpha^{(*)}|\alpha^{(t)}, \mathbf{x})) + \log(q(\alpha^{(t)}|\alpha^{(*)})) - \log(p(\alpha^{(t)}|\alpha^{(t)}, \mathbf{x})) - \log(q(\alpha^{(*)}|\alpha^{(t)}))$.

This method overcomes the problem, as it is easy to calculate the terms $n \log(\Gamma(\alpha))$ and $\log(\Gamma(n\alpha))$ by using the approximation

$$\log(\Gamma(x)) = x \log(x) - x - \frac{1}{2} \log\left(\frac{x}{12\pi}\right) + \frac{1}{12x} - \frac{1}{360x^3} + \frac{1}{1260x^5} + \dots$$

3.5 Simulation Study for Gamma Distribution

A simulation study was proposed to verify whether the posterior distribution presents good results when compared with the classical approach in terms of the mean relative errors (MRE) and the mean square errors (MSE), which are the two commonly used metrics to evaluate the parameter estimators. In practice, these metrics together allows the researcher to observe if the Bayes estimates obtained from the posterior distribution are, in general, close to the true parameters. The metrics are computed by

$$MRE = \frac{1}{N} \sum_{j=1}^N \frac{\hat{\theta}_j}{\theta} \quad \text{and} \quad MSE = \sum_{j=1}^N \frac{(\hat{\theta}_j - \theta)^2}{N},$$

where $\theta = C_{pk}$ and $N = 10.000$ is the number of estimates obtained through the posterior means of the parameters. Additionally, the 95% coverage probability ($CP_{95\%}$) of the credibility intervals are computed from the Bayesian credibility intervals (CI). Under

Table 3.1: MRE, MSE and CP from the estimates of C_{pk} considering different values of n with $N = 10.000$ simulated samples using Classical and Bayesian inference.

θ	n	MLE		Bayes		
		MRE	MSE	MRE	MSE	CP
$\alpha = 2, \beta = 0.5$	10	1.1636	0.0439	0.9742	0.0343	0.964
	20	1.1237	0.0272	1.0384	0.0185	0.953
	30	1.0821	0.0168	1.0305	0.0131	0.950
	40	1.0583	0.0118	1.0219	0.0099	0.951
	50	1.0470	0.0088	1.0191	0.0077	0.954
	60	1.0398	0.0072	1.0170	0.0064	0.952
	70	1.0321	0.0060	1.0128	0.0055	0.952
	80	1.0286	0.0053	1.0118	0.0049	0.947
	90	1.0284	0.0047	1.0135	0.0044	0.945
	100	1.0238	0.0041	1.0105	0.0039	0.948
	110	1.0198	0.0037	1.0078	0.0035	0.948
	120	1.0167	0.0033	1.0057	0.0031	0.951
	130	1.0160	0.0031	1.0058	0.0029	0.946
	140	1.0181	0.0029	1.0086	0.0028	0.948
	150	1.0160	0.0027	1.0073	0.0026	0.946
$\alpha = 2, \beta = 1$	10	1.0231	0.0039	0.9471	0.0097	0.950
	20	1.0175	0.0023	0.9943	0.0025	0.950
	30	1.0109	0.0011	1.0004	0.0011	0.953
	40	1.0081	0.0007	1.0021	0.0007	0.950
	50	1.0066	0.0005	1.0026	0.0005	0.947
	60	1.0052	0.0004	1.0023	0.0004	0.948
	70	1.0046	0.0004	1.0023	0.0003	0.946
	80	1.0042	0.0003	1.0022	0.0003	0.950
	90	1.0034	0.0003	1.0017	0.0002	0.948
	100	1.0030	0.0002	1.0015	0.0002	0.945
	110	1.0026	0.0002	1.0013	0.0002	0.949
	120	1.0026	0.0002	1.0014	0.0002	0.946
	130	1.0024	0.0002	1.0012	0.0002	0.950
	140	1.0021	0.0001	1.0010	0.0001	0.947
	150	1.0020	0.0001	1.0011	0.0001	0.944

this approach, the estimator is expected to return MREs and MSEs values closer to one and zero, respectively. In addition, considering a credibility level of 95% the ranges achieved are expected to cover the true values of the θ with a proportion of 0.95. The simulation was done with using the R software. An implemented function to compute the C_{pk} under the Bayesian approach is available in the supplemental material. Considering $n = (10, 20, 30, \dots, 150)$ the results were presented assuming different values of the C_{pk} .

Here, we considered four scenarios assuming different values for the parameters α , β , and the specification limits USL and LSL . For scenario 1 we assume $\alpha = 2$, $\beta = 0.5$, $USL = 10$, and $LSL = 0.5$ with the true $C_{pk} = 0.4599$. In the second scenario we use $\alpha = 2$, $\beta = 1$, $USL = 14.5$ and $LSL = 0.1$ and find the real $C_{pk} = 0.9710$. Similarly, we analyzed the third scenario with the respective values $\alpha = 1.1$, $\beta = 0.2$, $USL = 10$ and $LSL = 0.1$ given the true $C_{pk} = 0.1992$. Finally, in the fourth scenario we assume $\alpha = 7$, $\beta = 1.2$, $USL = 25$ and $LSL = 0.01$ and the real $C_{pk} = 1.3140$.

Table 3.2: MRE, MSE and Coverage Probability from the estimates of C_{pk} considering different values of n with $N = 10.000$ simulated samples using Classical and Bayesian inference.

θ	n	MLE		Bayes		
		MRE	MSE	MRE	MSE	CP
$\alpha = 1.1, \beta = 0.2$	10	1.3253	0.0299	1.0753	1.3957	0.956
	20	1.1497	0.0097	1.0895	0.0079	0.952
	30	1.0957	0.0055	1.0585	0.0048	0.948
	40	1.0692	0.0037	1.0425	0.0033	0.948
	50	1.0549	0.0028	1.0338	0.0026	0.951
	60	1.0455	0.0023	1.0283	0.0021	0.946
	70	1.0351	0.0019	1.0206	0.0018	0.948
	80	1.0333	0.0016	1.0205	0.0015	0.946
	90	1.0307	0.0014	1.0195	0.0014	0.948
	100	1.0234	0.0012	1.0134	0.0012	0.948
	110	1.0264	0.0012	1.0173	0.0011	0.945
	120	1.0219	0.0011	1.0136	0.0010	0.945
	130	1.0237	0.0010	1.0160	0.0009	0.947
	140	1.0192	0.0009	1.0122	0.0008	0.950
	150	1.0161	0.0008	1.0095	0.0008	0.947
$\alpha = 7, \beta = 1.2$	10	0.9450	0.0100	0.8520	0.0450	0.898
	20	1.0125	0.0094	0.9754	0.0096	0.960
	30	1.0179	0.0092	0.9949	0.0079	0.962
	40	1.0187	0.0083	1.0018	0.0070	0.953
	50	1.0167	0.0067	1.0033	0.0058	0.949
	60	1.0124	0.0054	1.0014	0.0049	0.946
	70	1.0102	0.0043	1.0009	0.0039	0.950
	80	1.0087	0.0038	1.0006	0.0035	0.950
	90	1.0089	0.0034	1.0016	0.0031	0.949
	100	1.0082	0.0031	1.0016	0.0028	0.948
	110	1.0078	0.0028	1.0019	0.0026	0.947
	120	1.0066	0.0024	1.0011	0.0023	0.950
	130	1.0054	0.0023	1.0004	0.0021	0.948
	140	1.0054	0.0021	1.0008	0.0020	0.946
	150	1.0046	0.0019	1.0002	0.0018	0.951

The MCMC approach was used to find the estimates a posterior distribution. To achieve that, we considered the Metropolis-Hastings algorithm that was described in detail in Section 3.4. For each simulated data we obtain two chains of size 5.500, the first 500 samples were discarded as a burn-in sample. As elements generated in the chain has usually a dependence between, we considered the thin 5 (jump between the elements) to decrease the autocorrelation. Further, we considered a convergence diagnostic criteria proposed by Geweke [13] to confirm the convergences of chains generated via MCMC assuming a 95% confidence level. The samples produced were used to estimate the posterior mean from the marginal distribution of C_{pk} . Here, we considered the posterior mean since the posterior moments are finite for $n \geq 2$ and have optimality under the Kullback-Leibler divergence. Tables 3.1 and 3.2 present the MREs, MSEs and $CP_{95\%}$ from the different estimators of C_{pk} .

As can be seen from the results above, the results obtained from the posterior distribution using the matching prior are more precise than the ones achieved through maximum likelihood in terms of MREs and MSEs. Under the classical approach, we were not able to construct the confidence intervals due to the difficulty to apply the delta method. On the other hand, we constructed the credibility intervals from the Bayesian estimates directly from its percentiles which returned accurate results, this is expected as the matching prior has frequentist coverage close to the nominal level. Overall, we conclude that the posterior distribution obtained with matching prior should be used to estimate the C_{pk} measure related to the Gamma distribution.

Weibull distribution

4.1 Introduction

The development of new technologies provided an industrial expansion in general, where competitive environment corroborates the need to produce high-quality products and services. In this sense, metrics have been proposed to measure in an objective way quality control specifications [62]. Evaluating if a process capability index (PCI) is in accordance with the specifications play a key role for companies optimize their production aiming to reduce eventual losses, for more details see [51]. The literature provides several capacity indexes of the process each one with its characteristics. These indexes are used widely to verify if generating items meet the desirable specifications by measuring the inherent process variability. Recent studies that focus on PCI can be seen in Kots and Johnson [18], Spiring et al [51], Sappakitkamjorn and Niwitpong [45], Yum and Kim [61], Piña Monarrez et al [34]. Zang [63] verified that the two indexes most used in capacity studies are part of the most varied indexes of process capacity available in the literature.

Process capability indexes were initially proposed assuming that the data is normally distributed, i.e., the underlying process that describes the data follows a Gaussian model. However, in many situations the data does not satisfy this condition, to overcome such problem, modified indexes were proposed to deal with asymmetric data [47, 34, 49, 48]. Clements [9] considered the use of non-normal data in process capability calculations. Hosseinifard et al. [16] discussed different non-normal process capability estimation methods for monitoring the capability of the leukocyte filtering process. Recent studies have focused on a wide range of applications of PCI using asymmetric models [17, 32, 20, 47].

In the present study, we overcome such problems by considering Bayesian methods, in this case, fully objective methods can be considered to construct posterior distributions that return precise estimates and credibility intervals. In the previous chapter we provided an objective Bayesian inference for the Gamma distribution, using formal rules, consider an objective prior that is minimally informative and has interesting properties, such as invariance under one-to-one transformations and accurate credibility intervals. From the constructed posterior distributions we derived Bayes estimators that are very precise in terms of mean square error (MSE), mean relative error (MRE), and accurate coverage probabilities. Although such a study has been considered for the Gamma model, no such analysis has been derived for more applied models such as the Weibull distribution. Therefore, here we will derive Bayes estimators for PCI of the Weibull distribution assuming objective priors, the obtained posterior distribution returned accurate results and is recommended to be used to estimate the PCI for the considered model.

This chapter is organized as follows: In Sections (4.2) discuss Bayesian methods for the parameters of the Weibull distribution and the steps to find the Bayes estimates using Markov Chain Monte Carlo (MCMC) methods. Section (4.3) present a simulation for estimating the C_{pk} of the distribution Weibull.

4.2 Bayesian Inference for the Weibull distribution

Assuming that $T_1, \dots, T_n$ are random variables where $T \sim W(\lambda, \alpha)$. Then the likelihood function is given by

$$L(\lambda, \alpha | t) = \lambda^n \alpha^n \prod_{i=1}^n t_i^{\alpha-1} \exp \left\{ -\lambda \sum_{i=1}^n t_i^\alpha \right\}. \quad (4.1)$$

As stated above the parameters are assumed as random variables and a prior must be assigned to conduct Bayesian inference. Based on the proposition 1, we will enunciate a theorem and demonstrate a theorem with the objective of finding the priors regarding the order of the Weibull distribution parameters. The most used prior for this model is known as reference prior [54] and is given by

$$\pi(\lambda, \alpha) \propto \frac{1}{\lambda \alpha}. \quad (4.2)$$

Using the prior distribution (4.2) and the information contained in the likelihood function (4.1) we obtain the posterior distribution for λ and α that is given by

$$\pi(\lambda, \alpha | \mathbf{t}) = \frac{(\lambda \alpha)^{n-1}}{c(\mathbf{t})} \prod_{i=1}^n t_i^{\alpha-1} \exp \left\{ -\lambda \sum_{i=1}^n t_i^\alpha \right\}, \quad (4.3)$$

where $c(\mathbf{t})$ is a normalizing constant

$$c(\mathbf{t}) = \int_{\mathcal{A}} (\lambda \alpha)^{n-1} \prod_{i=1}^n t_i^{\alpha-1} \exp \left\{ -\lambda \sum_{i=1}^n t_i^\alpha \right\} d\boldsymbol{\theta} \quad (4.4)$$

and $\mathcal{A} = \{(\iota, \infty) \times (\iota, \infty)\}$ it is defined as the $\boldsymbol{\theta}$ support. To get reliable inference, first we have to check whether the posterior distribution (4.3) is a proper posterior, i.e., $c(\mathbf{t}) < \infty$. Ramos et al. [38] proved that for samples of size $n \geq 2$ the posterior given in (4.3) is proper, therefore, we can use such posterior to obtain Bayes estimator.

The marginal posterior distribution for α is given by

$$\pi(\alpha | \mathbf{t}) \propto \alpha^{n-1} \prod_{i=1}^n t_i^{\alpha-1} \left\{ \sum_{i=1}^n t_i^\alpha \right\}^{-n}. \quad (4.5)$$

The conditional posterior distribution for λ is

$$\pi(\lambda | \alpha, \mathbf{t}) \sim \text{Gamma} \left(n, \sum_{i=1}^n t_i^\alpha \right), \quad (4.6)$$

where the gamma distribution is defined as Gamma and its PDF is given by $f(\alpha, \beta | x) = \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\beta x}$.

4.2.1 MCMC methods for Weibull Distribution

The marginal distributions (4.5) and (4.6) does not have closed-form expressions. Hence, it will be the Metropolis-Hastings (MH) algorithm to generate samples of α from the marginal posterior distribution. To sample from the marginal distribution of λ we can use the conditional distribution given in (4.6). Here, we considered the Gamma distribution $q(\alpha^{(j)}|\alpha^{(*)}, b)$ distribution as a proposal distribution to sample values of the parameter α , where b is defined as a hyperparameter that control the convergence rate of the algorithm. It is worth mentioning that other proposal distributions can be assumed instead of the Gamma model. The M-H algorithm is employed as follows

1. Introduce a value $\alpha^{(1)}$ to start and set a counter $j = 1$;
2. Create a random candidate from the proposed distribution $Gamma(\alpha^{(j)}, b)$;
3. Analyze the probability of acceptance defined by

$$\Delta(\alpha^{(j)}, \alpha^{(*)}) = \min \left(1, \frac{p(\alpha^{(*)}|\mathbf{t}) q(\alpha^{(j)}, \alpha^{(*)}|b)}{p(\alpha^{(j)}|\mathbf{t}) q(\alpha^{(*)}, \alpha^{(j)}|b)} \right),$$

where $p(\cdot)$ is the posterior conditional α distribution given through (4.5). Generates a random sample from an independent u (uniform) distribution in the interval $(0,1)$;

4. Analyze if $\Delta(\alpha^{(j)}, \alpha^{(*)}) \geq u(0, 1)$ then the $\alpha^{(*)}$ value is accepted and then $\alpha^{(j+1)} = \alpha^{(*)}$, but if $\Delta(\alpha^{(j)}, \alpha^{(*)}) < u(0, 1)$ then reject and define $\alpha^{(j+1)} = \alpha^{(j)}$;
5. From the later conditional distribution of λ defined in (4.6) generates a random value $\lambda^{(j+1)}$ through

$$\lambda^{(j+1)} \sim Gamma \left(n, \sum_{i=1}^n t_i^{\alpha^{(j+1)}} \right);$$

6. Consider

$$C_{pl} = \frac{USL - M(\alpha^{(j+1)}, \lambda^{(j+1)})}{U_p(\alpha^{(j+1)}, \lambda^{(j+1)}) - M(\alpha^{(j+1)}, \lambda^{(j+1)})}$$

$$C_{pu} = \frac{M(\alpha^{(j+1)}, \lambda^{(j+1)}) - LSL}{M(\alpha^{(j+1)}, \lambda^{(j+1)}) - L_p(\alpha^{(j+1)}, \lambda^{(j+1)})},$$

where $U_p(\alpha^{(j+1)}, \lambda^{(j+1)})$, $L_p(\alpha^{(j+1)}, \lambda^{(j+1)})$ e $M(\alpha^{(j+1)}, \lambda^{(j+1)})$ are defined as the 99.865, 0.135 and 50 percentiles and are calculated from the Weibull distribution, respectively;

7. Calculate

$$C_{pk}^{(j+1)} = \min(C_{pl}, C_{pu}); \quad (4.7)$$

8. Update the (j) counter to (j+1) and remake the steps until the chains converge.

4.3 Simulation Study for Weibull distribution

In this section, we present the results of a simulation study aiming to evaluate the performance of the posterior distributions for the parameters of the Weibull model. For the comparison, three metrics were considered to evaluate the Bayes estimators, the mean relative error (MRE) and the mean quadratic error (MSE) that are defined by

$$MRE(C_{pk}) = \frac{1}{N} \sum_{i=1}^N \frac{\hat{C}_{pk_i}}{C_{pk_i}} \quad \text{and} \quad MSE(C_{pk}) = \sum_{i=1}^N \frac{(\hat{C}_{pk_i} - C_{pk_i})^2}{N},$$

where N is the number of estimates from different samples. The third metric is the coverage probability assuming 95% of confidence obtained from the credibility intervals. From this approach, it is expected that the proposed Bayes estimators return the MREs closer to 1 and the MSEs closer to 0. Moreover, considering a 95% confidence level, it is expected that the proportion of estimates contained inside of the credibility intervals are 0.95. The codes were constructed using the R software (R Core Development Team) and can be obtained upon request to the corresponding author.

The choice of a decision rule to obtain the Bayes estimates is a fundamental problem in Bayesian statistics. Usually, this is a challenge addressed by several authors. The most commonly used risk function in obtaining Bayes estimates is the mean square error, considering this risk function, the Bayesian estimates achieved are the mean posterior. Alternatively, we can consider the posterior mode which is also known as the maximum probability a posterior (MAP) and is obtained through the loss function 0-1, while the median a posterior is obtained considering a linear loss function. Ramos et al. [38, 40] proved that the posterior means are not finite for some of the parameters of the Weibull and GW distributions. Hence, we consider the posterior median as Bayesian estimators since they are finite, these results are a direct corollary from the proof that posterior is proper.

For the present model, three scenarios assuming different values for the parameters and specifications were considered. The respective scenarios were considered: Scenario 1, $\lambda = 0.5$, $\alpha = 2$, $USL = 4$, $LSL = 1$ returning a $C_{pk} = 0.16$; Scenario 2, $\lambda = 0.5$, $\alpha = 2$, $USL = 4$, $LSL = 0.5$ where $C_{pk} = 0.60$; Scenario 3, $\lambda = 0.5$, $\alpha = 2$, $USL = 4$, $LSL = 0.05$ with $C_{pk} = 1.0$. For the simulation study, the sample sizes considered were $n = (50, 60, 70, \dots, 200)$ and $N = 10.000$ for the three scenarios.

Here, for each sample simulated we obtained the Bayes estimates for the three parameters in the case of the Weibull distribution, however, only the results for the C_{pk} are presented as it is the parameter of interest to be estimated. Therefore, three chains using the Metropolis-Hasting algorithm discussed in (4.2.1) were constructed containing 5.500 generated values, as burn-in 500 values were discarded and a thin of 5 was applied aiming to decrease the autocorrelation between the generated values, obtaining, in the end, a chain with 1.000 samples for C_{pk} . The convergence of the chains were confirmed using the Geweke diagnostic [13] assuming a confidence level of 95%. Figure 4.1 presents the MRE, MSE, and the CP for the C_{pk} obtained from the N posterior medians.

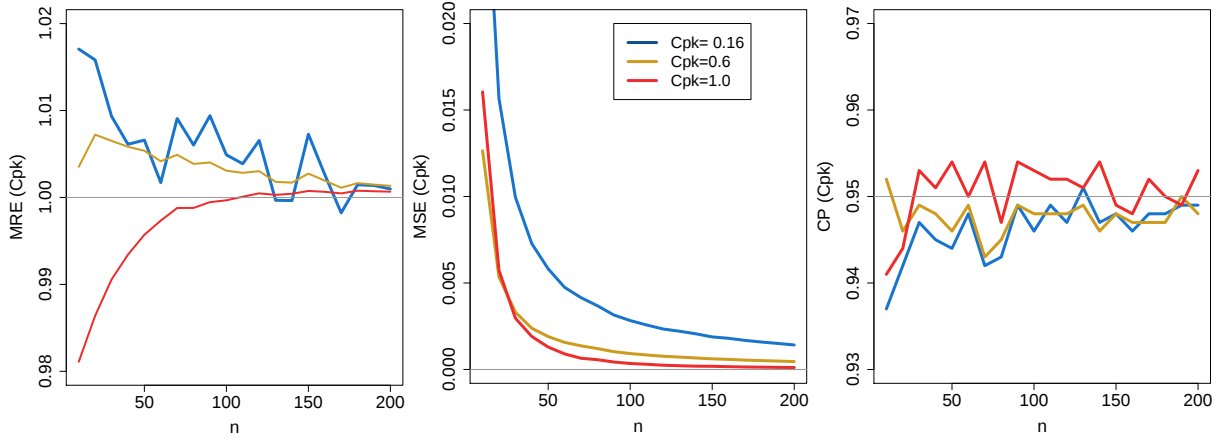


Figure 4.1: Simulation study for the Weibull distribution assuming different values for the C_{pk} , $N = 10.000$ and $n = 50, 60, \dots, 200$.

From the Figure 4.1 above we can see that, the Bayes estimators obtained from the sample of the reference posterior distribution (4.3) provides accurate estimates in terms of MREs, MSEs, additionally, the coverage probability of the C_{pk} are close to the nominal level, which shows that the proposed approach can return good estimates for the parameter of interested. Further readers are recommended to consider our proposed approach to estimate the C_{pk} when the Weibull distribution can be used to fit the data.

With these results obtained, we believe that the Bayesian estimators seen here, achieved with the prior reference, should be employed in the estimation of the C_{pk} index calculated from the Weibull model distribution.

Generalized Weibull distribution

This chapter presents the steps for estimating the C_{pk} index using the generalized Weibull distribution and is organized as follows: In section (5.1) we discuss Bayesian inference from the GW distribution using the prior reference to construct the posteriors and the step-by-step Bayesian estimates using (MCMC) methods. Finally, in section (5.2) a simulation study for estimating the capability index using the GW distribution is presented.

5.1 Bayesian Inference for the Generalized Weibull distribution

Assuming that $T_1, \dots, T_n$ are random variables $T \sim GW(\phi, \lambda, \alpha)$. Then The likelihood function that contains the information from the data is defined as

$$L(\phi, \lambda, \alpha | t) = \frac{\alpha^n}{\Gamma(\phi)^n} \lambda^{\alpha\phi n} \prod_{i=1}^n t_i^{\alpha\phi-1} \exp \left\{ -\lambda^\alpha \sum_{i=1}^n t_i^\alpha \right\}. \quad (5.1)$$

To conduct the Bayesian analysis a prior distribution must be assigned for the vector of parameters (ϕ, λ, α) . Ramos and Louzada [40] discussed different objective priors for the GW distribution assuming different parameters of interest. Assuming that $\boldsymbol{\theta} = (\phi, \lambda, \alpha)$ is the ordered vector of the parameters from the GW distribution. Then, the reference prior is given by

$$\pi(\boldsymbol{\theta}) \propto \frac{1}{\alpha\lambda} \sqrt{\frac{\phi^2 \psi'(\phi)^2 - \psi'(\phi) - 1}{\phi + \phi^2 \psi'(\phi) - 1}}, \quad (5.2)$$

where $\psi(k) = \frac{\partial}{\partial k} \log \Gamma(k) = \frac{\Gamma'(k)}{\Gamma(k)}$ is the digamma function and $\psi'(k) = \frac{\partial}{\partial k} \psi(k)$ is the trigamma function.

The joint posterior distribution of ϕ, λ, α using the prior reference (5.2) and the likelihood function (5.1) is given by

$$\pi(\boldsymbol{\theta} | t) = \frac{\pi(\phi)}{d(\mathbf{t})} \frac{\alpha^{n-1}}{\Gamma(\phi)^n} \lambda^{n\alpha\phi-1} \prod_{i=1}^n t_i^{\alpha\phi-1} \exp \left\{ -\lambda^\alpha \sum_{i=1}^n t_i^\alpha \right\}, \quad (5.3)$$

where $d(\mathbf{t})$ is defined as the normalizing constant and is given by

$$d(\mathbf{t}) = \int_{\mathcal{A}} \frac{\alpha^{n-1} \pi(\phi)}{\Gamma(\phi)^n} \lambda^{n\alpha\phi-1} \prod_{i=1}^n t_i^{\alpha\phi-1} \exp \left\{ -\lambda^\alpha \sum_{i=1}^n t_i^\alpha \right\} d\boldsymbol{\theta},$$

which $\pi(\phi) = \sqrt{\frac{\phi^2 \psi'(\phi)^2 - \psi'(\phi) - 1}{\phi + \phi^2 \psi'(\phi) - 1}}$ and $\mathcal{A} = \{(0, \infty) \times (0, \infty) \times (0, \infty)\}$ is the support for the parameter θ .

Notice that in order to carry out the Bayesian inference it is fundamental that the posterior distribution (5.3) is a proper distribution, the concern were discussed by Ramos and Louzada [40] that proved that the obtained posterior is proper, i.e., $d(\mathbf{t}) < \infty$. It is worth mentioning that it is possible to obtain 6 different reference priors depending on the ordered parameters of the GW distribution, however, only the presented reference prior returned a proper posterior.

Due to the consistent marginalization property of the reference prior, the reference marginal posterior distribution of ϕ and α is given by

$$\pi(\phi, \alpha | \mathbf{t}) \propto \alpha^{n-2} \frac{\pi(\phi) \Gamma(n\phi)}{\Gamma(\phi)^n} \left(\frac{\sqrt[n]{\prod_{i=1}^n t_i^\alpha}}{\sum_{i=1}^n t_i^\alpha} \right)^{n\phi}. \quad (5.4)$$

From the marginal posterior distribution (5.4) we find the conditional posterior distributions for α, λ and ϕ as follows

$$\pi(\alpha | \phi, \mathbf{t}) \propto \alpha^{n-2} \left(\frac{\sqrt[n]{\prod_{i=1}^n t_i^\alpha}}{\sum_{i=1}^n t_i^\alpha} \right)^{n\phi} \quad (5.5)$$

$$\pi(\phi | \alpha, \mathbf{t}) \propto \frac{\pi(\phi) \Gamma(n\phi)}{\Gamma(\phi)^n} \left(\frac{\sqrt[n]{\prod_{i=1}^n t_i^\alpha}}{\sum_{i=1}^n t_i^\alpha} \right)^{n\phi} \quad (5.6)$$

$$\pi(\lambda | \phi, \alpha, \mathbf{t}) \sim GW \left(n\phi, \left(\sum_{i=1}^n t_i^\alpha \right)^{\frac{1}{\alpha}}, \alpha \right), \quad (5.7)$$

where $GW(\phi, \lambda, \alpha)$ is the PDF of the GW distribution (2.17).

The conditional posterior distributions given in (5.5) and (5.6) are important expression to be used with MCMC method to sample from the marginal distribution of the parameters ϕ and α , respectively.

5.1.1 Markov Chain Monte Carlo Method for GW

Here, we follow similar steps as discussed for the Weibull distribution, however, a new parameter was included to be estimated which require an additional step in the M-H algorithm. Unlike the algorithm given in (4.2.1), here we use two proposed distributions $q_i(\theta_i^{(t)} | \theta_i^{(*)})$ to generate the values for the parameters (α, ϕ, λ) . Firstly, we use the auxiliary distribution defined through the Gamma distribution $q(\alpha^{(j)} | \alpha^{(*)}, b_1)$ and $q(\phi^{(j)} | \phi^{(*)}, b_2)$ to generate random samples for the α and ϕ marginal distributions, respectively. On the other hand, for the parameter λ we use the truncated Normal distribution in the interval $(0, \infty)$ with mean and standard deviation given by a_3 and b_3 and refereed to $Nt(a_3, b_3)$. The algorithm is executed as follows:

1. Start with the initial values $\theta^{(0)} = (\alpha^{(0)}, \phi^{(0)}, \lambda^{(0)})$ and set a counter $j = 1$;
2. Create a random candidate $\alpha^{(*)}$ from the proposed distribution $\text{Gamma}(\alpha^{(t)} \times b_1, b_1)$;

3. Compute the probability of acceptance defined by

$$\Delta_1(\alpha^{(j)}, \alpha^{(*)}) = \min \left(1, \frac{p_1(\alpha^{(*)}|\phi^{(j)}, \lambda^{(j)}, \mathbf{t}) q(\alpha^{(j)}, \alpha^{(*)}, b_1)}{p_1(\alpha^{(j)}|\phi^{(j)}, \lambda^{(j)}, \mathbf{t}) q(\alpha^{(*)}, \alpha^{(j)}, b_1)} \right)$$

where $p_1(\cdot)$ is the posterior conditional distribution defined in (5.5) and generates a random value from a uniform distribution u in interval (0,1);

4. If $\Delta_1(\alpha^{(j)}, \alpha^{(*)}) \geq u(0, 1)$ then the $\alpha^{(*)}$ value is accepted and $\alpha^{(j+1)} = \alpha^{(*)}$, but if $\Delta_1(\alpha^{(j)}, \alpha^{(*)}) < u(0, 1)$ then $\alpha^{(*)}$ is rejected and define $\alpha^{(j+1)} = \alpha^{(j)}$
5. Generate a random candidate $\phi^{(*)}$ from the proposed distribution $\text{Gamma}(\phi^{(t)} \times b_2, b_2)$;
6. Obtain the probability of acceptance defined by

$$\Delta_2(\phi^{(j)}, \phi^{(*)}) = \min \left(1, \frac{p_2(\phi^{(*)}|\alpha^{(j+1)}, \lambda^{(j)}, \mathbf{t}) q(\phi^{(j)}, \phi^{(*)}, b_2)}{p_2(\phi^{(j)}|\alpha^{(j+1)}, \lambda^{(j)}, \mathbf{t}) q(\phi^{(*)}, \phi^{(j)}, b_2)} \right)$$

where $p_2(\cdot)$ is the posterior conditional distribution defined in (5.6) and generates a random value from a uniform distribution u in interval (0,1);

7. If $\Delta_2(\phi^{(j)}, \phi^{(*)}) \geq u(0, 1)$ then the value $\phi^{(*)}$ is accepted and $\phi^{(j+1)} = \phi^{(*)}$, otherwise if $\Delta_2(\phi^{(j)}, \phi^{(*)}) < u(0, 1)$ then reject $\phi^{(*)}$ and define $\phi^{(j+1)} = \phi^{(j)}$;
8. Generate a random candidate $\lambda^{(*)}$ from the proposed distribution $\text{Nt}(a_3, b_3)$;
9. Obtain the probability of acceptance defined by

$$\Delta_3(\lambda^{(j)}, \lambda^{(*)}) = \min \left(1, \frac{p_3(\lambda^{(*)}|\alpha^{(j+1)}, \phi^{(j+1)}, \mathbf{t}) dNt(\lambda^{(j)}, a_3, b_3)}{p_3(\lambda^{(j)}|\alpha^{(j+1)}, \phi^{(j+1)}, \mathbf{t}) dNt(\lambda^{(*)}, a_3, b_3)} \right)$$

where $p_3(\cdot)$ is the posterior conditional distribution give in (5.7) and generates a random value from a uniform distribution u in interval (0,1);

10. If $\Delta_3(\lambda^{(j)}, \lambda^{(*)}) \geq u(0, 1)$ then $\lambda^{(*)}$ is accepted and $\lambda^{(j+1)} = \lambda^{(*)}$, but if $\Delta_3(\lambda^{(j)}, \lambda^{(*)}) < u(0, 1)$ then reject $\lambda^{(*)}$ and define $\lambda^{(j+1)} = \lambda^{(j)}$;
11. Consider

$$C_{pl} = \frac{USL - M(\phi^{(j+1)}, \lambda^{(j+1)}, \alpha^{(j+1)})}{U_p(\phi^{(j+1)}, \lambda^{(j+1)}, \alpha^{(j+1)}) - M(\phi^{(j+1)}, \lambda^{(j+1)}, \alpha^{(j+1)})}$$

$$C_{pu} = \frac{M(\phi^{(j+1)}, \lambda^{(j+1)}, \alpha^{(j+1)}) - LSL}{M(\phi^{(j+1)}, \lambda^{(j+1)}, \alpha^{(j+1)}) - L_p(\phi^{(j+1)}, \lambda^{(j+1)}, \alpha^{(j+1)})},$$

where $U_p(\phi^{(j+1)}, \lambda^{(j+1)}, \alpha^{(j+1)})$, $L_p(\phi^{(j+1)}, \lambda^{(j+1)}, \alpha^{(j+1)})$ and $M(\phi^{(j+1)}, \lambda^{(j+1)}, \alpha^{(j+1)})$ are defined as the 99.865, 0.135 and 50 percentiles, and are calculated from the generalized Weibull distribution, respectively;

12. Calculate

$$C_{pk}^{(j+1)} = \min(C_{pu}, C_{pl}); \quad (5.8)$$

13. Update the j counter to $j+1$ and repeat the steps until the converge is achieved.

5.2 Simulation Study Generalized Weibull distribution

In this section all simulation was performed following the steps of the Gamma and Weibull simulations presented in the previous sections. Here, we considered three scenarios under different specifications for the C_{pk} , $N = 31.000$ and sample sizes with $n = (50, 60, 70, \dots, 300)$ for the three scenarios. The first scenario, we assumed $\phi = 0.5$, $\mu = 1.5$, $\alpha = 4$, $USL = 4$ and $LSL = 0.3$ with $C_{pk}=0.37$. The second scenario consider that $\phi = 0.5$, $\mu = 1.5$, $\alpha = 4$, $USL = 4$ and $LSL = 0.1$ with $C_{pk} = 0.82$. Finally, the third scenario assumes that $\phi = 0.5$, $\mu = 1.5$, $\alpha = 4$, $USL = 4$ and $LSL = 0.01$ with $C_{pk} = 1.03$.

Different from the Weibull distribution, more iterations were necessary to construct the Bayes estimates due to the significant autocorrelation between the generated values. Hence, during the application of the Metropolis-Hasting algorithm 31.000 iterations were performed for the four chains in each simulated sample. The first 1.000 observations were discarded as a burn-in sample, the thin considered was 30 to reduce the autocorrelation between the chains. In the end, we achieve 1.000 proposed values of the marginal posterior distributions. The convergence of the chains was also confirmed using the Geweke diagnostic assuming a confidence level of 95%. Figure 5.1 provides the MREs, MSEs, achieved from the posterior medians and the coverage probabilities obtained from the credibility intervals with 95% credibility level.

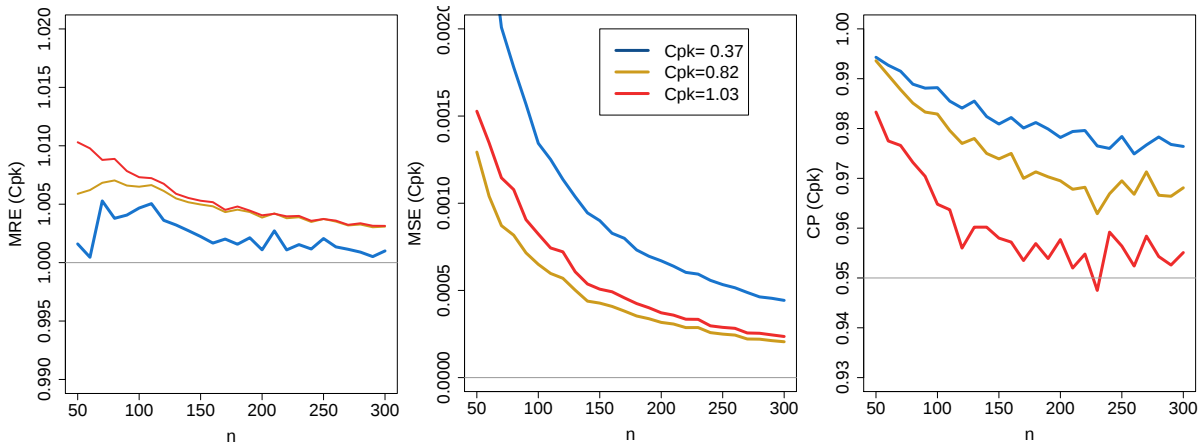


Figure 5.1: Simulation study for the Generalized Weibull distribution assuming different values for the C_{pk} , $N = 10.000$ and $n = 50, 60, \dots, 300$.

As shown in Figure 5.1, we observed the reference posterior distribution (5.3) obtained for the generalized model also provides accurate estimates in terms of MREs, MSEs, as well as the coverage probability of the C_{pk} are close to the nominal level, proving that our approach provides good estimates for the C_{pk} parameter. From these results, we conclude that the proposed Bayes estimators obtained from the reference priors should be used to estimate the process capability index calculated from asymmetric model such as GW distribution.

Observe that in chapter 3, we performed a comparison between the classical method and the Bayesian method using the Gamma distribution, but in chapter 4 and in the present chapter, we did not do this, because the objective in these chapters was to develop work with fully Bayesian results, without the use of comparisons from other methods.

Application

Over time technologies have played an extremely important role in all areas of knowledge. In general, one of its main functions today is to make work easier, more productive, and especially, to generate quality products that meet the wishes and needs of customers.

In this perspective, thinking about industries related manufacture in general, they are always prone to have failures in their manufacturing stages and at the end of the process, the generated products may not meet the specifications. In this sense, there are two possibilities: the customer ends up receiving a product that does not match their expectations or there is the possibility that the company is having losses due to lack of control manufacturing process. Another concern is related to the losses of customers due to the discontent of low quality of the products.

In this section, the interest is to verify if the proposing methodology can discriminate a process that is capable of producing units that meet the expectations of the consumers and company. Hence, we considered two data sets that describe the process of filling powdered juice bags with a desirable number of grams. The data is related to two different flavors, grape and strawberry, which were extracted from a study conducted by Molina [29] where the objective was to analyze using the C_{pk} index if the process is adequate with respect to the production tolerances. Table 6.1 presents 30 items with the weight of the juice packs (in grams) computed with an accuracy of three decimal number containing two different flavors, referred to as juice I and juice II. Further, in Figure 6.1 we present the control chart for both flavors assuming that $LSL = 18$ and $USL = 22$, such limits are considered due to the Institute of Metrology, Quality and Technology (INMETRO), that correspond to the regulatory agency in Brazil.

Table 6.1: Data set related to 30 packages of strawberry and grape juice

Juice I									
21.001	20.635	21.732	21.333	20.857	20.587	21.784	21.088	20.997	21.100
22.155	21.116	20.707	20.413	20.822	20.883	20.930	20.908	20.897	20.486
20.935	21.867	20.814	20.795	21.520	20.537	21.438	20.621	20.975	20.919
Juice II									
22.572	21.376	20.768	21.833	19.970	21.583	21.813	22.025	20.892	20.241
21.816	21.232	21.730	20.529	21.435	21.106	20.519	21.263	20.684	21.233
19.624	21.150	20.962	21.024	20.316	21.942	21.495	20.819	20.973	21.115

It can be seen from the data in Table 6.1 and Figure 6.1 that the process is not centered around its mean, in fact, the bags are receiving around 1 gram then it should be. More

importantly, there are values that out of the limits which indicate that the process is not adequate with respect to the production tolerances. After, identifying the problems in the productions that lead us to remove the 11-th value of Juice I and the 1-th and 8-th units of Juice II, we moved to estimate the C_{pk} index related to the process.

Using the obtained data, we performed the Kolmogorov-Smirnov (KS) test to investigate whether the data can be described by the Normal, Gamma, Weibull, and GW distributions. In all cases, the returned p-values were greater than 0.05 which assuming a 95% confidence level the data set can be described by these models. Note that the KS cannot be used to discriminate between the models, in this case, we need to consider discrimination methods such as the Akaike [2] information criterion (AIC) computed through $AIC = -2l(\boldsymbol{\theta}; x) + 2k$ and Bayesian information criterion (BIC) introduced by Schwarz [46] and given by $BIC = -2l(\boldsymbol{\theta}; x) + k \log(n)$, where k is the number of parameters to be fitted and $\boldsymbol{\theta}$ is the estimates of parameters assuming the different models. Among the selected models, the best is the one that provides the minimum values of those criteria. The estimation of the C_{pk} using the normal distribution was conducted by the standard MLE with confidence intervals estimated by bootstrap, as no objective Bayesian analysis was presented for this model. In the case of the Gamma distribution, we considered the same approach as discussed by Almeida et al. [3] which conducts an objective Bayesian analysis using the Gamma model. The inference using the Weibull and the GW is the same as discussed in the previous sections of the current work. Table 6.2 presents the estimated C_{pk} its respective 95% confidence/credibility intervals with estimated AIC and BIC for the different models.

Table 6.2: Estimates and 95% credibility intervals for the C_{pk} with its respective AIC and BIC.

	Juice I				Juice II			
Model	C_{pk}	$CI_{95\%}$	AIC	BIC	C_{pk}	$CI_{95\%}$	AIC	BIC
Normal	0.895	(0.688; 1.245)	30.13	32.86	0.553	(0.397; 0.800)	52.17	54.83
Gamma	0.849	(0.594; 1.125)	29.90	32.64	0.513	(0.323; 0.710)	52.52	55.18
Weibull	0.995	(0.751; 1.229)	39.28	42.01	0.765	(0.543; 0.967)	50.46	53.12
GW	0.928	(0.414; 1.419)	35.69	39.80	0.723	(0.320; 1.007)	52.58	56.58

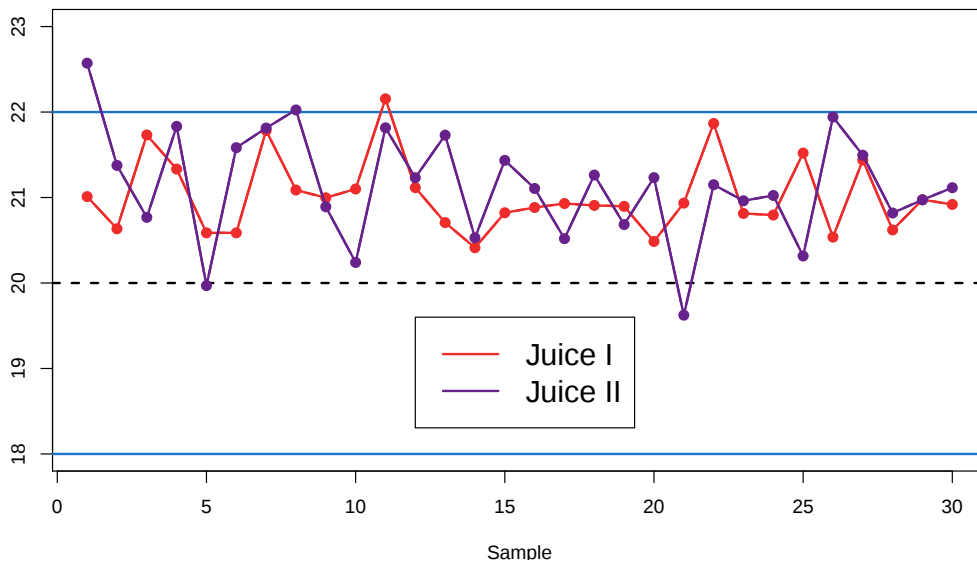


Figure 6.1: Control Chart for the Juice data set.

From the observed results, we notice that even removing the items that were out of the specification limits, the Gamma distribution returned better results for Juice I as returned smaller AIC and BIC values. In this case, the $C_{pk} = 0.849$ which is smaller than one. On the other hand, the credibility intervals contain values higher than one and, therefore, we cannot affirm that the process is not adequate with respect to the production tolerances. Now, for Juice II, we observed that the Weibull distribution returned a better fit when compared with its competitors, which justifies our proposed methodology. Here, the $C_{pk} = 0.765$ which is also smaller than one, however, the credibility intervals does contain values higher than one and, therefore, the process is not adequate with respect to the production tolerances.

A possible explanation for this might be that, in many cases, food companies are unable to adjust the filling machines according to the required specifications by the control standards, in this case, they consider the option of filling their products using the average range and the upper specification limit in order to not suffer penalties from the IPEN or INMETRO, which are the corresponding regulatory agencies in Brazil. This fact contributes to the occurrence of asymmetric data and leading to the waste of a significant amount of products. This led to significant losses for the company in financial aspects which should be strongly improved.

Final Considerations and Future Proposals

Process capability indexes usually assumes that the random process that generates the data follows a Gaussian model. However, in many situations the presence of asymmetric data is observe. In the proposed study, we considered the Weibull and Generalized Weibull distribution as asymmetric models to estimate the C_{pk} capability Index. The development of improved inferential procedures plays an important role in estimating PCI. Here, the parameters are assumed as random variables and are estimated from the Bayesian point of view. We proposed a Bayesian analysis through the gamma distribution together with an objective prior, called previous by matching prior that can return us a Bayesian estimator with good properties in terms of the mean relative errors, the mean square errors, and credibility intervals for the parameters and process capability index. For Weibull and WG we considered the objective reference priors to obtain posterior distributions with important properties, such as invariance under one-to-one transformation, consistency marginalization, and consistency sample properties. Hence, we provided an important review of the objective Bayesian methods and applied them successfully to estimate the PCI. As the posterior distributions do not have closed-form expressions, we considered Markov Chain Monte Carlo methods to sample from the posterior distributions, more specifically, the Metropolis-Hasting algorithms are derived for both models to estimate the C_{pk} .

A simulation study was presented to estimate the indexes assuming different values. For Gamma distribution we use posterior mean using the matching prior indicates better performance than the obtained with maximum likelihood estimator in terms of MREs and MSEs. For the Weibull and WG distributions we used Bayesian estimators based on a posterior medians that returned excellent results also in terms of MREs and MSEs. Additionally, the proposed posterior distributions achieved good results in terms of coverage probabilities obtained from the credibility intervals. The proposed approach allows us to estimate such intervals in a fast way with good precision, avoiding the need for resampling techniques that depend on maximization procedure methods and usually carried a high computational cost. A real dataset that describes the filling bags in a Brazilian company is present. We observed that the company has adopted a policy to include an additional gram of juice powder to avoid receive sanctions from the regulatory agency. From the tolerance limits, it can be noticed that the process is not adequate with respect to the production tolerances. After removing the damaged units and attesting that the samples were within the limits we computed the C_{pk} assuming that the data follows a Gaussian

model as well as the proposed asymmetric models. The discrimination methods showed that the random process that generates the data is not symmetric and the Weibull and Gamma distributions were the best models to describe the data, the adjusted C_{pk} showed that the estimates values under objective Bayesian analysis were smaller than 1.00 which indicates that the process is still not adequate with respect to the production tolerances and improvements in the process should be considered to decrease financial losses.

There is a large number of possible extensions of this current work. For example, the presence of covariates that may affect the quality of the product can be considered from a Bayesian perspective. In this case, new priors needs to be defined to take into account the additional parameters in the model. The development and comparison with other generalizations of the Weibull distribution should further be considered.

Bibliography

- [1] M Ahsanullah, M Maswadah, and Ali M Seham. Kernel inference on the generalized gamma distribution based on generalized order statistics. *Journal of Statistical Theory and Applications*, 12(2):152–172, 2013.
- [2] Hirotugu Akaike. A new look at the statistical model identification. *IEEE transactions on automatic control*, 19(6):716–723, 1974.
- [3] Marcello H Almeida, Pedro L Ramos, Gadde Srinivasa Rao, and Fernando Antonio Moala. Objective bayesian inference for the capability index of the gamma distribution. *Quality and Reliability Engineering International*, 2021.
- [4] Hassan S Bakouch, Sanku Dey, Pedro Luiz Ramos, and Francisco Louzada. Binomial-exponential 2 distribution: Different estimation methods with weather applications. *TEMA (São Carlos)*, 18(2):233–251, 2017.
- [5] James O Berger, Jose M Bernardo, Dongchu Sun, et al. Overall objective priors. *Bayesian Analysis*, 10(1):189–221, 2015.
- [6] Jose M Bernardo. Reference posterior distributions for bayesian inference. *Journal of the Royal Statistical Society: Series B (Methodological)*, 41(2):113–128, 1979.
- [7] José M Bernardo. Reference analysis. *Handbook of statistics*, 25:17–90, 2005.
- [8] George Casella and Roger L Berger. *Statistical inference*, volume 2. Duxbury Pacific Grove, CA, 2002.
- [9] John A Clements. Process capability calculations, for non-normal distributions. *Quality progress*, 22:95–100, 1989.
- [10] Guido Consonni, Dimitris Fouskakis, Brunero Liseo, Ioannis Ntzoufras, et al. Prior distributions for objective bayesian analysis. *Bayesian Analysis*, 13(2):627–679, 2018.
- [11] Sanku Dey, Devendra Kumar, Pedro L Ramos, and Francisco Louzada. Exponentiated chen distribution: Properties and estimation. *Communications in Statistics-Simulation and Computation*, 46(10):8118–8139, 2017.
- [12] LeRoy A Franklin and Wasserman Gary. Bootstrap confidence interval estimates of cpk: an introduction. *Communications in Statistics-Simulation and Computation*, 20(1):231–242, 1991.
- [13] John Geweke et al. *Evaluating the accuracy of sampling-based approaches to the calculation of posterior moments*, volume 196. Federal Reserve Bank of Minneapolis, Research Department Minneapolis, MN, 1991.

- [14] Patricia Ueda Gonçalves and Liane Werner. Comparação dos índices de capacidade do processo para distribuições não-normais. *Gestão & Produção*, 16(1):121–132, 2009.
- [15] Harold W Hager and Lee J Bain. Inferential procedures for the generalized gamma distribution. *Journal of the American Statistical Association*, 65(332):1601–1609, 1970.
- [16] Z Hosseinifard, B Abbasi, and STA Niaki. Process capability estimation for leukocyte filtering process in blood service: A comparison study. *IIE Transactions on Healthcare Systems Engineering*, 4(4):167–177, 2014.
- [17] Muhammad Kashif, Muhammad Aslam, Ali Hussein Al-Marshadi, and Chi-Hyuck Jun. Capability indices for non-normal distribution using gini’s mean difference as measure of variability. *IEEE Access*, 4:7322–7330, 2016.
- [18] Samuel Kotz and Norman L Johnson. Process capability indices—a review, 1992–2000. *Journal of quality technology*, 34(1):2–19, 2002.
- [19] Jerald F Lawless. *Statistical models and methods for lifetime data*, volume 362. John Wiley & Sons, 2011.
- [20] Víctor Leiva, Carolina Marchant, Helton Saulo, Muhammad Aslam, and Fernando Rojas. Capability indices for birnbaum–saunders processes applied to electronic and food industries. *Journal of Applied Statistics*, 41(9):1881–1902, 2014.
- [21] Francisco Louzada, Pedro L Ramos, and Gleici SC Perdoná. Different estimation procedures for the parameters of the extended exponential geometric distribution for medical data. *Computational and mathematical methods in medicine*, 2016, 2016.
- [22] Francisco Louzada, Pedro L Ramos, and Eduardo Ramos. A note on bias of closed-form estimators for the gamma distribution derived from likelihood equations. *The American Statistician*, 73(2):195–199, 2019.
- [23] Francisco Louzada and Pedro Luiz Ramos. Efficient closed-form maximum a posteriori estimators for the gamma distribution. *Journal of Statistical Computation and Simulation*, 88(6):1134–1146, 2018.
- [24] Mohamed Abdul Wahab Mahmoud, Neveen Mohamed Kilany, and Lobna Hisham El-Refai. Inference of the lifetime performance index with power rayleigh distribution based on progressive first-failure–censored data. *Quality and Reliability Engineering International*, 36(5):1528–1536, 2020.
- [25] Mazhar Ali Khan Malik. A note on the physical meaning of the weibull distribution. *IEEE Transactions on Reliability*, 24(1):95–95, 1975.
- [26] A Marani, I Lavagnini, and C Buttazzoni. Statistical study of air pollutant concentrations via generalized gamma distributions. *Journal of the Air Pollution Control Association*, 36(11):1250–1254, 1986.
- [27] Fanbing Meng, Jun Yang, and Shuo Huang. Hypothesis testing of process capability index c_{pk} from the perspective of generalized fiducial inference. *Quality and Reliability Engineering International*, 2020.
- [28] Robert B Miller. Bayesian analysis of the two-parameter gamma distribution. *Technometrics*, 22(1):65–69, 1980.

- [29] R.A. Molina. *Técnicas de Avaliação de Medidas e Verificação dos Índices de Capacidade*. Presidente Prudente - São Paulo, 2018.
- [30] Rahul Mukerjee and Dipak K Dey. Frequentist validity of posterior quantiles in the presence of a nuisance parameter: higher order asymptotics. *Biometrika*, 80(3):499–505, 1993.
- [31] Kurt Palmer and K-L Tsui. A review and interpretations of process capability indices. *Annals of Operations Research*, 87:31–47, 1999.
- [32] W Panichkitkosolkul. Confidence intervals for the process capability index cp based on confidence intervals for variance under non-normality. *Malays. J. Math. Sci*, 10(1):101–115, 2016.
- [33] T Pham and J Almhana. The generalized gamma distribution: its hazard rate and stress-strength model. *IEEE transactions on reliability*, 44(3):392–397, 1995.
- [34] Manuel R Piña-Monarrez, Jesús F Ortiz-Yañez, and Manuel I Rodríguez-Borbón. Non-normal capability indices for the weibull and lognormal distributions. *Quality and Reliability Engineering International*, 32(4):1321–1329, 2016.
- [35] Ross L Prentice. A log gamma model and its maximum likelihood estimation. *Biometrika*, 61(3):539–544, 1974.
- [36] R Core Team. *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria, 2020.
- [37] Balasundar I Raju and Mandayam A Srinivasan. Statistics of envelope of high-frequency ultrasonic backscatter from human skin in vivo. *IEEE transactions on ultrasonics, ferroelectrics, and frequency control*, 49(7):871–882, 2002.
- [38] Eduardo Ramos, Pedro L Ramos, and Francisco Louzada. Posterior properties of the weibull distribution for censored data. *Statistics & Probability Letters*, 166:108873, 2020.
- [39] Pedro Luiz Ramos, Dipak K. Dey, Francisco Louzada, and Eduardo Ramos. On posterior properties of the two parameter gamma family of distributions. *Anais da Academia Brasileira de Ciências*, 2020.
- [40] Pedro Luiz Ramos and Francisco Louzada. Bayesian reference analysis for the generalized gamma distribution. *IEEE Communications Letters*, 22(9):1950–1953, 2018.
- [41] Giovani Carrara Rodrigues, Francisco Louzada, and Pedro Luiz Ramos. Poisson-exponential distribution: different methods of estimation. *Journal of Applied Statistics*, 45(1):128–144, 2018.
- [42] Mahendra Saha, Sanku Dey, Abhimanyu Singh Yadav, Sajid Ali, et al. Confidence intervals of the index c_{pk} for normally distributed quality characteristics using classical and bayesian methods of estimation. *Brazilian Journal of Probability and Statistics*, 35(1):138–157.
- [43] Mahendra Saha, Sanku Dey, Abhimanyu Singh Yadav, and Sumit Kumar. Classical and bayesian inference of c py for generalized lindley distributed quality characteristic. *Quality and Reliability Engineering International*, 35(8):2593–2611, 2019.

- [44] Carlos Aparecido dos Santos et al. Dados de sobrevivência multivariados na presença de covariáveis e observações censuradas: uma abordagem bayesiana. 2010.
- [45] J Sappakitkamjorn and S.A Niwitpong. Biased estimation of the capability indices, c_p , c_{pk} and c_{pm} using bootstrap and jackknife methods. *Asia Pacific Management Review*, 10:1–7, 2006.
- [46] Gideon Schwarz et al. Estimating the dimension of a model. *Annals of statistics*, 6(2):461–464, 1978.
- [47] Bahar Sennaroglu and Ozlem Senvar. Performance comparison of box-cox transformation and weighted variance methods with weibull distribution. *J. Aeronaut. Space Technol*, 8(2):49–55, 2015.
- [48] Ozlem Senvar and Cengiz Kahraman. Fuzzy process capability indices using clements’ method for non-normal processes. *Journal of Multiple-Valued Logic & Soft Computing*, 22:95–121, 2014.
- [49] Ozlem Senvar and Cengiz Kahraman. Type-2 fuzzy process capability indices for non-normal processes. *Journal of Intelligent & Fuzzy Systems*, 27(2):769–781, 2014.
- [50] Ozlem Senvar and Bahar Sennaroglu. Comparing performances of clements, box-cox, johnson methods with weibull distributions for assessing process capability. *Journal of Industrial Engineering and Management*, 9(3):634–656, 2016.
- [51] Fred Spiring, Bartholomew Leung, Smiley Cheng, and Anthony Yeung. A bibliography of process capability papers. *Quality and Reliability Engineering International*, 19(5):445–460, 2003.
- [52] E Webb Stacy and G Arthur Mihram. Parameter estimation for a generalized gamma distribution. *Technometrics*, 7(3):349–358, 1965.
- [53] UPORABO Statisti and N Tehnike. Improving the process capability of a turning operation by the application of statistical techniques. *Materiali in tehnologije*, 43(1):55–59, 2009.
- [54] Dongchu Sun. A note on noninformative priors for Weibull distributions. *Journal of Statistical Planning and Inference*, 61(2):319–338, 1997.
- [55] Dongchu Sun and Keying Ye. Frequentist validity of posterior quantiles for a two-parameter exponential family. *Biometrika*, 83(1):55–65, 1996.
- [56] Mahdi Teimouri, Seyed M Hoseini, and Saralees Nadarajah. Comparison of estimation methods for the weibull distribution. *Statistics*, 47(1):93–109, 2013.
- [57] Robert Tibshirani. Noninformative priors for one parameter of many. *Biometrika*, 76(3):604–608, 1989.
- [58] Gary S Wasserman and Leroy A Franklin. Standard bootstrap confidence interval estimates of cpk. *Computers & industrial engineering*, 22(2):171–176, 1992.
- [59] Waloddi Weibull. Wide applicability. *Journal of applied mechanics*, 103(730):293–297, 1951.
- [60] Chia-Huang Wu. Modified processes capability assessment with dynamic mean shift. *Quality and Reliability Engineering International*, 36(4):1258–1271, 2020.

-
- [61] Bong-Jin Yum and Kwan-Woo Kim. A bibliography of the literature on process capability indices: 2000–2009. *Quality and Reliability Engineering International*, 27(3):251–268, 2011.
 - [62] Zhiguo Zeng, Meilin Wen, and Rui Kang. Belief reliability: A new metrics for products' reliability. *Fuzzy Optimization and Decision Making*, 12(1):15–27, 2013.
 - [63] Jianchun Zhang. Conditional confidence intervals of process capability indices following rejection of preliminary tests. 2010.