



Universidade Federal da Paraíba  
Centro de Informática  
Programa de Pós-Graduação em Modelagem Matemática e Computacional

## DISTRIBUIÇÃO QUI-QUADRADO INF: UMA NOVA ABORDAGEM PARA O APERFEIÇOAMENTO DO TESTE DA RAZÃO DE VEROSSIMILHANÇAS.

Antônio Rubens de Sousa

Dissertação de Mestrado apresentada ao  
Programa de Pós-Graduação em Modelagem  
Matemática e Computacional, UFPB, da  
Universidade Federal da Paraíba, como parte  
dos requisitos necessários à obtenção do título  
de Mestre em Modelagem Matemática e  
Computacional.

Orientadores: Claudio Javier Tablada  
Pedro Rafael Diniz Marinho

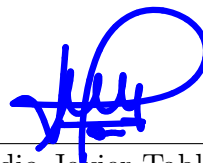
João Pessoa  
Maio de 2020

DISTRIBUIÇÃO QUI-QUADRADO INF: UMA NOVA ABORDAGEM PARA O  
APERFEIÇOAMENTO DO TESTE DA RAZÃO DE VEROSSIMILHANÇAS.

Antônio Rubens de Sousa

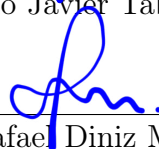
DISSERTAÇÃO SUBMETIDA AO CORPO DOCENTE DO PROGRAMA DE  
PÓS-GRADUAÇÃO EM MODELAGEM MATEMÁTICA E COMPUTACIONAL  
(PPGMMC) DO CENTRO DE INFORMÁTICA DA UNIVERSIDADE FEDERAL  
DA PARAÍBA COMO PARTE DOS REQUISITOS NECESSÁRIOS PARA A  
OBTENÇÃO DO GRAU DE MESTRE EM CIÊNCIAS EM MODELAGEM  
MATEMÁTICA E COMPUTACIONAL.

Examinada por:



---

Prof. Claudio Javier Tablada, D.Sc.



---

Prof. Pedro Rafael Diniz Marinho, D.Sc.



---

Prof. Tatiene Correia de Souza, D.Sc.



---

Prof. Renilma Pereira da Silva, D.Sc.

.

**Catálogo na publicação**  
**Seção de Catalogação e Classificação**

S725d Sousa, Antônio Rubens de.

Distribuição qui-quadrado inf: uma nova abordagem para o aperfeiçoamento do teste da razão de verossimilhanças / Antônio Rubens de Sousa. - João Pessoa, 2020.  
81 f. : il.

Orientação: Claudio Javier Tablada.

Coorientação: Pedro Rafael Diniz Marinho.

Dissertação (Mestrado) - UFPB/CI/PPGMMC.

1. Estatística matemática. 2. Teste de hipótese. 3. Teste da razão de verossimilhanças. 4. Simulação de Monte-Carlo. 5. Pacote LikRatioTest. I. Tablada, Claudio Javier. II. Marinho, Pedro Rafael Diniz. III. Título.

UFPB/BC

CDU 519.2(043)

*Aos meus pais.*

# Agradecimentos

A Deus por ter me dado saúde e força para superar todos os obstáculos da vida.

Aos meus pais José Lindenor e Maria da Consolação, pelo incentivo, apoio incondicional e educação dados ao longo dessa vida.

Ao meu irmão Carlos Reuber pelo apoio durante esses anos.

A meus orientadores Dr. Claudio Javier Tablada e Dr. Pedro Rafael Diniz Marinho, pelo suporte dado para a realização desta pesquisa, correções e incentivos. Além disso, agradeço pelo laço de amizade que fizemos durante toda a essa trajetória.

Ao meu colega Alexandre Gomes por me apoiar e me acompanhar durante todo esse mestrado. Agradecer por sempre estar disposto ajudar nas correções e edições desta dissertação no LaTeX.

À Maria Cristiane Magalhães Brandão, minha ex-orientadora de graduação, por todo apoio dado para que eu pudesse ingressar neste mestrado e por toda sua contribuição em minha vida acadêmica.

Ao Flávio Alexandre Falcão Nascimento, meu ex-professor de graduação, que juntamente com Mariana Maia me ajudaram no processo inicial da minha chegada em João Pessoa-PB.

À Mylenna Sá por me acolher em seu apartamento durante os primeiros meses, assim que cheguei em João Pessoa-PB para cursar o mestrado.

A todos os meus colegas que me acompanharam e batalharam junto comigo nessa longa trajetória árdua. Saibam que lembrarei de todos sempre.

A esta universidade, seu corpo docente, direção e administração.

A todos que direta ou indiretamente fizeram parte da minha formação, o meu muito obrigado.

À CAPES pela contribuição financeira de forma parcial para me manter durante o mestrado.

Resumo da Dissertação apresentada ao PPGMMC/CI/UFPB como parte dos requisitos necessários para a obtenção do grau de Mestre em Ciências (M.Sc.)

## DISTRIBUIÇÃO QUI-QUADRADO INF: UMA NOVA ABORDAGEM PARA O APERFEIÇOAMENTO DO TESTE DA RAZÃO DE VEROSSIMILHANÇAS.

Antônio Rubens de Sousa

Maio/2020

Orientadores: Claudio Javier Tablada

Pedro Rafael Diniz Marinho

Programa: Modelagem Matemática e Computacional

Um dos principais objetivos da estatística é realizar inferência em uma população ou fenômeno a partir de um subconjunto de dados desta, denominado de amostra. Uma das maneiras de realizar inferência é realizando teste de hipóteses. De forma geral, podemos dizer que é a partir de uma amostra da população que estabeleceremos uma regra de decisão segundo a qual rejeitaremos ou não rejeitaremos a hipótese proposta, denominada hipótese nula. Um procedimento geral que produz testes razoáveis é o Teste da Razão de Verossimilhanças (TRV). Para aplicar o TRV, precisamos conhecer a verdadeira distribuição da estatística da razão de verossimilhanças  $\lambda^*(\mathbf{x})$  que, geralmente, não é fácil de ser obtida. Todavia, é conhecido na literatura que a estatística  $RV = -2\log(\lambda^*(\mathbf{x}))$  segue distribuição aproximada qui-quadrado quando o teste é baseado em uma amostra de tamanho grande. No entanto, o uso da distribuição qui-quadrado como uma aproximação à verdadeira distribuição da estatística  $RV$  pode levar a inferências imprecisas quando o tamanho da amostra é pequeno. O objetivo desta pesquisa é aperfeiçoar o TRV quando o teste for baseado em amostra de tamanho pequeno ou moderado. Para isso, fazemos uso de novas famílias de distribuições, denotadas de sup e inf. A partir de algumas propriedades da família inf de distribuições, propomos uma nova abordagem para o aperfeiçoamento do TRV. Especificamente, utilizamos a distribuição qui-quadrado inf como uma distribuição de correção da qui-quadrado para a obtenção do quantil que determina a região crítica do teste. Além disso, neste trabalho é criado o pacote computacional `LikRatioTest`, escrito na linguagem R, com o objetivo de avaliar o desempenho do TRV aperfeiçoado ( $TRV^*$ ) e compará-lo com o TRV assintótico clássico. Este pacote permite fazer simulações de Monte Carlo, impondo vários

cenários, e assim, calcular a taxa de rejeição da hipótese nula utilizando as duas abordagens (clássica e aperfeiçoada). Tal pacote é open source e encontra-se disponível no GitHub para instalação no R. Para ilustrar a utilidade da nossa proposta, mostramos dois exemplos numéricos usando conjuntos de dados reais.

Abstract of Dissertation presented to PPGMMC/CI/UFPB as a partial fulfillment of the requirements for the degree of Master of Science (M.Sc.)

## CHI-SQUARE INF DISTRIBUTION: A NEW APPROACH TO IMPROVE THE LIKELIHOOD RATIO TEST.

Antônio Rubens de Sousa

May/2020

Advisors: Claudio Javier Tablada

Pedro Rafael Diniz Marinho

Program: Computational Mathematical Modelling

One of the main objectives of statistics is to make inference in a population or phenomenon from a subset of data from this, called a sample. One way to make inference is to perform hypothesis testing. In general, we can say that it is from a sample of the population that we will establish a decision rule according to which we will reject or not reject the proposed hypothesis, called null hypothesis. A general procedure that produces reasonable tests is the Likelihood Ratio test (TRV). To apply the TRV, we need to know the true distribution of the likelihood ratio  $\lambda^*(\mathbf{x})$  which is generally not easy to obtain. However, it is known in the literature that the  $RV = -2\log(\lambda^*(\mathbf{x}))$  statistic follows an approximate chi-square distribution when the test is based on a large sample size. However, using the chi-square distribution as an approximation to the true distribution of the  $RV$  statistic can lead to inaccurate inferences when the sample size is small. The objective of this research is to improve the TRV when the test is based on a small or moderate sample. For this, we make use of new families of distributions, denoted sup and inf. Based on some properties of the inf distribution family, we propose a new approach for the improvement of TRV. Specifically, we used the chi-square inf distribution as a chi-square correction distribution to obtain the quantile that determines the critical region of the test. In addition, this work creates the computational package `LikRatioTest`, written in the R language, with the objective of evaluating the performance of the improved TRV (TRV\*) and comparing it with the classic asymptotic TRV. This package allows you to do Monte Carlo simulations, imposing various scenarios, and thus, calculate the rejection rate of the null hypothesis using the two approaches (classical and improved). This package is open source and is available on GitHub for installation



on R. To illustrate the usefulness of our proposal, we show two numerical examples using real data sets.

# Sumário

|                                                                                                       |             |
|-------------------------------------------------------------------------------------------------------|-------------|
| <b>Lista de Figuras</b>                                                                               | <b>xii</b>  |
| <b>Lista de Tabelas</b>                                                                               | <b>xiii</b> |
| <b>Lista de Símbolos</b>                                                                              | <b>xiv</b>  |
| <b>Lista de Abreviaturas</b>                                                                          | <b>xvi</b>  |
| <b>1 Introdução</b>                                                                                   | <b>1</b>    |
| <b>2 Preliminares</b>                                                                                 | <b>4</b>    |
| 2.1 Conceitos básicos na teoria das probabilidades . . . . .                                          | 4           |
| 2.2 Conceitos básicos na teoria da inferência estatística . . . . .                                   | 7           |
| <b>3 Famílias de distribuições sup e inf</b>                                                          | <b>12</b>   |
| 3.1 O método $T-X$ na geração de novas famílias de distribuições contínuas                            | 12          |
| 3.2 Famílias $G_c^{\text{sup}}$ e $G_c^{\text{inf}}$ . . . . .                                        | 14          |
| 3.2.1 Obtendo as Famílias $G_c^{\text{sup}}$ e $G_c^{\text{inf}}$ . . . . .                           | 14          |
| 3.2.2 Propriedades das famílias $G_c^{\text{sup}}$ e $G_c^{\text{inf}}$ . . . . .                     | 18          |
| 3.2.3 Interpretação Física dos Modelos sup e inf . . . . .                                            | 24          |
| 3.3 Gerando números pseudo-aleatórios das distribuições $G_c^{\text{sup}}$ e $G_c^{\text{inf}}$ . . . | 26          |
| 3.3.1 Algoritmo para gerar números pseudo-aleatórios da família $G_c^{\text{sup}}$                    | 26          |
| 3.3.2 Algoritmo para gerar números pseudo-aleatórios da família $G_c^{\text{inf}}$                    | 27          |
| 3.4 Estimadores de máxima verossimilhança das famílias $G_c^{\text{sup}}$ e $G_c^{\text{inf}}$ . .    | 28          |
| 3.4.1 Função de log-verossimilhança da família $G_c^{\text{sup}}$ . . . . .                           | 28          |
| 3.4.2 Função de log-verossimilhança da família $G_c^{\text{inf}}$ . . . . .                           | 29          |
| <b>4 Uma nova abordagem para o aperfeiçoamento do teste da razão de verossimilhanças</b>              | <b>31</b>   |
| 4.1 A nova abordagem . . . . .                                                                        | 32          |
| 4.1.1 A distribuição qui-quadrado inf . . . . .                                                       | 33          |

|          |                                                                        |           |
|----------|------------------------------------------------------------------------|-----------|
| <b>5</b> | <b>O pacote LikRatioTest</b>                                           | <b>39</b> |
| 5.1      | Função 'lrt' . . . . .                                                 | 40        |
| 5.2      | Função 'est_q ' . . . . .                                              | 42        |
| 5.3      | Função 'mc' . . . . .                                                  | 44        |
| <b>6</b> | <b>Simulações e Exemplos Numéricos</b>                                 | <b>50</b> |
| 6.1      | Simulações . . . . .                                                   | 50        |
| 6.1.1    | Testes envolvendo a Dist. Weibull exponencializada . . . . .           | 51        |
| 6.1.2    | Testes envolvendo a Dist. Gama . . . . .                               | 54        |
| 6.1.3    | Testes envolvendo a Dist. Weibull inversa de dois parâmetros . . . . . | 56        |
| 6.2      | Exemplos Numéricos . . . . .                                           | 58        |
| <b>7</b> | <b>Conclusões</b>                                                      | <b>62</b> |
|          | <b>Referências Bibliográficas</b>                                      | <b>64</b> |

# Lista de Figuras

|     |                                                                                                                                                                                   |    |
|-----|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 3.1 | Sistema de componentes em série para interpretação física do modelo sup. . . . .                                                                                                  | 25 |
| 3.2 | Sistema de componentes em paralelo para interpretação física do modelo inf. . . . .                                                                                               | 26 |
| 4.1 | Plots das f.d.p da distribuição qui-quadrado (linha contínua) e qui-quadrado inf (linhas tracejadas) . . . . .                                                                    | 35 |
| 4.2 | Histograma de observações da estatística $RV$ (para $n = 10$ ). Curvas das densidades da distribuição qui-quadrado (linha contínua) e qui-quadrado inf (linha tracejada). . . . . | 38 |
| 6.1 | Distorções dos tamanhos dos testes envolvendo a distribuição Weibull Exponencializada: $\gamma = 10\%$ . . . . .                                                                  | 54 |
| 6.2 | Distorções dos tamanhos dos testes envolvendo a distribuição Gama: $\gamma = 10\%$ . . . . .                                                                                      | 57 |
| 6.3 | Distorções dos tamanhos dos testes envolvendo a distribuição Weibull Inversa de dois parâmetros: $\gamma = 10\%$ . . . . .                                                        | 60 |

# Lista de Tabelas

|     |                                                                                                                                         |    |
|-----|-----------------------------------------------------------------------------------------------------------------------------------------|----|
| 2.1 | Tipos de erros em testes de hipóteses. . . . .                                                                                          | 9  |
| 6.1 | Teste envolvendo a distribuição Weibull Exponencializada. Taxas de rejeição nulas (%), inferências sobre $\lambda$ . . . . .            | 52 |
| 6.2 | Teste envolvendo a distribuição Weibull Exponencializada. Taxas de rejeição nulas (%), inferências sobre $\alpha$ e $\lambda$ . . . . . | 53 |
| 6.3 | Teste envolvendo a distribuição Gama. Taxas de rejeição nulas (%), inferências sobre $\alpha$ . . . . .                                 | 55 |
| 6.4 | Teste envolvendo a distribuição Gama. Taxas de rejeição nulas (%), inferências sobre $\alpha$ e $\beta$ . . . . .                       | 56 |
| 6.5 | Teste envolvendo a distribuição Weibull Inversa. Taxas de rejeição nulas (%), inferências sobre $\beta$ . . . . .                       | 58 |
| 6.6 | Teste envolvendo a distribuição Weibull Inversa. Taxas de rejeição nulas (%), inferências sobre $\alpha$ e $\beta$ . . . . .            | 59 |
| 6.7 | Dados I: Vinte itens testados até a falha. . . . .                                                                                      | 60 |
| 6.8 | Dados II: Fluxos anuais do rio Weldon em Mill Grove, Missouri (1930-59). . . . .                                                        | 61 |

# Lista de Símbolos

|                                      |                                                                                                                      |
|--------------------------------------|----------------------------------------------------------------------------------------------------------------------|
| $H_0$                                | Hipótese nula, p. <a href="#">1</a>                                                                                  |
| $H_1$                                | Hipótese alternativa, p. <a href="#">1</a>                                                                           |
| $\Theta$                             | Espaço paramétrico, p. <a href="#">1</a>                                                                             |
| $\Theta_0$                           | Espaço paramétrico sob a hipótese nula, p. <a href="#">1</a>                                                         |
| $\Theta_1$                           | Diferença entre $\Theta$ e $\Theta_0$ , p. <a href="#">1</a>                                                         |
| $\lambda^*(\mathbf{x})$              | Estatística da razão de verossimilhanças, p. <a href="#">1</a>                                                       |
| $\mathbf{A}_1^*$                     | Região crítica que faz o uso do quantíl da verdadeira distribuição de $\lambda^*(\mathbf{x})$ , p. <a href="#">1</a> |
| $RV$                                 | Estatística do teste $-2 \log(\lambda^*(\mathbf{x}))$ , p. <a href="#">1</a>                                         |
| $\chi_k^2$                           | Distribuição qui-quadrado com $k$ graus de liberdade, <a href="#">2</a>                                              |
| $\Omega$                             | Espaço amostral, <a href="#">3</a>                                                                                   |
| $\mathcal{F}$                        | $\sigma$ -álgebra de eventos, <a href="#">3</a>                                                                      |
| $\mathcal{P}$                        | Medida de probabilidade, <a href="#">3</a>                                                                           |
| $(\Omega, \mathcal{F}, \mathcal{P})$ | Espaço de probabilidade, <a href="#">3</a>                                                                           |
| $Q(\cdot)$                           | Função quantílica, <a href="#">5</a>                                                                                 |
| $L(\cdot)$                           | Função de Verossimilhança, <a href="#">7</a>                                                                         |
| $l(\cdot)$                           | Função de log-verossimilhança, <a href="#">7</a>                                                                     |
| $\gamma$                             | Nível de significância ou nível nominal ou erro tipo I, <a href="#">8</a>                                            |
| $\beta$                              | Erro tipo II, <a href="#">8</a>                                                                                      |
| $q$                                  | Quantil $1 - \gamma$ através da distribuição $\chi_k^2$ , <a href="#">10</a>                                         |
| $\mathbf{B}_1^*$                     | Região crítica que faz o uso de $q$ , <a href="#">10</a>                                                             |

|                        |                                                                                                                                   |
|------------------------|-----------------------------------------------------------------------------------------------------------------------------------|
| $G_c^{\text{sup}}$     | Distribuição generalizada sup que tem $G$ como distribuição de base, <a href="#">13</a>                                           |
| $G_c^{\text{inf}}$     | Distribuição generalizada inf que tem $G$ como distribuição de base, <a href="#">13</a>                                           |
| $\sim$                 | Segue distribuição, <a href="#">25</a>                                                                                            |
| $\mathcal{U}(0, 1)$    | Distribuição uniforme no intervalo $(0, 1)$ , <a href="#">25</a>                                                                  |
| $\mathbb{E}(X)$        | Esperança de uma variável aleatória $X$ , <a href="#">32</a>                                                                      |
| $q'$                   | Quantil $1 - \gamma$ através da verdadeira distribuição de $RV$ , <a href="#">32</a>                                              |
| $\gamma(\cdot, \cdot)$ | Função Gama incompleta inferior, <a href="#">32</a>                                                                               |
| $\Gamma(\cdot)$        | Função Gama, <a href="#">32</a>                                                                                                   |
| $q_{\text{inf}}$       | Quantil $1 - \gamma$ através da distribuição qui-quadrado inf com $k$ graus de liberdade e $c = 0.5 \log(n)$ , <a href="#">33</a> |
| $\mathbf{C}_1^*$       | Região crítica que faz o uso de $q_{\text{inf}}$ , <a href="#">35</a>                                                             |
| $\tilde{\theta}$       | Estimadores de máxima verossimilhança restrito, <a href="#">36</a>                                                                |
| $\hat{\theta}$         | Estimador de máxima verossimilhança irrestrito, <a href="#">36</a>                                                                |

# Lista de Abreviaturas

|       |                                                                       |
|-------|-----------------------------------------------------------------------|
| TRV   | Teste da Razão de Verossimilhanças, p. <a href="#">1</a>              |
| TRV*  | Teste da Razão de Verossimilhanças Aperfeiçoado, p. <a href="#">1</a> |
| V.A.  | Variável Aleatória, p. <a href="#">3</a>                              |
| f.d.a | função de distribuição acumulada, p. <a href="#">4</a>                |
| f.d.p | função densidade de probabilidade, p. <a href="#">4</a>               |
| MMV   | Método de Máxima Verossimilhança, p. <a href="#">6</a>                |
| EMV   | Estimador de Máxima Verossimilhança, p. <a href="#">6</a>             |
| KS    | Kolmogorov - Smirnov, p. <a href="#">60</a>                           |



# Capítulo 1

## Introdução

Um dos principais objetivos da estatística é realizar inferência em uma população ou fenômeno a partir de um subconjunto de dados desta, denominado de amostra. Para isso, precisamos considerar o modelo que represente o fenômeno de interesse. Segundo [Cordeiro e Cribari-Neto \(2014\)](#) um modelo é uma representação simplificada de uma realidade mais abrangente. Um modelo estatístico representa um fenômeno que ocorre de maneira incerta, e normalmente é indexado por quantidades fixas e que, em geral, são desconhecidas, chamadas de parâmetros. Desta forma, a inferência estatística é realizada para tais parâmetros usando dados coletados previamente. Realizar inferência para os parâmetros do modelo equivale a fazermos inferência para o modelo e, por consequência, no fenômeno que ele descreve.

A inferência pode ser realizada de três maneiras diferentes: estimação pontual, estimação intervalar e teste de hipóteses. ([CORDEIRO e CRIBARI-NETO, 2014](#); [LEHMANN e ROMANO, 2006](#)). Em nosso trabalho daremos ênfase em testes de hipóteses.

Vale lembrar, que em um teste de hipóteses, denominamos de parâmetros de interesse àqueles parâmetros que estão sendo testados na hipótese nula, e o restante são chamados de parâmetros de perturbação. Um procedimento geral que produz testes razoáveis e que pode ser utilizado em muitos casos, sem muita dificuldade, é o Teste da Razão de Verossimilhanças (TRV).

O TRV pode ser usado como um procedimento para testar hipóteses compostas do tipo  $H_0 : \theta \in \Theta_0$  versus  $H_1 : \theta \in \Theta_1$ , em que  $\{\Theta_0, \Theta_1\}$  é uma partição do espaço paramétrico  $\Theta$ . Tal procedimento consiste em definir a região crítica

$$\mathbf{A}_1^* = \left\{ \mathbf{x}; \lambda^*(\mathbf{x}) = \frac{\sup_{\theta \in \Theta_0} L(\theta; \mathbf{x})}{\sup_{\theta \in \Theta} L(\theta; \mathbf{x})} \leq c \right\},$$

sendo  $\lambda^*(\mathbf{x})$  a estatística da razão de verossimilhanças com base na amostra observada  $\mathbf{x} = (x_1, x_2, \dots, x_n)$ .

Para aplicar o TRV, precisamos conhecer a verdadeira distribuição de  $\lambda^*(\mathbf{x})$  para

obtermos o valor de  $c$  que determina a região crítica do teste. Porém, geralmente, esta não é fácil de ser obtida. Todavia, é conhecido na literatura que, sob hipótese nula, a estatística  $RV = -2 \log(\lambda^*(\mathbf{x}))$  tem distribuição assintótica qui-quadrado com  $k$  graus de liberdade ( $\chi_k^2$ ), sendo  $k$  o número de parâmetros de interesse a ser testados na hipótese nula. No entanto, usar a distribuição qui-quadrado como uma aproximação à verdadeira distribuição da estatística  $RV$  pode levar a inferências imprecisas quando o tamanho da amostra é pequeno.

O objetivo deste trabalho é aperfeiçoar o teste da razão de verossimilhanças quando o teste for baseado em uma amostra de tamanho pequeno ou moderado. Para isso, fazemos uso de novas famílias de distribuições, denotadas de sup e inf. A partir de algumas propriedades da família inf de distribuições, propomos uma nova abordagem para o aperfeiçoamento do TRV. Especificamente, utilizamos a distribuição qui-quadrado inf como uma distribuição de correção da qui-quadrado para a obtenção do quantil que determina a região crítica do teste.

Este trabalho está organizado em sete capítulos. O Capítulo 2 foi reservado para leitores que não são da área de probabilidade, inferência estatística ou áreas afins. Assim sendo, na Seção 2.1, abordamos conceitos básicos na teoria das probabilidades. Já na Seção 2.2, é realizado um breve estudo sobre resultados da inferência estatística, que dão suporte a nossa proposta de trabalho.

No Capítulo 3, fazemos um estudo aprofundado sobre as famílias de distribuições sup e inf (TABLADA, 2017). Na Seção 3.1 abordamos o método  $T-X$  (ALZAA-TREH *et al.*, 2013), o qual foi utilizado em Tablada (2017) para construir tais famílias de distribuições. Na Seção 3.2 fazemos todo um processo construtivo para se chegar às famílias sup e inf, finalizando a seção dando uma interpretação física aos modelos. Além disso, enunciamos e demonstramos propriedades que os mesmos possuem e obtivemos novas propriedades que dão suporte a nossa proposta. Na Seção 3.3 mostramos os algoritmos para gerar números pseudo-aleatórios de ambas as famílias. Na Seção 3.4 abordamos as funções de log-verossimilhança de cada um dos modelos e mostramos resultados limite de tais distribuições.

O Capítulo 4 é destinado a nossa proposta de trabalho, em que mostramos uma nova abordagem para o aperfeiçoamento do teste da razão de verossimilhanças. Iniciamos com uma problematização do tema. Em seguida abordamos a distribuição qui-quadrado inf, que embasa esta pesquisa.

No Capítulo 5, abordamos sobre o pacote **LikRatioTest**, versão 0.1, construído utilizando a linguagem R, em que instruímos o leitor de como manuseá-lo. O pacote foi utilizado para realizar todas as simulações deste trabalho.

O Capítulo 6 é reservado para tratarmos de todas as simulações e exemplos numéricos, tendo em vista validar nossa proposta. Na Seção 6.1, objetivamos avaliar o desempenho do TRV aperfeiçoado (TRV\*) e compará-lo com o TRV clássico. Para

isso, realizamos simulações de Monte Carlo e utilizamos como critério de comparação a taxa de rejeição empírica da hipótese nula. Na Seção 6.2, mostramos através de exemplos numéricos, envolvendo conjuntos de dados reais, a utilidade da nossa proposta. Por fim, destinamos o Capítulo 7 para conter as considerações finais deste trabalho.

# Capítulo 2

## Preliminares

### 2.1 Conceitos básicos na teoria das probabilidades

A teoria das probabilidades é desenvolvida com base em experimentos aleatórios. De modo geral, um experimento aleatório é caracterizado por três fatos. O primeiro é que se for possível repetir as mesmas condições do experimento, os resultados, em diferentes realizações, podem ser diferentes. O segundo é que é possível descrever o conjunto de todos os possíveis resultados do experimento. Já o terceiro é que se o experimento é repetido um grande número de vezes, uma configuração definida ou regularidade surgirá.

Os resultados de um experimento aleatório são caracterizados, matematicamente, por três componentes, a saber: espaço amostral, coleção de conjuntos de interesse e uma probabilidade. O espaço amostral é o conjunto de resultados possíveis do experimento aleatório e denotamos por  $\Omega$ . A coleção de conjuntos de resultados de interesse é conhecida como  $\sigma$ -álgebra de eventos e denotada por  $\mathcal{F}$ , em que tais elementos são subconjuntos de  $\Omega$ . A probabilidade é um valor numérico  $0 \leq p \leq 1$  da chance de ocorrência de cada um dos conjuntos de resultados de interesse. Tal probabilidade é calculada com base numa função  $\mathcal{P} : \mathcal{F} \rightarrow \mathbb{R}$  chamada de medida de probabilidade. Vale ressaltar que  $\mathcal{F}$  e  $\mathcal{P}$  satisfazem algumas condições as quais não entraremos em detalhes. Para mais detalhes consultar [Magalhães \(2006\)](#). A terna  $(\Omega, \mathcal{F}, \mathcal{P})$  é conhecida como espaço de probabilidade.

**Definição 2.1** (Variável aleatória). *Seja  $(\Omega, \mathcal{F}, \mathcal{P})$  um espaço de probabilidade. Denominamos de variável aleatória (V.A.), qualquer função  $X : \Omega \rightarrow \mathbb{R}$  tal que*

$$X^{-1}(I) = \{\omega \in \Omega; X(\omega) \in I\} \in \mathcal{F},$$

*para todo intervalo  $I \subset \mathbb{R}$ . É comum denotar o conjunto  $X^{-1}(I)$  por  $[X \in I]$ .*

Assim, vemos que uma variável aleatória  $X$  é tal que sua imagem inversa de inter-

valores da reta real pertencem a  $\mathcal{F}$ . Em outras palavras, uma variável aleatória é uma função do espaço amostral nos reais, para a qual é possível calcular a probabilidade de ocorrência de seus valores.

**Definição 2.2** (Função de Distribuição Acumulada). *Seja  $X$  uma variável aleatória em um espaço de probabilidade  $(\Omega, \mathcal{F}, \mathcal{P})$ . Sua função de distribuição acumulada (f.d.a) é definida por*

$$F_X(x) = P(X \leq x) = P(\{\omega \in \Omega; X(\omega) \leq x\}),$$

em que  $x$  percorre todos os reais.

Segundo [Magalhães \(2006\)](#), conhecer a função de distribuição acumulada nos permite obter qualquer informação sobre a variável. Mesmo que a variável só assuma valores num subconjunto dos reais, a função de distribuição é definida em toda a reta. Assim, ressaltamos que, em alguns momentos, iremos nos referir a mesma como sendo apenas função de distribuição.

**Proposição 2.1** (Propriedades da Função de Distribuição). *Uma função de distribuição de uma variável  $X$  em  $(\Omega, \mathcal{F}, \mathcal{P})$  obedece às seguintes propriedades:*

(F1)  $x \leq y \Rightarrow F_X(x) \leq F_X(y)$ ,  $\forall x, y \in \mathbb{R}$ . Isto é,  $F_X$  é não-decrescente;

(F2)  $F_X$  é contínua à direita, ou seja, para  $h > 0$ ,  $\lim_{h \rightarrow 0^+} F_X(x + h) = F_X(x)$ ;

(F3)  $\lim_{x \rightarrow -\infty} F_X(x) = 0$  e  $\lim_{x \rightarrow +\infty} F_X(x) = 1$ .

*Demonstração.* Ver prova em [Magalhães \(2006, p. 66\)](#). □

**Teorema 2.1.** *Se uma função  $F : \mathbb{R} \rightarrow \mathbb{R}$ , satisfaz as propriedades (F1), (F2) e (F3), então existe uma variável aleatória  $X$  definida em algum espaço de probabilidade  $(\Omega, \mathcal{F}, \mathcal{P})$  tal que  $F$  é função de distribuição acumulada da variável aleatória  $X$ .*

*Demonstração.* Para uma ideia da prova, veja [Barry \(1981, p. 39\)](#). □

**Definição 2.3** (Função Densidade de Probabilidade). *Uma variável aleatória  $X$  em  $(\Omega, \mathcal{F}, \mathcal{P})$ , com função de distribuição  $F_X$ , será classificada como contínua, se existir uma função não negativa  $f_X$  tal que:*

$$F_X(x) = \int_{-\infty}^x f_X(t) dt, \quad \text{para todo } x \in \mathbb{R}.$$

A função  $f_X$  é denominada função densidade de probabilidade (f.d.p) ou somente função densidade.

A partir da definição acima, podemos notar a relação entre a função de distribuição e a função densidade. Além disso, vemos que a função de densidade define se a variável é contínua. Assim, se temos a função densidade, obtemos a função de distribuição por integração. Por outro lado, pelo teorema fundamental do cálculo, se derivarmos a função de distribuição, obtemos a densidade.

**Proposição 2.2** (Propriedades da Função Densidade). *A função densidade de  $X$  em  $(\Omega, \mathcal{F}, \mathcal{P})$  satisfaz:*

$$(f1) \quad f_X(x) \geq 0, \forall x \in \mathbb{R};$$

$$(f2) \quad \int_{-\infty}^{+\infty} f_X(x) dx = 1.$$

*Demonstração.* A demonstração pode ser consultada em [Magalhães \(2006, p. 68\)](#). □

Uma variável aleatória contínua possui função de distribuição acumulada contínua e, além disso, tal função é estritamente crescente no suporte  $S_X = \{x \in \mathbb{R}; f_X(x) > 0\}$ . Salientamos que em nosso trabalho focaremos em variáveis aleatórias contínuas.

**Definição 2.4** (Função Quantílica). *Seja  $F$  uma função de distribuição qualquer. A função quantílica da distribuição  $F$ , denotada por  $Q = F^{-1}$  é definida da seguinte forma:*

$$Q(y) = F^{-1}(y) = \inf\{x \in \mathbb{R} : F(x) \geq y\}$$

*Perceba, se a inversa de  $F$  existe no sentido usual, isto é,  $F^{-1}$  possui forma fechada, ela coincide com a função quantílica.*

**Definição 2.5.** *Seja  $X$  uma variável aleatória contínua com função de distribuição  $F$ . Dado uma probabilidade  $0 < \alpha < 1$ , o quantil  $1 - \alpha$  com base na distribuição de  $X$  é o valor  $q$ , tal que*

$$F(q) = P(X \leq q) = 1 - \alpha$$

*ou  $Q(1 - \alpha) = q$ , em que  $Q$  é função quantílica da distribuição  $F$ .*

Mais a diante, veremos como gerar números pseudo-aleatórios de uma distribuição  $F$ . Desta forma, devemos ter conhecimento da proposição a seguir que será muito útil na geração de tais números.

**Proposição 2.3.** *Seja  $X$  uma variável aleatória contínua com função de distribuição  $F$ . Então  $Y = F(X)$  tem distribuição Uniforme Contínua em  $[0, 1]$ . Vale a recíproca, sendo  $U$  uma variável aleatória contínua que segue distribuição uniforme no intervalo  $[0, 1]$ , então  $X = F^{-1}(U)$  tem função de distribuição  $F$ .*

*Demonstração.* A prova pode ser vista em [Magalhães \(2006, p. 151\)](#).  $\square$

**Definição 2.6** (Vetor aleatório). *Um vetor  $\mathbf{X} = (X_1, X_2, \dots, X_n)$  cujos componentes são variáveis aleatórias definidas no mesmo espaço de probabilidade  $(\Omega, \mathcal{F}, \mathcal{P})$ , é chamado vetor aleatório (ou variável aleatória  $n$ -dimensional).*

Este, por sua vez, possui função de distribuição denotada por  $F = F_{\mathbf{X}}$  e definida por

$$F_{\mathbf{X}}(x_1, x_2, \dots, x_n) = P(X_1 \leq x_1, \dots, X_n \leq x_n),$$

para todo  $(x_1, x_2, \dots, x_n)$  pertencente a  $\mathbb{R}^n$ . Tal função é chamada de função de distribuição conjunta das variáveis aleatórias  $X_1, X_2, \dots, X_n$ .

**Proposição 2.4** (Independência de V.A.'s).  $X_1, X_2, \dots, X_n$  são variáveis aleatórias independentes se, e somente se

$$F_{\mathbf{X}}(x_1, x_2, \dots, x_n) = \prod_{j=1}^n F_{X_j}(x_j), \quad \forall (x_1, x_2, \dots, x_n) \in \mathbb{R}^n$$

em que  $\mathbf{X} = (X_1, X_2, \dots, X_n)$ .

*Demonstração.* A demonstração pode ser consultada em [Barry \(1981, p. 60\)](#).  $\square$

**Proposição 2.5.** *Sejam  $X_1, X_2, \dots, X_n$  variáveis aleatórias i.i.d (independentes e identicamente distribuídas), isto é, além de independentes possuem mesma função de distribuição  $F(x)$ . As expressões da função de distribuição de  $X_{(1)} = \min(X_1, X_2, \dots, X_n)$  e  $X_{(n)} = \max(X_1, X_2, \dots, X_n)$  são dadas, nessa ordem, por*

$$F_{X_{(1)}}(x) = P(X_{(1)} \leq x) = 1 - [1 - F(x)]^n$$

e

$$F_{X_{(n)}}(x) = P(X_{(n)} \leq x) = [F(x)]^n.$$

*Demonstração.* A prova pode ser vista em [Magalhães \(2006, p. 159\)](#).  $\square$

## 2.2 Conceitos básicos na teoria da inferência estatística

Como já visto na introdução, a inferência estatística pode ser realizada de três maneiras diferentes: estimação pontual, estimação intervalar e teste de hipóteses. Vários métodos para estimação pontual de parâmetros foram propostas na literatura, mas o que mais se destaca é o método de máxima verossimilhança (MMV). O estimador de máxima verossimilhança (EMV) possui propriedades desejáveis e pode ser usado

na construção de intervalos e regiões de confiança e também nas estatísticas dos testes.

**Definição 2.7.** *Sejam  $X_1, X_2, \dots, X_n$  variáveis aleatórias i.i.d em que cada variável tem função de densidade de probabilidade  $f(x; \theta)$ , sendo  $\theta$  um vetor  $p$ -dimensional de parâmetros, ou seja,  $\theta = (\theta_1, \dots, \theta_p)^\top$  pertencente ao espaço paramétrico  $\Theta$ . Os valores observados das variáveis aleatórias são denotados como  $x_1, x_2, \dots, x_n$ . A função de verossimilhança é uma função do vetor de parâmetros  $\theta$ , correspondente à amostra observada dada por*

$$L(\theta; \mathbf{x}) = \prod_{i=1}^n f(x_i; \theta),$$

em que  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$ .

**Definição 2.8.** *O Estimador de máxima verossimilhança do parâmetro  $\theta$  é o vetor  $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_p)^\top \in \Theta$  que maximiza a função de verossimilhança  $L(\theta; \mathbf{x})$ , se esse vetor existir.*

O logaritmo natural da função de verossimilhança de  $\theta$  é chamado de log-verossimilhança e denotado por

$$l(\theta; \mathbf{x}) = \log[L(\theta; \mathbf{x})].$$

Vale ressaltar que o EMV também maximiza a função de log-verossimilhança. Para simplificar, em alguns momentos, faremos o uso de  $L(\theta)$  e  $l(\theta)$  para representar a verossimilhança e a log-verossimilhança, respectivamente.

Chamamos de hipótese estatística qualquer afirmação que se faça sobre um parâmetro populacional desconhecido. De forma geral, podemos dizer que é a partir de uma amostra da população que estabeleceremos uma regra de decisão segundo a qual rejeitaremos ou aceitaremos a hipótese proposta. Chamamos essa regra de decisão de teste. Geralmente, existe uma hipótese que é mais importante para o pesquisador, denotada por  $H_0$  e chamada hipótese nula. Qualquer que seja outra hipótese diferente de  $H_0$  será chamada de hipótese alternativa e será denotada por  $H_1$ .

Sendo  $X$  uma variável aleatória com f.d.p  $f(x; \theta)$ , em que  $\theta \in \Theta$ , dizemos que a distribuição de  $X$  está totalmente especificada quando conhecemos  $f(x; \theta)$  e  $\theta$ . Quando conhecemos a f.d.p, mas não  $\theta$ , dizemos que a distribuição de  $X$  está parcialmente especificada. Associados às hipóteses  $H_0$  e  $H_1$ , definimos os subconjuntos disjuntos  $\Theta_0$  e  $\Theta_1$  de  $\Theta$ , em que  $\Theta_0 \cup \Theta_1 = \Theta$ , ou seja,  $H_0$  afirma que  $\theta \in \Theta_0$  (notação:  $H_0 : \theta \in \Theta_0$ ) e  $H_1$  afirma que  $\theta \in \Theta_1$  (notação:  $H_1 : \theta \in \Theta_1$ ). Quando  $\dim(\Theta_0) = 1$  dizemos que  $H_0$  é simples. Caso contrário, isto é,  $\dim(\Theta_0) > 1$ , dizemos que  $H_0$  é composta. O mesmo vale para a hipótese alternativa.



Um subconjunto que contém os valores de  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$  para os quais  $H_0$  será rejeitada é chamado região crítica do teste, e o outro que contém os valores de  $\mathbf{x}$  para os quais  $H_0$  será aceita é chamado região de aceitação do teste. De um modo geral, um teste fica determinado quando especificamos sua região crítica.

Todo teste está associado à algum tipo de erro. Assim, existem dois tipos,  $\gamma$  e  $\beta$ , que são conhecidos na literatura como probabilidades dos erros dos tipos I e II, respectivamente. Além disso,  $\gamma$  é também chamado de nível de significância do teste ou nível nominal. Associado a decisão de rejeitar ou não rejeitar a hipótese nula ( $H_0$ ), podemos citar as seguintes consequências dadas na tabela abaixo:

Tabela 2.1: Tipos de erros em testes de hipóteses.

| Decisão            | $H_0$ é verdadeira       | $H_0$ é falsa            |
|--------------------|--------------------------|--------------------------|
| Não rejeitar $H_0$ | Decisão Correta          | Erro tipo II ( $\beta$ ) |
| rejeitar $H_0$     | Erro tipo I ( $\gamma$ ) | Decisão Correta          |

Vale ressaltar, que antes de iniciar qualquer teste de hipótese é preciso especificar um nível de significância  $\gamma$ , isto é, quanto estamos dispostos a cometer o erro tipo I.

**Definição 2.9** (Poder do teste). *O poder do teste é a probabilidade de rejeitar corretamente a hipótese nula quando esta é falsa sob a hipótese alternativa, ou seja*

$$\pi(\theta) = 1 - \beta = 1 - P(\text{erro tipo II}).$$

**Definição 2.10** (Função do Poder). *A função  $\pi(\theta)$  é chamada função de poder do teste. Assim, denotando por  $\mathbf{A}_1^*$  a região crítica, a função de poder é definida como*

$$\pi(\theta) = P_\theta(\mathbf{x} \in \mathbf{A}_1^* | \theta), \quad \forall \theta \in \Theta.$$

*Ou seja, é a probabilidade de rejeitar  $H_0$  para  $\theta \in \Theta$ . Perceba que  $\pi(\theta_0) = \gamma$ , em que  $\theta_0$  é vetor de parâmetros especificado em  $H_0$ .*

Um procedimento geral que produz testes razoáveis e que pode ser utilizado em muitos casos, sem muita dificuldade, é o Teste da Razão de Verossimilhanças (TRV). (DEKKING *et al.*, 2005; LEHMANN e ROMANO, 2006; ROUSSAS, 2003; WASSERMAN, 2013)

**Definição 2.11** (TRV). *O TRV para testar  $H_0 : \theta \in \Theta_0$  versus  $H_1 : \theta \in \Theta_1$  pode ser definido como o teste com região crítica dada por*

$$\mathbf{A}_1^* = \left\{ \mathbf{x}; \lambda^*(\mathbf{x}) = \frac{\sup_{\theta \in \Theta_0} L(\theta; \mathbf{x})}{\sup_{\theta \in \Theta} L(\theta; \mathbf{x})} \leq c \right\}$$

*em que  $\lambda^*(\mathbf{x})$  é conhecida como a estatística da razão de verossimilhanças.*

Notemos que  $0 \leq \lambda^*(\mathbf{x}) \leq 1$ , pois seu numerador é o valor máximo atingido pela função de verossimilhança quando  $\theta \in \Theta_0$  ( $\Theta_0 \subset \Theta$ ), enquanto o denominador é o valor máximo atingido quando  $\theta$  pertence a todo espaço paramétrico, isto é,  $\theta \in \Theta$ .

Desta forma, dado um nível de significância  $\gamma$ , é razoável decidir rejeitar  $H_0$  se  $\lambda^*(\mathbf{x}) \leq c$ , em que a constante  $c$  é obtida de modo que

$$\sup_{\theta \in \Theta_0} P(\lambda^*(\mathbf{x}) \leq c) = \gamma.$$

Para isso, precisamos conhecer a verdadeira distribuição da estatística  $\lambda^*(\mathbf{x})$  que, geralmente, não é fácil de ser obtida. Note que,  $c$  é o quantil  $\gamma$  com base na distribuição de  $\lambda^*(\mathbf{x})$ . Logo, rejeitaremos a hipótese nula se o valor da estatística  $\lambda^*(\mathbf{x})$  baseado nos dados amostrais  $\mathbf{x}$  for menor ou igual ao quantil  $\gamma$  obtido através da distribuição de  $\lambda^*(\mathbf{x})$ .

A seguir, mostraremos um resultado bastante conhecido na literatura, que nos ajuda nos casos em que não conhecemos a verdadeira distribuição de  $\lambda^*(\mathbf{x})$ . Mas, vale ressaltar que o mesmo é aplicado quando o nosso teste de hipóteses é feito com base em uma amostra de tamanho grande.

**Teorema 2.2.** *Suponhamos que o vetor paramétrico está particionado como  $\theta = (\psi_1^\top, \psi_2^\top)^\top$ , em que  $\dim(\psi_1) + \dim(\psi_2) = \dim(\theta)$ . Aqui,  $\psi_1$  representa o vetor de parâmetros de interesse e  $\psi_2$  o vetor de parâmetros de perturbação. A estatística da razão de verossimilhanças (RV) para testar a hipótese nula  $H_0 : \psi_1 = \psi_1^{(0)}$  versus a hipótese alternativa  $H_1 : \psi_1 \neq \psi_1^{(0)}$  é dada por*

$$\begin{aligned} RV = -2 \log(\lambda^*(\mathbf{x})) &= -2 \log \left( \frac{L(\tilde{\theta})}{L(\hat{\theta})} \right) = -2 \{ \log(L(\tilde{\theta})) - \log(L(\hat{\theta})) \} \\ &= 2 \{ l(\hat{\theta}) - l(\tilde{\theta}) \}, \end{aligned}$$

em que  $\tilde{\theta} = (\psi_1^{(0)\top}, \tilde{\psi}_2^\top)^\top$  e  $\hat{\theta} = (\hat{\psi}_1^\top, \hat{\psi}_2^\top)^\top$  são os estimadores de máxima verossimilhança sob a hipótese nula e todo espaço paramétrico, respectivamente. Baseado na teoria assintótica de primeira ordem, sabemos que a estatística RV segue distribuição assintoticamente qui-quadrado  $\chi_k^2$  com  $k$  graus de liberdade, em que  $k = \dim(\psi_1)$ .

*Demonstração.* A prova deste resultado pode ser consultada em Severini (2000, p. 114–115).  $\square$

Desta forma, através do TRV, dado na Definição 2.11, a região crítica para testar

$H_0 : \theta \in \Theta_0$  versus  $H_1 : \theta \in \Theta_1$  pode ser dada por

$$\mathbf{B}_1^* = \{\mathbf{x}; RV > q\}.$$

Assim, dado um nível de significância  $\gamma$ , é plausível decidir rejeitar  $H_0$  se  $RV > q$  em que a constante  $q$  é obtida de modo que

$$P(\chi_k^2 > q) = \gamma.$$

Note que, fixado uma probabilidade  $\gamma$ , o número  $q$  é o quantil  $1 - \gamma$  obtido da distribuição qui-quadrado com  $k$  graus de liberdade, em que  $k$  é dimensão do vetor de parâmetros sob a hipótese nula. Neste trabalho, chamaremos de TRV o teste da razão de verossimilhanças que utiliza o quantil da distribuição qui-quadrado para determinar a região crítica do teste, isto é,  $\mathbf{B}_1^*$ .

Veja que, este resultado é válido para  $n$  (tamanho da amostra) grande. No entanto, o uso da distribuição qui-quadrado como uma aproximação à verdadeira distribuição da estatística  $RV$ , pode levar a inferências imprecisas quando o tamanho da amostra é pequeno. ([CORDEIRO e CRIBARI-NETO, 2014](#))

## Capítulo 3

# Famílias de distribuições sup e inf

### 3.1 O método $T-X$ na geração de novas famílias de distribuições contínuas

Visando uma melhor compreensão do leitor para com o nosso trabalho, destinamos esta seção para abordarmos o método  $T-X$  ([ALZAATREH et al., 2013](#)), o qual foi utilizado em [Tablada \(2017\)](#) para construir as famílias de distribuições generalizadas sup e inf que serão utilizadas em nosso trabalho.

Seja  $r(t)$  a f.d.p de uma variável aleatória  $T$ , tendo como suporte o intervalo  $[a, b]$ , para  $-\infty \leq a < b \leq \infty$ . Seja  $W(G(x))$  a composição da f.d.a  $G(x)$  de qualquer variável aleatória  $X$  com uma função  $W$ , de maneira que  $W(G(x))$  satisfaça as seguintes condições:

$$\begin{aligned} W(G(x)) &\text{ tem como suporte o intervalo } [a, b]; \\ W(G(x)) &\text{ é diferenciável e monótona não-decrescente;} \\ W(G(x)) &\rightarrow a \text{ quando } x \rightarrow -\infty \text{ e } W(G(x)) \rightarrow b \text{ quando } x \rightarrow \infty. \end{aligned} \tag{3.1}$$

Desta forma, dadas as condições anteriores, o método  $T-X$  é definido a seguir.

**Definição 3.1.** *Seja  $X$  uma variável aleatória com f.d.p  $g(x)$  e f.d.a  $G(x)$ . Seja  $T$  uma variável aleatória definida em  $[a, b]$  com f.d.p  $r(t)$  e f.d.a  $R(t)$ . A f.d.a de uma*

nova família de distribuições é definida como

$$\begin{aligned}
H(x) &= \int_a^{W(G(x))} r(t) dt = R(t) \Big|_a^{W(G(x))} \\
&= R(W(G(x))) - R(a) \\
&= P\{T \leq W(G(x))\} - P\{T \leq a\} \\
&= P\{T \leq W(G(x))\} - 0 \\
&= R(W(G(x)))
\end{aligned} \tag{3.2}$$

em que  $W(G(x))$  satisfaz as condições em (3.1). Daí, segue que a f.d.p, denotada por  $h(x)$ , associada a função de distribuição  $H(x)$  é dada por

$$h(x) = \frac{d}{dx} [R(W(G(x)))] = r(W(G(x))) \frac{d}{dx} [W(G(x))]$$

Em seguida, [Alzaatreh, Lee, e Famoye \(2013\)](#) observam que a nova família de distribuições em (3.2) é uma composição  $(R \circ W \circ G)(x)$ , em que a f.d.p  $r(t)$  em (3.2) é “transformada” em uma nova f.d.a  $H(x)$  através da função,  $W(G(x))$ , que atua como um “transformador”.

**Observação 3.1.** A variável aleatória  $X$  pode ser discreta e, nesse caso,  $H(x)$  é a f.d.a. de uma família de distribuições discretas.

**Observação 3.2.** A definição de  $W(G(x))$  depende do suporte da variável aleatória  $T$ . A seguir os autores deste método expõem alguns exemplos de  $W(\cdot)$ .

- Quando o suporte de  $T$  é limitado: sem perda de generalidade, assumimos o suporte de  $T$  como sendo o intervalo  $[0, 1]$ . As distribuições para tal  $T$ , incluem  $\mathcal{U}(0, 1)$ , beta, Kumaraswamy e outros tipos de distribuições beta generalizadas. Assim, sob essas condições,  $W(G(x))$  pode ser definida como  $G(x)$  ou  $G^\alpha(x)$ , em que  $\alpha > 0$ .
- Quando o suporte de  $T$  é  $[a, \infty)$ ,  $a \geq 0$ : sem perda de generalidade, assumimos  $a = 0$ .  $W(G(x))$  pode ser definido como  $-\log(1 - G(x))$ ,  $G(x)/(1 - G(x))$ ,  $-\log(1 - G^\alpha(x))$ , e  $G^\alpha(x)/(1 - G^\alpha(x))$ , em que  $\alpha > 0$ .
- Quando o suporte de  $T$  é  $(-\infty, \infty)$ :  $W(G(x))$  pode ser definido como  $\log[-\log(1 - G(x))]$ ,  $\log[G(x)/(1 - G(x))]$ ,  $\log[-\log(1 - G^\alpha(x))]$ , e  $\log[G^\alpha(x)/(1 - G^\alpha(x))]$ .

De forma resumida, o método consiste em: dado  $R$  e  $G$ , funções de distribuição das variáveis aleatórias  $T$  e  $X$ , respectivamente, uma nova família de distribuições

pode ser definida como uma composição  $(R \circ W \circ G)(x)$ , em que  $(W \circ G)(x)$  deve satisfazer as três condições dadas em 3.1, sendo  $W(\cdot)$  uma função que é usada para ligar o suporte de  $T$  ao de  $X$ .

## 3.2 Famílias $G_c^{\sup}$ e $G_c^{\inf}$

### 3.2.1 Obtendo as Famílias $G_c^{\sup}$ e $G_c^{\inf}$

Destinamos esta seção para fazer um estudo detalhado e mostrarmos todo o processo de construção e as propriedades que as famílias de distribuições sup e inf possuem. Essas famílias são obtidas em [Tablada \(2017\)](#) e são a base para nossa proposta de trabalho.

Inicialmente, definiremos duas funções  $R_c^{\sup}(z)$  e  $R_c^{\inf}(z)$  e provaremos que existem variáveis aleatórias  $T^{\sup}$  e  $T^{\inf}$ , tais que  $R_c^{\sup}(z)$  e  $R_c^{\inf}(z)$  são funções de distribuição destas variáveis aleatórias, respectivamente. Posteriormente, definiremos  $W(G(x))$  a depender dos suportes de  $T^{\sup}$  e  $T^{\inf}$ , em quem  $G(x)$  é a função de distribuição de base de uma variável aleatória  $X$ .

Definimos as funções a seguir

$$R_c^{\sup}(z) = \begin{cases} 0 & , z < 0 \\ z + z^c - z^{c+1} & , 0 \leq z \leq 1, c > 0 \text{ constante.} \\ 1 & , z > 1 \end{cases} \quad (3.3)$$

$$R_c^{\inf}(z) = \begin{cases} 0 & , z < 0 \\ z[1 - (1 - z)^c] & , 0 \leq z \leq 1, c > 0 \text{ constante.} \\ 1 & , z > 1 \end{cases} \quad (3.4)$$

Provaremos agora, que ambas são funções de distribuição acumulada. Para isso, verificaremos se estas satisfazem as três condições (F1), (F2) e (F3) enunciadas na proposição 2.1. Desta forma, provaremos inicialmente para a função (3.3).

(F1) Note que

$$\begin{aligned} \frac{d}{dz} R_c^{\sup}(z) &= 1 + cz^{c-1} - (c+1)z^c \\ &= 1 + cz^{c-1} - z^c - cz^c \\ &= c(z^{c-1} - z^c) + 1 - z^c. \end{aligned}$$

Temos que  $z^{c-1} \geq z^c$ , o qual equivale a  $z^{c-1} - z^c \geq 0$ . Além disso, se  $0 \leq z^c \leq 1$ ,

então  $0 \leq 1 - z^c \leq 1$ . Logo,

$$\frac{d}{dz} R_c^{\text{sup}}(z) \geq 0, \quad \forall z \in [0, 1].$$

Daí, segue  $R_c^{\text{sup}}$  é monótona não-decrescente.

- (F2) Veja que  $R_c^{\text{sup}}$  é diferenciável no intervalo  $(0, 1)$ ,  $R_c^{\text{sup}}(0) = 0$  e  $R_c^{\text{sup}}(1) = 1$ . Logo é contínua, e em particular é contínua à direita.
- (F3) Inicialmente perceba que, como  $R_c^{\text{sup}}$  é monótona não-decrescente, se  $0 \leq z \leq 1$ , então

$$\begin{aligned} R_c^{\text{sup}}(0) &\leq R_c^{\text{sup}}(z) \leq R_c^{\text{sup}}(1) \\ 0 &\leq R_c^{\text{sup}}(z) \leq 1, \quad \forall z \in [0, 1]. \end{aligned}$$

Agora, veja que

$$\lim_{z \rightarrow -\infty} R_c^{\text{sup}}(z) = \lim_{z \rightarrow 0^+} R_c^{\text{sup}}(z) = \lim_{z \rightarrow 0} R_c^{\text{sup}}(z).$$

Pela continuidade de  $R_c^{\text{sup}}$ , temos que

$$\lim_{z \rightarrow 0} R_c^{\text{sup}}(z) = R_c^{\text{sup}}(0) = 0.$$

Por outro lado,

$$\lim_{z \rightarrow +\infty} R_c^{\text{sup}}(z) = \lim_{z \rightarrow 1^-} R_c^{\text{sup}}(z) = \lim_{z \rightarrow 1} R_c^{\text{sup}}(z).$$

Novamente, pela continuidade de  $R_c^{\text{sup}}$ , temos que

$$\lim_{z \rightarrow 1} R_c^{\text{sup}}(z) = R_c^{\text{sup}}(1) = 1.$$

logo, as condições (F1), (F2) e (F3) estão satisfeitas. Portanto, pelo Teorema (2.1) existe uma variável aleatória  $T^{\text{sup}}$  definida em algum espaço de probabilidade  $(\Omega, \mathcal{F}, \mathcal{P})$ , tal que  $R_c^{\text{sup}}(z)$  é função de distribuição da variável aleatória  $T^{\text{sup}}$ , ou seja,  $R_c^{\text{sup}}(z) = P(T^{\text{sup}} \leq z)$ .

Provaremos agora para a função (3.4).

(F1) De fato,

$$\begin{aligned}
0 \leq x \leq y \leq 1 &\Rightarrow 1 - x \geq 1 - y \Rightarrow (1 - x)^c \geq (1 - y)^c \\
&\Rightarrow 1 - (1 - x)^c \leq 1 - (1 - y)^c \\
&\Rightarrow x(1 - (1 - x)^c) \leq x(1 - (1 - y)^c) \leq y(1 - (1 - y)^c) \\
&\Rightarrow R_c^{\text{inf}}(x) \leq R_c^{\text{inf}}(y).
\end{aligned}$$

Logo,  $R_c^{\text{inf}}$  é monótona não-decrescente.

(F2) Perceba que  $R_c^{\text{inf}}$  é contínua, e em particular é contínua á direita.

(F3) Inicialmente note que  $R_c^{\text{inf}}$  é limitada. Com efeito,

$$\begin{aligned}
0 \leq z \leq 1 &\Leftrightarrow 0 \leq 1 - z \leq 1 \Leftrightarrow 0 \leq (1 - z)^c \leq 1 \\
&\Leftrightarrow 0 \leq 1 - (1 - z)^c \leq 1 \Leftrightarrow 0 \leq z(1 - (1 - z)^c) \leq z \leq 1 \\
&\Leftrightarrow 0 \leq R_c^{\text{inf}}(z) \leq 1, \quad \forall z \in [0, 1].
\end{aligned}$$

Agora, veja que

$$\lim_{z \rightarrow -\infty} R_c^{\text{inf}}(z) = \lim_{z \rightarrow 0^+} R_c^{\text{inf}}(z) = \lim_{z \rightarrow 0} R_c^{\text{inf}}(z).$$

Pela continuidade de  $R_c^{\text{inf}}$ , temos que

$$\lim_{z \rightarrow 0} R_c^{\text{inf}}(z) = R_c^{\text{inf}}(0) = 0.$$

Por outro lado,

$$\lim_{z \rightarrow +\infty} R_c^{\text{inf}}(z) = \lim_{z \rightarrow 1^-} R_c^{\text{inf}}(z) = \lim_{z \rightarrow 1} R_c^{\text{inf}}(z).$$

Novamente, pela continuidade de  $R_c^{\text{inf}}$ , temos que

$$\lim_{z \rightarrow 1} R_c^{\text{inf}}(z) = R_c^{\text{inf}}(1) = 1.$$

logo, as condições (F1), (F2) e (F3) estão satisfeitas. Portanto, pelo Teorema 2.1 existe uma variável aleatória  $T^{\text{inf}}$  definida em algum espaço de probabilidade  $(\Omega, \mathcal{F}, \mathcal{P})$ , tal que  $R_c^{\text{inf}}(z)$  é função de distribuição acumulada da variável aleatória  $T^{\text{inf}}$ , ou seja,  $R_c^{\text{inf}}(z) = P(T^{\text{inf}} \leq z)$ .

Assim, provamos que  $R_c^{\text{sup}}$  e  $R_c^{\text{inf}}$  são funções de distribuição acumuladas das variáveis aleatórias  $T^{\text{sup}}$  e  $T^{\text{inf}}$ , respectivamente. Então, as funções de densidades



destas, são dadas, nessa ordem, por

$$r_c^{\text{sup}}(z) = \frac{d}{dz} R_c^{\text{sup}}(z) = \begin{cases} 0 & , \quad z < 0 \text{ ou } z > 1 \\ 1 + cz^{c-1} - (c+1)z^c, & 0 \leq z \leq 1 \end{cases}$$

e

$$r_c^{\text{inf}}(z) = \frac{d}{dz} R_c^{\text{inf}}(z) = \begin{cases} 0 & , \quad z < 0 \text{ ou } z > 1 \\ 1 - (1-z)^c + zc(1-z)^{c-1}, & 0 \leq z \leq 1 \end{cases}.$$

Sabendo que o suporte de  $r_c^{\text{sup}}$  e  $r_c^{\text{inf}}$  é o intervalo  $[0, 1]$ , podemos definir  $W(G(x)) = G(x)$ , como visto no primeiro ponto da Observação 3.2. Denotamos por  $G(x)$  a função de distribuição acumulada de base de uma variável aleatória  $X$ . Posto isto, definimos, fazendo o uso da Definição 3.1 do método  $T$ - $X$  visto na Seção 3.1, novas famílias de distribuições generalizadas. Estas são dadas por

$$\begin{aligned} G_c^{\text{sup}}(x) &= \int_0^{G(x)} r_c^{\text{sup}}(t) dt = R_c^{\text{sup}}(t) \Big|_0^{G(x)} \\ &= R_c^{\text{sup}}(G(x)) - R_c^{\text{sup}}(0) = R_c^{\text{sup}}(G(x)) \\ &= G(x) + G^c(x) - G^{c+1}(x). \end{aligned}$$

e

$$\begin{aligned} G_c^{\text{inf}}(x) &= \int_0^{G(x)} r_c^{\text{inf}}(t) dt = R_c^{\text{inf}}(t) \Big|_0^{G(x)} \\ &= R_c^{\text{inf}}(G(x)) - R_c^{\text{inf}}(0) = R_c^{\text{inf}}(G(x)) \\ &= G(x)[1 - (1 - G(x))^c]. \end{aligned} \tag{3.5}$$

Daí, segue que as funções de densidade de probabilidade de  $G_c^{\text{sup}}$  e  $G_c^{\text{inf}}$ , são respectivamente,

$$\begin{aligned} g_c^{\text{sup}}(x) &= \frac{d}{dx} G_c^{\text{sup}}(x) = \frac{d}{dx} G(x) + cG^{c-1}(x) \frac{d}{dx} G(x) - (c+1)G^c(x) \frac{d}{dx} G(x) \\ &= g(x) + cG^{c-1}(x)g(x) - (c+1)G^c(x)g(x) \\ &= g(x)[1 + cG^{c-1}(x) - (c+1)G^c(x)] \end{aligned}$$

e

$$\begin{aligned}
g_c^{\inf}(x) &= \frac{d}{dx} G_c^{\inf}(x) = \frac{d}{dx} G(x)[1 - (1 - G(x))^c] + G(x) \frac{d}{dx} [1 - (1 - G(x))^c] \\
&= g(x)[1 - (1 - G(x))^c] + G(x)c(1 - G(x))^{c-1}g(x) \\
&= g(x)[1 - (1 - G(x))^c + cG(x)(1 - G(x))^{c-1}] \quad (3.6)
\end{aligned}$$

em que  $g$  é f.d.p da variável aleatória  $X$  que segue distribuição  $G$ .

Portanto, terminamos de definir as duas famílias de distribuições de probabilidade, tal como foi comentado no início desta seção. A partir de agora, sabemos que dado uma função de distribuição de base  $G$ , de uma variável aleatória contínua  $X$  com suporte  $\mathcal{X}$ , podemos obter duas novas famílias de distribuições  $G_c^{\sup}$  e  $G_c^{\inf}$ , adicionando-se um parâmetro de forma  $c > 0$  a distribuição  $G$ . Além disso, vale ressaltar que as novas famílias de distribuições possuem também suporte  $\mathcal{X}$ .

### 3.2.2 Propriedades das famílias $G_c^{\sup}$ e $G_c^{\inf}$

Após realizarmos todo o processo construtivo de tais famílias, mostraremos, a partir de agora, algumas das propriedades que estas possuem. Além disso, após realizarmos estudos mais a fundo sobre estas, obtemos três resultados que contribuem de forma significativa tanto para nosso estudo, quanto para tais famílias. Tais propriedades são enunciadas a seguir como proposições e demonstradas.

**Proposição 3.1.** *Seja  $X$  uma variável aleatória que segue distribuição de base  $G$ . Então,*

$$G_c^{\inf}(x) \leq G(x) \leq G_c^{\sup}(x), \quad \forall x \in \mathbb{R}.$$

*Demonstração.* Inicialmente, lembre que para todo  $x$  temos  $0 \leq G(x) \leq 1$ . Assim, para  $c > 0$  segue que  $0 \leq 1 - (1 - G(x))^c \leq 1$ . Logo,

$$G(x)[1 - (1 - G(x))^c] \leq G(x) \quad \Leftrightarrow \quad G_c^{\inf}(x) \leq G(x), \quad \forall x.$$

Por outro lado, perceba que para  $c > 0$  temos que  $G^c(x) \geq G^{c+1}(x)$ , o que equivale a  $G^c(x) - G^{c+1}(x) \geq 0$ . Desta forma,

$$G(x) \leq G(x) + G^c(x) - G^{c+1}(x) \quad \Leftrightarrow \quad G(x) \leq G_c^{\sup}(x), \quad \forall x$$

Portanto,

$$G_c^{\inf}(x) \leq G(x) \leq G_c^{\sup}(x), \quad \forall x$$

□

**Proposição 3.2.** *Seja  $X$  uma variável aleatória que segue distribuição de base  $G$ . Então para todo  $j > 0$ ,*

$$G_{c+j}^{\sup}(x) \leq G_c^{\sup}(x) \quad (3.7)$$

e

$$G_{c+j}^{\inf}(x) \geq G_c^{\inf}(x) \quad (3.8)$$

para todo  $x \in \mathbb{R}$ . Isto é  $G_c^{\sup}$  e  $G_c^{\inf}$  são, nessa ordem, famílias monótonas não-crescentes e não-decrescentes em relação ao parâmetro  $c$ .

*Demonstração.* Inicialmente provaremos (3.7). Sabemos que

$$G_c^{\sup}(x) = G(x) + G^c(x) - G^{c+1}(x) \quad \text{e} \quad G_{c+j}^{\sup}(x) = G(x) + G^{c+j}(x) - G^{c+j+1}(x).$$

Temos que  $0 \leq G(x) \leq 1$ . Assim, para todo  $j > 0$ ,

$$\begin{aligned} G^{c+j}(x) &\leq G^c(x) \\ \Rightarrow G^{c+j}(x)[1 - G(x)] &\leq G^c(x)[1 - G(x)] \\ \Rightarrow G^{c+j}(x) - G^{c+j+1}(x) &\leq G^c(x) - G^{c+1}(x) \\ \Rightarrow G(x) + G^{c+j}(x) - G^{c+j+1}(x) &\leq G(x) + G^c(x) - G^{c+1}(x) \\ \Rightarrow G_{c+j}^{\sup}(x) &\leq G_c^{\sup}(x), \quad \forall j > 0. \end{aligned}$$

Prova de (3.8): Sabemos que

$$G_c^{\inf}(x) = G(x)[1 - (1 - G(x))^c] \quad \text{e} \quad G_{c+j}^{\inf}(x) = G(x)[1 - (1 - G(x))^{c+j}].$$

Temos que  $0 \leq 1 - G(x) \leq 1$ . Assim, para todo  $j > 0$

$$\begin{aligned} (1 - G(x))^{c+j} &\leq (1 - G(x))^c \\ \Rightarrow 1 - (1 - G(x))^{c+j} &\geq 1 - (1 - G(x))^c \\ \Rightarrow G(x)[1 - (1 - G(x))^{c+j}] &\geq G(x)[1 - (1 - G(x))^c] \\ \Rightarrow G_{c+j}^{\inf}(x) &\geq G_c^{\inf}(x), \quad \forall j > 0. \end{aligned}$$

□

**Proposição 3.3.** *Seja  $X$  uma variável aleatória que segue distribuição de base  $G$ . Então*

$$\lim_{c \rightarrow \infty} G_c^{\sup}(x) = G(x) \quad (3.9)$$

e

$$\lim_{c \rightarrow \infty} G_c^{\inf}(x) = G(x) \quad (3.10)$$

*Demonstração.* Provaremos inicialmente (3.9). Temos que

$$G_c^{\text{sup}}(x) = G(x) + G^c(x) - G^{c+1}(x)$$

em que  $0 \leq G(x) \leq 1$ .

Se  $G(x) = 0$ , então  $G_c^{\text{sup}}(x) = 0$ . Logo, o limite é válido.

Se  $G(x) = 1$ , então  $G_c^{\text{sup}}(x) = 1$ . Daí segue que o limite também é válido.

Se  $0 < G(x) < 1$ , então

$$\begin{aligned} \lim_{c \rightarrow \infty} G_c^{\text{sup}}(x) &= \lim_{c \rightarrow \infty} [G(x) + G^c(x) - G^{c+1}(x)] \\ &= G(x) + \lim_{c \rightarrow \infty} G^c(x) - \lim_{c \rightarrow \infty} G^{c+1}(x) \\ &= G(x) + 0 - 0 = G(x). \end{aligned}$$

Prova de (3.10): Temos que

$$G_c^{\text{inf}}(x) = G(x)[1 - (1 - G(x))^c]$$

em que  $0 \leq G(x) \leq 1$ . Se  $G(x) = 0$ , então  $G_c^{\text{inf}}(x) = 0$ . Logo, o limite é válido.

Se  $G(x) = 1$ , então  $G_c^{\text{inf}}(x) = 1$ . Daí segue, que o limite também é válido.

Se  $0 < G(x) < 1$ , então

$$\begin{aligned} \lim_{c \rightarrow \infty} G_c^{\text{inf}}(x) &= \lim_{c \rightarrow \infty} \{G(x)[1 - (1 - G(x))^c]\} \\ &= G(x) \lim_{c \rightarrow \infty} [1 - (1 - G(x))^c] \\ &= G(x)[1 - \lim_{c \rightarrow \infty} (1 - G(x))^c]. \end{aligned}$$

Como  $0 < G(x) < 1$ , então  $0 < (1 - G(x))^c < 1$ . Assim,

$$\lim_{c \rightarrow \infty} (1 - G(x))^c = 0.$$

Portanto,  $\lim_{c \rightarrow \infty} G_c^{\text{inf}}(x) = G(x)$ . □

**Corolário 3.1.** *Seja  $X_1, X_2, \dots, X_n, \dots$  sequência de variáveis aleatórias.*

(i) *Se  $X_n \sim G_n^{\text{sup}}$  para todo  $n$ , então  $X_n$  converge em distribuição para  $X$ , em que  $X \sim G$*

(ii) *Se  $X_n \sim G_n^{\text{inf}}$  para todo  $n$ , então  $X_n$  converge em distribuição para  $X$ , em que  $X \sim G$ .*

*Demonstração.* Pela proposição 3.3 temos que, para  $c > 0$

$$\lim_{c \rightarrow \infty} G_c^{\text{sup}}(x) = G(x) \quad \text{e} \quad \lim_{c \rightarrow \infty} G_c^{\text{inf}}(x) = G(x)$$

Assim, para  $n \in \mathbb{N}$  vale

$$\lim_{n \rightarrow \infty} G_n^{\sup}(x) = G(x) \quad \text{e} \quad \lim_{n \rightarrow \infty} G_n^{\inf}(x) = G(x)$$

Perceba que isto é válido para todo ponto  $x$  de continuidade de  $G$ , pois  $G$ ,  $G_c^{\sup}$  e  $G_c^{\inf}$  possuem o mesmo suporte. Portanto,  $X_n$  converge em distribuição para  $X$  em ambos os casos.  $\square$

Todos os resultados enunciados até o momento são provados em [Tablada \(2017\)](#). As proposições a seguir, [3.4](#), [3.5](#) e [3.6](#) são contribuições teóricas deste trabalho.

**Proposição 3.4.** *Seja  $X$  uma variável aleatória contínua com f.d.a de base  $G(x)$  e f.d.p  $g(x)$ . Para  $0 < G(x) < 1$ , existe  $A > 0$  tal que  $g_c^{\sup}(x) \leq g(x) \leq g_c^{\inf}(x)$ , sempre que  $x > A$ .*

*Demonstração.* Sabemos que

$$\begin{aligned} g_c^{\sup}(x) &= g(x)[1 + cG^{c-1}(x) - (c+1)G^c(x)] \\ &= g(x)[1 - G^c(x) + c(G^{c-1}(x) - G^c(x))]. \end{aligned}$$

Pela condição  $F3$  dada na proposição [2.1](#), temos que  $\lim_{x \rightarrow +\infty} G(x) = 1$ . Daí, vem que

$$\lim_{x \rightarrow +\infty} (1 - G^c(x)) = 0 \tag{3.11}$$

e

$$\lim_{x \rightarrow +\infty} (G^{c-1}(x) - G^c(x)) = 0. \tag{3.12}$$

Assim, pela definição de limite no infinito, segue de [\(3.11\)](#) e [\(3.12\)](#), respectivamente, que dado  $0 < \varepsilon < 1$

$$\exists N_1 > 0; \ x > N_1 \quad \Rightarrow \quad |1 - G^c(x)| < \frac{\varepsilon}{2}; \tag{3.13}$$

$$\exists N_2 > 0; \ x > N_2 \quad \Rightarrow \quad |G^{c-1}(x) - G^c(x)| < \frac{\varepsilon}{2c}. \tag{3.14}$$

Como  $0 < 1 - G^c(x) < 1$  e  $G^{c-1}(x) > G^c(x)$ , fazendo  $N = \max\{N_1, N_2\}$  segue de [\(3.13\)](#) e [\(3.14\)](#), que dado  $0 < \varepsilon < 1$

$$\begin{aligned} \exists N > 0; \ x > N &\Rightarrow 1 - G^c(x) + c(G^{c-1}(x) - G^c(x)) < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon \\ &\Rightarrow g(x)[1 - G^c(x) + c(G^{c-1}(x) - G^c(x))] < g(x)\varepsilon \leq g(x) \\ &\Rightarrow g_c^{\sup} \leq g(x), \quad \forall x > N. \end{aligned} \tag{3.15}$$

Por outro lado, temos que

$$g_c^{\inf}(x) = g(x)[1 - (1 - G(x))^c + cG(x)(1 - G(x))^{c-1}].$$

Agora, perceba que

$$\lim_{x \rightarrow +\infty} \frac{G(x)}{1 - G(x)} = +\infty$$

ou seja, dado  $M_1 > 0$ , existe  $M > 0$  tal que

$$\left| \frac{G(x)}{1 - G(x)} \right| \geq M_1, \quad \forall x > M.$$

Seja  $M_1 \geq 1/c$ . Logo, existe  $M > 0$  tal que, para todo  $x > M$ , vale

$$\begin{aligned} \frac{G(x)}{1 - G(x)} \geq M_1 \geq \frac{1}{c} &\Leftrightarrow \frac{cG(x)}{1 - G(x)} \geq 1 \\ &\Leftrightarrow cG(x)(1 - G(x))^{c-1} \geq (1 - G(x))^c \\ &\Leftrightarrow cG(x)(1 - G(x))^{c-1} - (1 - G(x))^c \geq 0 \\ &\Leftrightarrow 1 + cG(x)(1 - G(x))^{c-1} - (1 - G(x))^c \geq 1 \\ &\Leftrightarrow g(x)[1 + cG(x)(1 - G(x))^{c-1} - (1 - G(x))^c] \geq g(x) \\ &\Leftrightarrow g_c^{\inf}(x) \geq g(x), \quad \forall x > M. \end{aligned} \tag{3.16}$$

Portanto, segue de (3.15) e (3.16) que, existe  $A = \max\{N, M\}$  tal que

$$g_c^{\sup}(x) \leq g(x) \leq g_c^{\inf}(x), \quad \forall x > A.$$

□

**Proposição 3.5.** *Seja  $X$  uma variável aleatória contínua com f.d.a de base  $G(x)$  e f.d.p  $g(x)$ . Para  $0 < G(x) < 1$ , existe  $A < 0$  tal que  $g_c^{\inf}(x) \leq g(x) \leq g_c^{\sup}(x)$ , sempre que  $x < A$ .*

*Demonstração.* Inicialmente, lembre que

$$g_c^{\inf}(x) = g(x)[1 - (1 - G(x))^c + cG(x)(1 - G(x))^{c-1}].$$

Pela condição  $F3$  dada na proposição 2.1, temos que  $\lim_{x \rightarrow -\infty} G(x) = 0$ . Daí, vem que

$$\lim_{x \rightarrow -\infty} [1 - (1 - G(x))^c] = 0 \tag{3.17}$$

e

$$\lim_{x \rightarrow -\infty} [G(x)(1 - G(x))^{c-1}] = 0. \tag{3.18}$$

Note que, o limite em (3.18) é igual a zero, pois  $G(x) \rightarrow 0$  e  $(1 - G(x))^{c-1} \rightarrow 1$  quando  $x \rightarrow -\infty$ . Assim, pela definição de limite no infinito, segue de (3.17) e (3.18), respectivamente, que dado  $0 < \varepsilon < 1$

$$\exists N_1 < 0; \ x < N_1 \quad \Rightarrow \quad |1 - (1 - G(x))^c| < \frac{\varepsilon}{2}; \quad (3.19)$$

$$\exists N_2 < 0; \ x < N_2 \quad \Rightarrow \quad |G(x)(1 - G(x))^{c-1}| < \frac{\varepsilon}{2c}. \quad (3.20)$$

Como  $1 - (1 - G(x))^c > 0$  e  $G(x)(1 - G(x))^{c-1} > 0$ , fazendo  $N = \min\{N_1, N_2\}$  segue de (3.19) e (3.20), que dado  $0 < \varepsilon < 1$

$$\begin{aligned} \exists N < 0; \ x < N \quad &\Rightarrow \quad 1 - (1 - G(x))^c + cG(x)(1 - G(x))^{c-1} < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon \\ &\Rightarrow \quad g(x)[1 - (1 - G(x))^c + cG(x)(1 - G(x))^{c-1}] < g(x)\varepsilon \leq g(x) \\ &\Rightarrow \quad g_c^{\inf}(x) \leq g(x), \quad \forall x < N. \end{aligned} \quad (3.21)$$

Por outro lado, lembre que

$$g_c^{\sup}(x) = g(x)[1 - G^c(x) + c(G^{c-1}(x) - G^c(x))].$$

Agora, observe que

$$\lim_{x \rightarrow -\infty} \frac{1 - G(x)}{G(x)} = +\infty,$$

ou seja, dado  $M_1 > 0$ , existe  $M < 0$  tal que

$$\left| \frac{1 - G(x)}{G(x)} \right| \geq M_1, \quad \forall x < M.$$

Seja  $M_1 \geq 1/c$ . Logo, existe  $M < 0$  tal que, para todo  $x < M$ , vale

$$\begin{aligned} \frac{1 - G(x)}{G(x)} \geq M_1 \geq \frac{1}{c} \quad &\Leftrightarrow \quad \frac{1}{G(x)} - 1 \geq \frac{1}{c} \\ &\Leftrightarrow \quad G^{c-1}(x) - G^c(x) \geq \frac{G^c(x)}{c} \\ &\Leftrightarrow \quad c(G^{c-1}(x) - G^c(x)) \geq G^c(x) \\ &\Leftrightarrow \quad 1 - G^c(x) + c(G^{c-1}(x) - G^c(x)) \geq 1 \\ &\Leftrightarrow \quad g(x)[1 - G^c(x) + c(G^{c-1}(x) - G^c(x))] \geq g(x) \\ &\Leftrightarrow \quad g_c^{\sup}(x) \geq g(x), \quad \forall x < M. \end{aligned} \quad (3.22)$$

Portanto, segue de (3.21) e (3.22) que, existe  $A = \min\{N, M\}$  tal que

$$g_c^{\inf}(x) \leq g(x) \leq g_c^{\sup}(x), \quad \forall x < A.$$

□

O seguinte resultado mostra que as caudas das distribuições  $G(x)$  e  $G_c^{\text{inf}}(x)$  são assintoticamente da mesma ordem.

**Proposição 3.6.** *Seja  $X$  uma variável aleatória contínua com f.d.a  $G(x)$ . Se  $X^{\text{inf}}$  é uma variável aleatória com f.d.a  $G_c^{\text{inf}}(x)$ , com  $c > 1$ , em que  $G(x)$  é a distribuição de base, então*

$$\lim_{x \rightarrow +\infty} \frac{P(X > x)}{P(X^{\text{inf}} > x)} = 1.$$

*Demonstração.* De fato,

$$\begin{aligned} \lim_{x \rightarrow +\infty} \frac{P(X^{\text{inf}} > x)}{P(X > x)} &= \lim_{x \rightarrow +\infty} \frac{1 - P(X^{\text{inf}} \leq x)}{1 - P(X \leq x)} = \lim_{x \rightarrow +\infty} \frac{1 - G_c^{\text{inf}}(x)}{1 - G(x)} \\ &= \lim_{x \rightarrow +\infty} \frac{1 - G(x)[1 - (1 - G(x))^c]}{1 - G(x)} \\ &= \lim_{x \rightarrow +\infty} \frac{1 - G(x) + G(x)(1 - G(x))^c}{1 - G(x)} \\ &= 1 + \lim_{x \rightarrow +\infty} \frac{G(x)(1 - G(x))^c}{1 - G(x)} \\ &= 1 + \lim_{x \rightarrow +\infty} G(x)(1 - G(x))^{c-1} \\ &= 1 + 0 = 1. \end{aligned}$$

Portanto, segue que

$$\lim_{x \rightarrow +\infty} \frac{P(X^{\text{inf}} > x)}{P(X > x)} = 1.$$

□

### 3.2.3 Interpretação Física dos Modelos sup e inf

Dado uma f.d.a  $G(x)$ , considere um sistema em série composto por dois componentes independentes. O sistema falha quando algum dos componentes falhar. Suponha que  $Y$  e  $Z$  denotam sua vida útil. Digamos que  $Y \sim G$ .

Por exemplo, se  $n \in \mathbb{N}$  e o componente com vida útil  $Z$  for composto por  $n$  subcomponentes paralelos independentes, de modo que o componente falhe apenas se todos os subcomponentes falharem, onde cada subcomponente tenha vida útil  $Z_j \sim G$ ,  $j = 1, \dots, n$ . Neste caso,  $Z = \max(Z_1, \dots, Z_n)$  (Ver Figura 3.1). Seja  $X = \min(Y, Z)$ . Então,  $X$  representa a vida útil de todo o sistema.

Observe que, pela Proposição 2.5

$$F_Z(x) = P(Z \leq x) = G^n(x).$$



Agora, note que

$$\begin{aligned} P(X > x) &= P(Y > x, Z > x) = P(Y > x)P(Z > x) \\ &= [1 - P(Y \leq x)][1 - P(Z \leq x)] = [1 - G(x)][1 - G^n(x)] \\ &= 1 - G^n(x) - G(x) + G^{n+1}(x). \end{aligned}$$

Assim,

$$F_X(x) = P(X \leq x) = 1 - P(X > x) = G(x) + G^n(x) - G^{n+1}(x),$$

portanto  $X \sim G_n^{\text{sup}}$ .

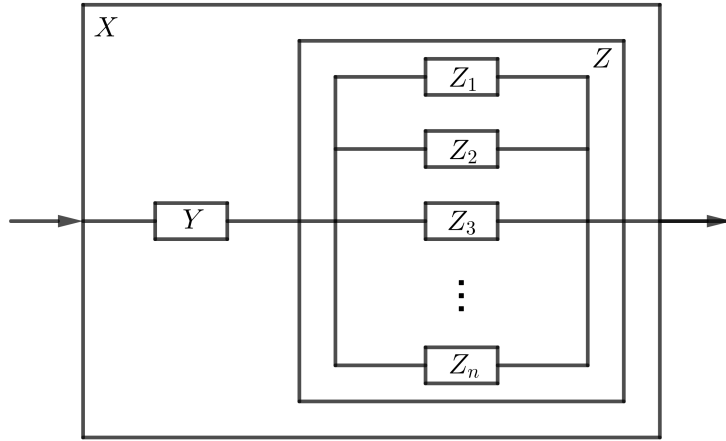


Figura 3.1: Sistema de componentes em série para interpretação física do modelo sup.

Por outro lado, considere um sistema paralelo composto por dois componentes independentes. O sistema falha quando todos os componentes falharem. Suponha que  $Y$  e  $Z$  denotam sua vida útil. Digamos que  $Y \sim G$ .

Por exemplo, se  $n \in \mathbb{N}$  e o componente com vida útil  $Z$  for composto por  $n$  subcomponentes independentes em série, de modo que o componente falhe se um de seus subcomponentes falharem, onde cada subcomponente tenha vida útil  $Z_j \sim G$ ,  $j = 1, \dots, n$ . Neste caso,  $Z = \min(Z_1, \dots, Z_n)$ . Seja  $X = \min(Y, Z)$ . Então,  $X$  representa a vida útil de todo o sistema.

Perceba novamente que, pela Proposição 2.5

$$F_Z(x) = P(Z \leq x) = 1 - [1 - G(x)]^n.$$

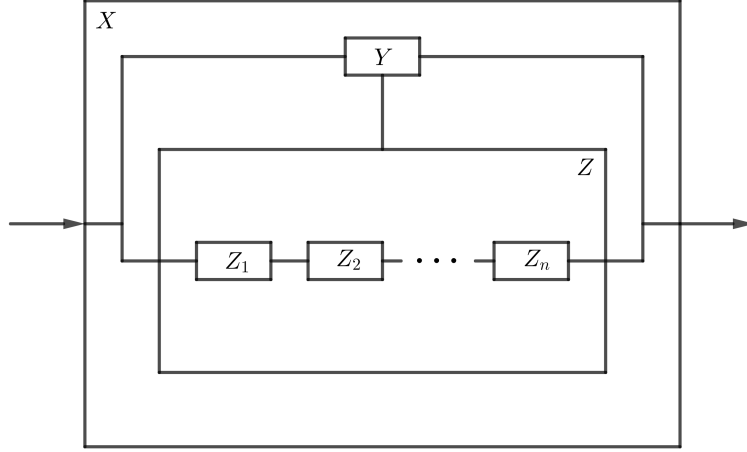


Figura 3.2: Sistema de componentes em paralelo para interpretação física do modelo inf.

Logo,

$$\begin{aligned} P(X \leq x) &= P(Y \leq x, Z \leq x) = P(Y \leq x)P(Z \leq x) \\ &= G(x)[1 - [1 - G(x)]^n]. \end{aligned}$$

Portanto,  $X \sim G_n^{\text{inf}}$ . Desta forma, podemos observar que tais famílias de distribuições podem ser obtidas como máximos e mínimos de sequências de variáveis aleatórias independentes.

### 3.3 Gerando números pseudo-aleatórios das distribuições $G_c^{\text{sup}}$ e $G_c^{\text{inf}}$

Nesta seção, iremos mostrar de acordo com [Tablada \(2017\)](#), como obter uma observação das distribuições  $G_c^{\text{sup}}$  e  $G_c^{\text{inf}}$ , para posteriormente chegarmos ao algoritmo para gerar números pseudo-aleatórios que provenham destas distribuições.

#### 3.3.1 Algoritmo para gerar números pseudo-aleatórios da família $G_c^{\text{sup}}$

Vimos na Subseção 3.2.3, que se  $Y \sim G$  e  $Z_{(n)} \sim G^n$ , onde  $Z_{(n)}$  é o máximo de uma amostra aleatória *i.i.d.* que segue distribuição  $G$ , fazendo  $X = \min\{Y, Z_{(n)}\}$ , obtemos que  $X \sim G_n^{\text{sup}}$ .

Sejam  $u$  e  $v$  observações independentes com distribuição uniforme  $\mathcal{U}(0, 1)$  e seja  $Q = F^{-1}$  a função quantílica da distribuição  $G$ . Assim, pela Proposição 2.3 temos que

Se  $u = G(y)$ , então  $y = Q(u)$  segue distribuição  $G$ .

Se  $v = G^c(z)$ , então  $z = Q(v^{\frac{1}{c}})$  segue distribuição  $G^c$ .

Logo, se fizermos  $x = \min\{y, z\}$ , teremos uma observação  $x$  que segue distribuição  $G_c^{\sup}$ .

Portanto segue o algoritmo para gerar uma observação da distribuição  $G_c^{\sup}$ :

1. Gerar observações  $u$  e  $v$  independentes com distribuição  $\mathcal{U}(0, 1)$ .
2. Calcular  $y = Q(u)$  da distribuição  $G$ .
3. Calcular  $z = Q(v^{\frac{1}{c}})$  da distribuição  $G^c$ .
4. Calcular  $x = \min\{y, z\}$ , a qual será uma observação da distribuição  $G_c^{\sup}$ .

### 3.3.2 Algoritmo para gerar números pseudo-aleatórios da família $G_c^{\inf}$

Temos pela Subseção 3.2.3, que se  $Y \sim G$  e  $Z_{(1)} \sim 1 - [1 - G]^n$ , em que  $Z_{(1)}$  é o mínimo de uma amostra aleatória *i.i.d.* que segue distribuição  $G$ , fazendo  $X = \max\{Y, Z_{(1)}\}$ , obtemos que  $X \sim G_n^{\inf}$ .

Novamente, sejam  $u$  e  $v$  observações independentes com distribuição uniforme  $\mathcal{U}(0, 1)$  e seja  $Q = F^{-1}$  a função quantílica da distribuição  $G$ . Assim, pela Proposição 2.3 temos que,

Se  $u = G(y)$ , então  $y = Q(u)$  segue distribuição  $G$ .

$$\begin{aligned} \text{Se } v = 1 - [1 - G(z)]^c &\Rightarrow [1 - G(z)]^c = 1 - v \\ &\Rightarrow G(z) = 1 - (1 - v)^{\frac{1}{c}} \\ &\Rightarrow z = G^{-1}(1 - (1 - v)^{\frac{1}{c}}) = Q(1 - (1 - v)^{\frac{1}{c}}). \end{aligned}$$

em que  $z$  segue distribuição  $1 - [1 - G]^c$ .

Logo, se fizermos  $x = \max\{y, z\}$ , teremos uma observação  $x$  que segue distribuição  $G_c^{\inf}$ .

Portanto, segue o algoritmo para gerar uma observação da distribuição  $G_c^{\inf}$ :

1. Gerar observações  $u$  e  $v$  independentes com distribuição  $\mathcal{U}(0, 1)$ .
2. Calcular  $y = Q(u)$  da distribuição  $G$ .
3. Calcular  $z = Q(1 - (1 - v)^{\frac{1}{c}})$  da distribuição  $1 - [1 - G]^c$ .
4. Calcular  $x = \max\{y, z\}$ , a qual será uma observação da distribuição  $G_c^{\inf}$ .

### 3.4 Estimadores de máxima verossimilhança das famílias $G_c^{\sup}$ e $G_c^{\inf}$

Nesta seção, faremos um estudo inferencial nas famílias  $G_c^{\sup}$  e  $G_c^{\inf}$ .

#### 3.4.1 Função de log-verossimilhança da família $G_c^{\sup}$

Seja  $X_1, X_2, \dots, X_n$  uma amostra aleatória *i.i.d.* da variável aleatória  $X$  que segue distribuição  $G_c^{\sup}(x, \theta)$ , em que  $\theta = (c, \delta^\top)^\top$  é o vetor de parâmetros da distribuição  $G_c^{\sup}$  e  $\delta$  é o vetor de parâmetros da distribuição de base  $G$ . Denotaremos por  $l_{g_c^{\sup}}(\theta; \mathbf{x})$  e  $l_g(\delta; \mathbf{x})$  a função de log-verossimilhança de  $G_c^{\sup}$  e  $G$ , respectivamente, sendo  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$  observação das variáveis aleatórias.

Temos que, a função de verossimilhança de  $G_c^{\sup}$  é dada por

$$\begin{aligned} L^{\sup}(\theta; \mathbf{x}) &= \prod_{i=1}^n g_c^{\sup}(x_i; \theta) \\ &= \prod_{i=1}^n g(x_i; \delta) [1 + cG^{c-1}(x_i; \delta) - (c+1)G^c(x_i; \delta)] \\ &= \prod_{i=1}^n g(x_i; \delta) h_1(x_i; \theta) \end{aligned}$$

em que  $h_1(x; \theta) = [1 + cG^{c-1}(x; \delta) - (c+1)G^c(x; \delta)]$ . Daí, segue que a função de log-verossimilhança de  $G_c^{\sup}$  é dada por

$$\begin{aligned} l_{g_c^{\sup}}(\theta; \mathbf{x}) &= \log[L^{\sup}(\theta; \mathbf{x})] = \sum_{i=1}^n \log [g(x_i; \delta) h_1(x_i; \theta)] \\ &= \sum_{i=1}^n \log (g(x_i; \delta)) + \sum_{i=1}^n \log (h_1(x_i; \theta)) \\ &= l_g(\delta; \mathbf{x}) + \sum_{i=1}^n \log (h_1(x_i; \theta)). \end{aligned}$$

Provaremos agora que para  $0 \leq G(x; \delta) < 1$ , que

$$\lim_{c \rightarrow \infty} \sum_{i=1}^n \log (h_1(x_i; \theta)) = 0. \quad (3.23)$$

Note que, provar a igualdade (3.23) equivale a provar que  $\lim_{c \rightarrow \infty} h_1(x; \theta) = 1$ . De fato, se  $G(x; \delta) = 0$ , então é claro que  $h_1(x; \theta) = 1$ . Para  $0 < \overset{c \rightarrow \infty}{G}(x; \delta) < 1$ , temos que

$$\begin{aligned}\lim_{c \rightarrow \infty} h_1(x; \theta) &= \lim_{c \rightarrow \infty} [1 + c G^{c-1}(x; \delta) - (c+1) G^c(x; \delta)] \\ &= 1 + \lim_{c \rightarrow \infty} \frac{c}{[G(x; \delta)^{-1}]^{c-1}} - \lim_{c \rightarrow \infty} \frac{c+1}{[G(x; \delta)^{-1}]^c}\end{aligned}$$

Visto que, os dois limites acima, na última igualdade, são do tipo  $\frac{+\infty}{+\infty}$ , utilizaremos a regra de L'Hôpital para calculá-los. Assim, derivaremos o numerador e denominador de ambos os limites, em relação a  $c$ .

Com efeito,

$$\begin{aligned}\lim_{c \rightarrow \infty} h_1(x; \theta) &= 1 + \lim_{c \rightarrow \infty} \frac{1}{[G(x; \delta)^{-1}]^{c-1} \log(G(x; \delta)^{-1})} - \lim_{c \rightarrow \infty} \frac{1}{[G(x; \delta)^{-1}]^c \log(G(x; \delta)^{-1})} \\ &= 1 + 0 - 0 = 1.\end{aligned}$$

pois  $G(x; \delta)^{-1} > 1$  visto que  $0 < G(x; \delta) < 1$ .

Isto prova (3.23).

Portanto, segue que para todo  $x_1, x_2, \dots, x_n$  no suporte de  $G$

$$\lim_{c \rightarrow \infty} l_{g_c^{\sup}}(\theta; \mathbf{x}) = l_g(\delta; \mathbf{x}).$$

### 3.4.2 Função de log-verossimilhança da família $G_c^{\inf}$

Seja  $X_1, X_2, \dots, X_n$  uma amostra aleatória *i.i.d.* da variável aleatória  $X$  que segue distribuição  $G_c^{\inf}(x, \theta)$ , em que  $\theta = (c, \delta^\top)^\top$  é o vetor de parâmetros da distribuição  $G_c^{\inf}$  e  $\delta$  é o vetor de parâmetros da distribuição de base  $G$ . Denotaremos por  $l_{g_c^{\inf}}(\theta; \mathbf{x})$  e  $l_g(\delta; \mathbf{x})$  a função de log-verossimilhança de  $G_c^{\inf}$  e  $G$ , respectivamente, sendo  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$  uma observação das variáveis aleatórias  $X_1, X_2, \dots, X_n$ .

Temos que a função de verossimilhança de  $G_c^{\inf}$  é dada por

$$\begin{aligned}L^{\inf}(\theta; \mathbf{x}) &= \prod_{i=1}^n g_c^{\inf}(x_i; \theta) \\ &= \prod_{i=1}^n g(x_i; \delta) [1 - (1 - G(x_i; \delta))^c + c G(x_i; \delta) (1 - G(x_i; \delta))^{c-1}] \\ &= \prod_{i=1}^n g(x_i; \delta) h_2(x_i; \theta)\end{aligned}$$

em que  $h_2(x; \theta) = [1 - (1 - G(x; \delta))^c + c G(x; \delta) (1 - G(x; \delta))^{c-1}]$ . Daí, segue que

a função de log-verossimilhança de  $G_c^{\text{inf}}$  é dada por

$$\begin{aligned} l_{g_c^{\text{inf}}}(\theta; \mathbf{x}) &= \log[L^{\text{inf}}(\theta; \mathbf{x})] = \sum_{i=1}^n \log [g(x_i; \delta)h_2(x_i; \theta)] \\ &= \sum_{i=1}^n \log (g(x_i; \delta)) + \sum_{i=1}^n \log (h_2(x_i; \theta)) \\ &= l_g(\delta; \mathbf{x}) + \sum_{i=1}^n \log (h_2(x_i; \theta)). \end{aligned}$$

Provaremos agora, para  $0 < G(x; \delta) \leq 1$ , que

$$\lim_{c \rightarrow \infty} \sum_{i=1}^n \log (h_2(x_i; \theta)) = 0. \quad (3.24)$$

Perceba que, provar a igualdade (3.24) é equivalente a provar que

$$\lim_{c \rightarrow \infty} h_2(x; \theta) = 1 \quad (3.25)$$

De fato, se  $G(x; \delta) = 1$ , então  $h_2(x; \theta) = 1$ , satisfazendo assim, o limite (3.25).

Se  $0 < G(x; \delta) < 1$ , então

$$\begin{aligned} \lim_{c \rightarrow \infty} h_2(x; \theta) &= \lim_{c \rightarrow \infty} [1 - (1 - G(x; \delta))^c + cG(x; \delta)(1 - G(x; \delta))^{c-1}] \\ &= 1 - \lim_{c \rightarrow \infty} (1 - G(x; \delta))^c + G(x; \delta) \lim_{c \rightarrow \infty} \frac{c}{[(1 - G(x; \delta))^{-1}]^{c-1}} \end{aligned}$$

Temos que  $\lim_{c \rightarrow \infty} (1 - G(x; \delta))^c = 0$ , pois  $0 < (1 - G(x; \delta)) < 1$ . Utilizando novamente a regra de L'Hôpital, temos

$$\lim_{c \rightarrow \infty} \frac{c}{[(1 - G(x; \delta))^{-1}]^{c-1}} = \lim_{c \rightarrow \infty} \frac{1}{[(1 - G(x; \delta))^{-1}]^{c-1} \log ((1 - G(x; \delta))^{-1})} = 0$$

pois  $[(1 - G(x; \delta))^{-1}]^{c-1}$  vai para  $+\infty$  quando  $c \rightarrow +\infty$ .

Logo,

$$\lim_{c \rightarrow \infty} h_2(x; \theta) = 1.$$

Isto prova (3.24).

Portanto, segue que para todo  $x_1, x_2, \dots, x_n$  no suporte de  $G$

$$\lim_{c \rightarrow \infty} l_{g_c^{\text{inf}}}(\theta; \mathbf{x}) = l_g(\delta; \mathbf{x}).$$

## Capítulo 4

# Uma nova abordagem para o aperfeiçoamento do teste da razão de verossimilhanças

Como foi comentado na seção 2.2, o TRV pode ser usado como um procedimento para testar hipóteses compostas do tipo

$$H_0 : \theta \in \Theta_0 \quad \text{versus} \quad H_1 : \theta \in \Theta_1 \quad (4.1)$$

em que  $\{\Theta_0, \Theta_1\}$  é uma partição do espaço paramétrico  $\Theta$ . Tal procedimento consiste em definir a região crítica

$$\mathbf{A}_1^* = \left\{ \mathbf{x}; \lambda^*(\mathbf{x}) = \frac{\sup_{\theta \in \Theta_0} L(\theta; \mathbf{x})}{\sup_{\theta \in \Theta} L(\theta; \mathbf{x})} \leq c \right\},$$

sendo  $\lambda^*(\mathbf{x})$  a estatística da razão de verossimilhanças com base na amostra observada  $\mathbf{x} = (x_1, x_2, \dots, x_n)$ .

Em inferência estatística, para testar uma hipótese usando o TRV, geralmente é usada a estatística  $RV = -2 \log(\lambda^*(\mathbf{x}))$ . Cordeiro e Cribari-Neto (2014) afirmam que a principal dificuldade de testar uma hipótese nula usando a estatística  $RV$  não consiste em obter sua expressão de forma fechada, quando ela tem uma, mas em encontrar sua distribuição nula exata, ou pelo menos uma boa aproximação desta. Nesse sentido, Wilks (1938) demonstrou que, sob hipótese nula, a estatística  $RV$  tem distribuição assintótica qui-quadrado com  $k$  graus de liberdade ( $\chi_k^2$ ), sendo  $k$  o número de parâmetros de interesse a ser testados na hipótese nula (Teorema 2.2). Esse importante resultado viabiliza o uso do quantil  $1 - \gamma$  da distribuição  $\chi_k^2$  para definir a região crítica do TRV quando o tamanho  $n$  da amostra considerada é grande.

Por se tratar de um resultado assintótico, o Teorema 2.2 não garante que a

distribuição  $\chi_k^2$  seja uma boa aproximação à distribuição nula da estatística  $RV$  quando a amostra é pequena. De fato, como foi comentado no final da Seção 2.2, se  $n$  é pequeno, o TRV pode levar a inferências imprecisas quando consideramos a distribuição  $\chi_k^2$ . Sendo assim, para reverter esse quadro, na literatura são propostas diferentes correções na estatística  $RV$ , sendo [Bartlett \(1937\)](#) o pioneiro na correção dessa estatística.

Já [Cordeiro e Ferrari \(1991\)](#), derivaram uma fórmula geral para correções do tipo Bartlett para melhorar qualquer estatística de teste que, sob a hipótese nula, seja assintoticamente distribuída como a qui-quadrado. A correção padrão de Bartlett é um caso especial do resultado geral. Ainda, segundo [Cordeiro e Ferrari \(1991\)](#), as correções do tipo Bartlett e Bartlett pretendem aproximar os tamanhos empíricos dos testes assintóticos aos tamanhos nominais correspondentes. Na maioria dos casos, elas fazem isso com bastante eficiência. Vale ressaltar que essas correções podem levar a uma perda de poder do teste.

A implementação das correções Bartlett e tipo Bartlett dependem de certos cumulantes – quantidades que envolvem o cálculo do valor esperado de derivadas da função de verossimilhança. Geralmente, estas quantidades são de difícil obtenção e, dependendo da expressão da função de verossimilhança, os cálculos desses cumulantes tornam-se inviáveis.

## 4.1 A nova abordagem

Nesta pesquisa propomos uma correção na distribuição a qual a estatística  $RV$  é distribuída quando  $n \rightarrow \infty$ , isto é, a distribuição qui-quadrado, ao contrário de corrigir a estatística  $RV$ , como acontece com os métodos que foram mencionados anteriormente. O objetivo desse procedimento é obter uma correção na cauda da distribuição qui-quadrado de tal forma que a nova distribuição forneça um aperfeiçoamento da região crítica do TRV, conduzindo a resultados inferenciais mais acurados quando  $n$  é pequeno ou moderado.

[Wilks \(1938\)](#) mostrou que, sob a hipótese nula, a estatística  $\lambda^*(\mathbf{x})$  pode ser distribuída como

$$\lambda^*(\mathbf{x}) \sim \exp\left(-\frac{1}{2}\chi_k^2\right)(1 + \phi)$$

sendo  $\chi_k^2$  uma variável aleatória com distribuição qui-quadrado com  $k$  graus de liberdade e  $\phi = \mathcal{O}(1/\sqrt{n})$  é um termo de ordem  $1/\sqrt{n}$ , ou seja,  $\lim_{n \rightarrow \infty} \sqrt{n}\phi$  é uma constante. Portanto, com exceção de termos de ordem  $1/\sqrt{n}$ , temos que

$$RV = -2\log(\lambda^*(\mathbf{x})) \sim \chi_k^2.$$



Podemos observar que, quando o tamanho da amostra  $n$  é pequeno, o termo  $\phi$  não é desprezível. Assim, para  $n$  pequeno, temos que o valor esperado da estatística  $RV$  excede o valor esperado da distribuição  $\chi_k^2$ , ou seja

$$\mathbb{E}_{H_0}[-2 \log(\lambda^*(\mathbf{x}))] > \mathbb{E}(\chi_k^2) = k, \quad n \text{ pequeno},$$

em que a esperança do lado esquerdo é calculada supondo  $H_0$  verdadeira. Dessa forma, para  $n$  pequeno, os valores observados da estatística  $RV$  devem ser maiores do que  $k$ , em média. Como consequência, se  $q$  e  $q'$  são os quantis  $1 - \gamma$  da distribuição  $\chi_k^2$  e da distribuição nula de  $RV$  respectivamente, os quais satisfazem

$$P_{H_0}(RV > q') = \gamma = P(\chi_k^2 > q),$$

então esperamos que  $q' > q$ .

Dessa forma percebemos que usar o quantil  $q$  como uma aproximação de  $q'$  (desconhecido) para definir a região crítica do TRV, quando  $n$  é pequeno, deve ocasionar um aumento na taxa de rejeição em relação ao nível nominal  $\gamma$  especificado (o teste rejeita “mais do que deveria”).

Com o objetivo de aperfeiçoar o TRV para pequenas amostras, e após o discutido anteriormente, parece razoável propor uma correção na cauda da distribuição  $\chi_k^2$ , de tal forma que o quantil  $1 - \gamma$  da distribuição “corrigida” esteja mais próximo do quantil nulo exato  $q'$ . Como mostraremos a seguir, essa correção é dada naturalmente pela distribuição qui-quadrado inf.

#### 4.1.1 A distribuição qui-quadrado inf

Seja  $X$  uma variável aleatória que segue distribuição qui-quadrado com  $k$  graus de liberdade, e que denotamos por  $X \sim \chi_k^2$ . Isto é,  $X$  tem f.d.p dada por

$$g(x; k) = \frac{1}{2^{\frac{k}{2}} \Gamma(\frac{k}{2})} x^{\frac{k}{2}-1} e^{-\frac{x}{2}}, \quad x > 0.$$

A f.d.a pode expressada como

$$G(x; k) = \int_{-\infty}^x g(t; k) dt = \int_0^x \frac{1}{2^{\frac{k}{2}} \Gamma(\frac{k}{2})} t^{\frac{k}{2}-1} e^{-\frac{t}{2}} dt = \frac{\gamma(\frac{k}{2}, \frac{x}{2})}{\Gamma(\frac{k}{2})}, \quad x > 0,$$

sendo  $\gamma(\alpha, t) = \int_0^t s^{\alpha-1} e^{-s} ds$  a função gama incompleta inferior e  $\Gamma(\alpha) = \int_0^\infty s^{\alpha-1} e^{-s} ds$  a função gama.

Seja  $X^{\text{inf}}$  uma variável aleatória que segue distribuição qui-quadrado inf de parâmetros  $k$  e  $c > 0$ . Por (3.5) e (3.6), f.d.a e f.d.p são dadas, respectivamente,

por

$$G_c^{\text{inf}}(x; k, c) = \frac{\gamma(\frac{k}{2}, \frac{x}{2})}{\Gamma(\frac{k}{2})} \left[ 1 - \left( 1 - \frac{\gamma(\frac{k}{2}, \frac{x}{2})}{\Gamma(\frac{k}{2})} \right)^c \right]$$

e

$$g_c^{\text{inf}}(x; k, c) = \frac{1}{2^{\frac{k}{2}} \Gamma(\frac{k}{2})} x^{\frac{k}{2}-1} e^{-\frac{x}{2}} \left[ 1 - \left( 1 - \frac{\gamma(\frac{k}{2}, \frac{x}{2})}{\Gamma(\frac{k}{2})} \right)^c + c \frac{\gamma(\frac{k}{2}, \frac{x}{2})}{\Gamma(\frac{k}{2})} \left( 1 - \frac{\gamma(\frac{k}{2}, \frac{x}{2})}{\Gamma(\frac{k}{2})} \right)^{c-1} \right]. \quad (4.2)$$

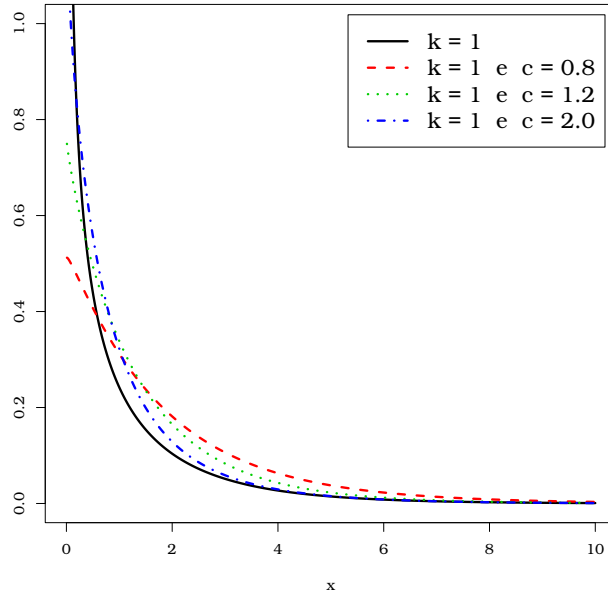
Na Figura 4.1 é ilustrado o comportamento da distribuição qui-quadrado inf em relação à distribuição qui-quadrado. Na primeira imagem fixamos o parâmetro  $k = 1$ , na segunda fixamos  $k = 2$  e variamos o parâmetro  $c$  em ambas.

Analizando a Figura 4.1, podemos verificar que os resultados discutidos no Capítulo 3 são satisfeitos. De fato, vemos que para  $k$  fixo, quando o parâmetro  $c$  aumenta, a distribuição qui-quadrado inf converge para a distribuição qui-quadrado (Proposições 3.2 e 3.3). Ainda, observe que para cada distribuição qui-quadrado inf, existe um valor real  $A > 0$  tal que a cauda dessa distribuição fica por cima da distribuição qui-quadrado para  $x > A$  (proposição 3.4). Finalmente, pela proposição 3.6, pode-se observar que, para  $x$  grande, a cauda da distribuição qui-quadrado inf “coincide” com a cauda da distribuição qui-quadrado (as caudas são assintoticamente da mesma ordem).

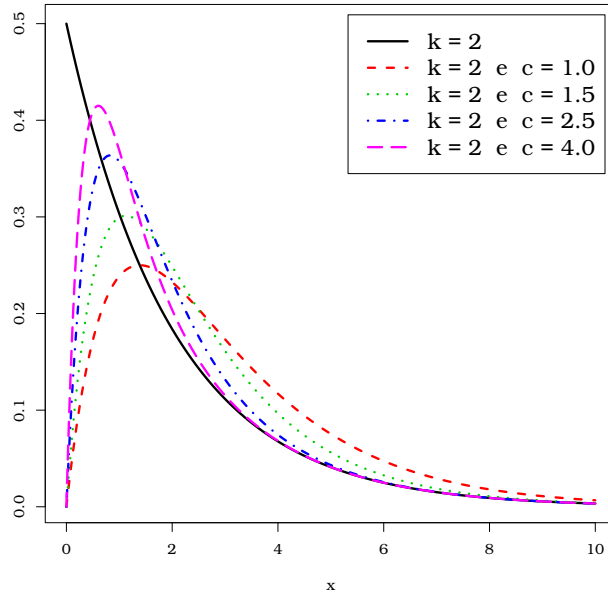
Denotemos por  $q$  e  $q_{\text{inf}}$  os quantis  $1 - \gamma$  das distribuições qui-quadrado e qui-quadrado inf, respectivamente. O fato do modelo qui-quadrado inf ser uma distribuição adequada para ser usada como correção da distribuição qui-quadrado no TRV, considerando pequenas amostras, decorre do seguinte:

- Pela Proposição 3.4, a distribuição qui-quadrado inf possui uma cauda mais leve (maior área na cauda) do que a distribuição qui-quadrado. Como consequência, temos que  $q_{\text{inf}} > q$ .
- Embora  $q_{\text{inf}} > q$ , a Proposição 3.6 garante que o quantil  $q_{\text{inf}}$  não seja excessivamente maior do que  $q$ , pois as caudas de ambas distribuições são assintoticamente da mesma ordem.
- Os comentários feitos no início da seção e as Proposições 3.2 e 3.3 sugerem que existe um valor do parâmetro  $c$  para o qual  $q_{\text{inf}} \approx q'$ , sendo  $q'$  o quantil exato  $1 - \gamma$  da distribuição nula da estatística  $RV$ .

Geralmente, o TRV baseia-se em uma única observação da estatística  $RV$ . Sendo assim, a estimação do parâmetro  $c$  da distribuição qui-quadrado inf para garantir que  $q_{\text{inf}} \approx q'$  é inviável pelos métodos de estimação usuais. No entanto, as propriedades



(a)



(b)

Figura 4.1: Plots das f.d.p da distribuição qui-quadrado (linha contínua) e qui-quadrado inf (linhas tracejadas)

discutidas no Capítulo 3 referentes às famílias  $G_c^{\text{inf}}$  sugerem que o parâmetro  $c$  deva ser uma função crescente do tamanho amostral  $n$ , pelos seguintes motivos:

- Pelo discutido no início da seção, esperamos que o valor  $|q' - q|$  aumente na medida em que o valor de  $n$  diminua. Analogamente, pela Proposição 3.2, o

valor  $|q_{\inf} - q|$  também aumenta na medida em que o valor de  $c$  diminui (se aproxima de zero).

- Dado que a distribuição nula da estatística  $RV$  é assintoticamente qui-quadrado, esperamos que  $q' \rightarrow q$  quando  $n \rightarrow \infty$ . Da mesma forma, a Proposição 3.3 garante que  $q_{\inf} \rightarrow q$  quando  $c \rightarrow \infty$ .

Por simplicidade, neste trabalho propomos a seguinte relação

$$c = \rho \log(n), \quad (4.3)$$

sendo  $\rho > 0$  uma constante. A função  $\log(\cdot)$  é introduzida para “atenuar” o valor de  $n$ .

Após um estudo computacional de diversos casos, foi comprovado que  $\rho = 0.5$  é um valor adequado para se considerar, fornecendo uma boa correção na cauda da distribuição qui-quadrado para  $n$  pequeno. Dessa forma, no que segue sempre consideraremos a relação (4.3) com  $\rho = 0.5$ .

Em resumo, a nova abordagem para o aperfeiçoamento do TRV é apresentada a seguir.

**Aperfeiçoamento do TRV:** suponha que, com base em  $n$  observações, deseje-se testar a hipótese composta (4.1), e seja  $k = \dim(\Theta) - \dim(\Theta_0)$  (número de parâmetros não especificados em  $\Theta$  menos o número de parâmetros não especificados em  $\Theta_0$ ). A região crítica do TRV aperfeiçoado, ao nível de significância  $\gamma$ , fica determinada por

$$\mathbf{C}_1^* = \{\mathbf{x}; RV = -2 \log(\lambda^*(\mathbf{x})) \geq q_{\inf}\},$$

sendo  $q_{\inf}$  o quantil  $1 - \gamma$  da distribuição qui-quadrado inf de parâmetros  $k$  e  $c = 0.5 \log(n)$   $\square$

Pode-se observar que o uso da relação (4.3) garante a consistência da metodologia proposta, no sentido que

$$q_{\inf} \rightarrow q, \quad \text{quando } n \rightarrow \infty,$$

ou seja, quando  $n \rightarrow \infty$ , a região crítica aperfeiçoada do TRV coincide com a região crítica usual do TRV, considerando o quantil  $q$  da distribuição qui-quadrado.

A partir de agora, chamaremos de  $\text{TRV}^*$  o teste da razão de verossimilhanças que utiliza o quantil da distribuição qui-quadrado inf para determinar a região crítica do teste, isto é,  $\mathbf{C}_1^*$ .

Dado que a função quantílica da distribuição qui-quadrado inf não possui forma fechada, o quantil  $q_{\inf}$  deve ser obtido numericamente (ver Capítulo 5) através da

equação

$$\gamma = \int_{q_{\inf}}^{\infty} g_c^{\inf}(x; k, c) dx$$

em que  $g_c^{\inf}(x; k, c)$  é dada por (4.2) e  $c = 0.5 \log(n)$ .

A seguir, mostraremos um exemplo envolvendo o TRV baseado em uma amostra de tamanho pequeno, em que o mesmo será replicado um número grande de vezes para observarmos o comportamento do histograma das observações da estatística  $RV$  juntamente com as densidades das distribuições qui-quadrado e qui-quadrado inf.

**Exemplo:** Seja  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$   $n$  observações da variável aleatória  $X \sim \text{Exp} - \text{Weibull}(\alpha, \beta, \lambda)$  (Weibull Exponencializada), com f.d.p dada por

$$f(x; \alpha, \beta, \lambda) = \frac{\beta \lambda}{\alpha} \left(\frac{x}{\alpha}\right)^{\beta-1} \left[1 - \exp\left\{-\left(\frac{x}{\alpha}\right)^\beta\right\}\right]^{\lambda-1} \exp\left\{-\left(\frac{x}{\alpha}\right)^\beta\right\} \quad (4.4)$$

em que  $\alpha$ ,  $\beta$  e  $\lambda$  são parâmetros positivos e  $x > 0$ . Considere que o interesse seja testar  $H_0 : \lambda = 1$  versus  $H_1 : H_0$  é falsa, ou seja, queremos testar a hipótese nula a favor do modelo Weibull contra a hipótese alternativa a favor do modelo Weibull Exponencializado. Fixemos os parâmetros de perturbação em  $\alpha = 0.6$  e  $\beta = 1.4$ . Consideremos como nível nominal do teste  $\gamma = 10\%$  e o tamanho da amostra  $n = 10$ .

Neste exemplo, temos que

$$\Theta = \{(\alpha, \beta, \lambda); \alpha > 0, \beta > 0, \lambda > 0\}, \quad \Theta_0 = \{(\alpha, \beta, \lambda); \alpha > 0, \beta > 0, \lambda = 1\},$$

portanto  $k = 3 - 2 = 1$ .

A estatística da razão de verossimilhanças é dada por

$$\lambda^*(\mathbf{x}) = \frac{L(\tilde{\theta}; \mathbf{x})}{L(\hat{\theta}; \mathbf{x})},$$

em que  $\tilde{\theta} = (\tilde{\alpha}, \tilde{\beta}, 1)$  e  $\hat{\theta} = (\hat{\alpha}, \hat{\beta}, \hat{\lambda})$  são os estimadores de máxima verossimilhança restrito e irrestrito, respectivamente.

A Figura 4.2 mostra o histograma correspondente a 20000 observações da estatística  $RV$  com base no exemplo acima. Perceba o melhor ajuste na cauda dada pela distribuição qui-quadrado inf.

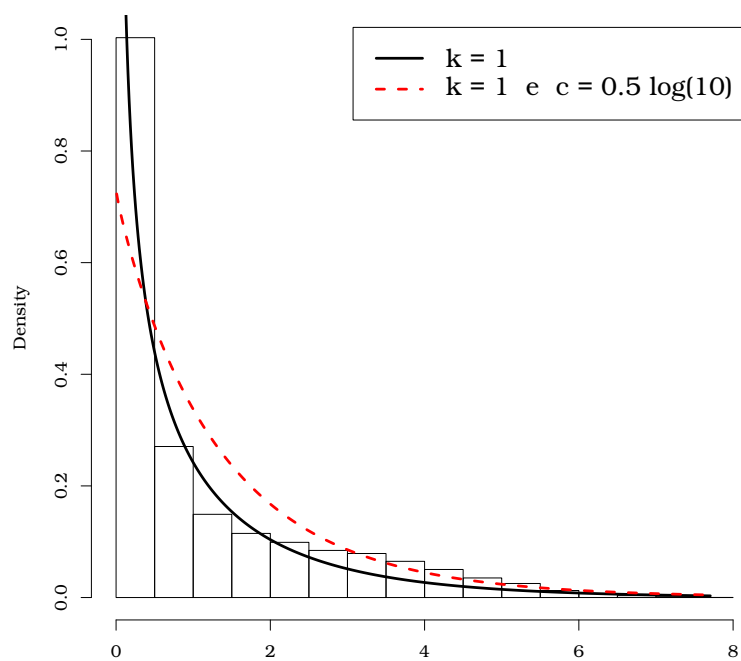


Figura 4.2: Histograma de observações da estatística  $RV$  (para  $n = 10$ ). Curvas das densidades da distribuição qui-quadrado (linha contínua) e qui-quadrado inf (linha tracejada).

## Capítulo 5

# O pacote LikRatioTest

As simulações deste trabalho foram realizadas pelo pacote **LikRatioTest**, versão 0.1, construído no decorrer desta pesquisa utilizando a linguagem R ([R CORE TEAM, 2019](#)). Sendo assim, destinamos esse capítulo para explicar ao leitor como manuseá-lo.

O pacote foi desenvolvido para obtenção dos resultados desta pesquisa, podendo este servir para que outros pesquisadores façam uso do mesmo impondo cenários diversos. Vale ressaltar, que o pacote ainda não encontra-se no CRAN (Comprehensive R Archive Network), mas pode ser facilmente instalado diretamente do repositório do GitHub, podendo ser consultado em <https://github.com/prdm0/LikRatioTest>. Instalar diretamente um pacote R que está sendo mantido no GitHub dará a possibilidade de instalar versões mais recentes do pacote que ainda não encontram-se no CRAN.

Vale salientar, que todos os resultados obtidos nesta pesquisa através das simulações poderão ser reproduzidos fazendo o devido uso deste pacote. Desta forma, garantimos a veracidade dos resultados. Para realizar a instalação, o usuário deverá ter instalado, previamente, o pacote **devtools** ([WICKHAM \*et al.\*, 2019](#)). A instalação do pacote **LikRatioTest** se dará fazendo

```
devtools::install_github(repo = "prdm0/LikRatioTest", force = TRUE)
```

Este pacote foca na realização do teste da razão de verossimilhanças. O mesmo é composto por três funções, a saber: `est_q()`, `lrt()` e `mc()`. A função `mc()` faz o uso das funções `lrt()` e `est_q()`. Vale salientar que a função `mc()` faz o uso de computação paralela por meio do pacote **parallel** ([R CORE TEAM, 2019](#)). Assim, o usuário irá usufruir de um código com maior performance, podendo o usuário controlar o número de cores disponível em sua máquina.

## 5.1 Função 'lrt'

A função `lrt()` é responsável por calcular a estatística da razão de verossimilhanças, isto é, calcula a estatística  $RV = -2\log(\lambda^*(\mathbf{x}))$  como dado no Teorema 2.2. A forma geral de uso da função é fornecida abaixo com a descrição de cada um de seus argumentos.

```
lrt(f, data, kicks, par0 = NULL, ...)
```

em que

- **f**: Função densidade de probabilidade considerada no teste da razão de verossimilhanças. Essa função deverá ser implementada conforme o exemplo;
- **data**: Um conjunto de dados que será considerado para a realização do teste;
- **kicks**: Chutes iniciais utilizados no processo de otimização;
- **par0**: Uma lista contendo como primeiro elemento um vetor com o nomes dos parâmetros que serão fixados sob hipótese nula e um vetor com os valores fixados para cada um dos respectivos parâmetros;
- **...**: Lista de argumentos adicionais que serão passados à função `optim()` utilizada no processo de otimização. Por exemplo, é possível escolher o método de otimização a ser utilizado.

Note que se a função for implementada segundo o exemplo, não haverá a necessidade de implementar a função sob a hipótese nula, visto que são as mesmas funções. Para especificar os parâmetros que serão fixados sob hipótese nula, i.e, para especificar a distribuição sob a hipótese nula, utiliza-se o argumento `par0`. A seguir, fornecemos um exemplo do uso da função `lrt()`.

**Exemplo:** Seja  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$   $n$  observações da variável aleatória  $X \sim Weibull(\alpha, \beta)$ , com f.d.p

$$f(x; \alpha, \beta) = \frac{\alpha}{\beta} \left(\frac{x}{\beta}\right)^{\alpha-1} \exp\left(-\left(\frac{x}{\beta}\right)^\alpha\right),$$

em que  $\alpha$  e  $\beta$  são parâmetros positivos e  $x > 0$ . Suponhamos que o interesse seja testar  $H_0 : \beta = 1$  versus  $H_1 : \beta \neq 1$ . Para fazermos o uso da função `lrt()` fixemos o parâmetro de perturbação  $\alpha = 1$ .

Neste exemplo, temos que o espaço paramétrico e o espaço restrito são dados, respectivamente por

$$\Theta = \{(\alpha, \beta); \alpha > 0, \beta > 0\}, \quad \Theta_0 = \{(\alpha, \beta); \alpha > 0, \beta = 1\}.$$



Sabemos que a estatística da razão de verossimilhanças é dada por

$$\lambda^*(\mathbf{x}) = \frac{\sup_{\theta \in \Theta_0} L(\theta; \mathbf{x})}{\sup_{\theta \in \Theta} L(\theta; \mathbf{x})}.$$

No entanto, computacionalmente o calculo de  $\lambda^*(\mathbf{x})$ , para este exemplo, é dado por

$$\lambda^*(\mathbf{x}) = \frac{L(\tilde{\theta}; \mathbf{x})}{L(\hat{\theta}; \mathbf{x})},$$

em que  $\tilde{\theta} = (\tilde{\alpha}, 1)$  e  $\hat{\theta} = (\hat{\alpha}, \hat{\beta})$  são os EMV restrito e irrestrito, respectivamente. Assim, teremos como retorno da função `lrt()` o valor da estatística do teste *RV*.

Inicialmente implementamos a função de densidade  $f$  que será atribuída a função `lrt()`. Em seguida, é implementado um gerador de números pseudo-aleatórios da distribuição  $Weibull(\alpha, \beta)$  que será usado para gerar um conjunto de dados para a realização do teste. Por fim, foi feito o uso da função, em que obtivemos como retorno o valor da estatística do teste, 0.9705, com base no conjunto de dados passado como argumento à `data`.

```
R> # Funcao densidade de probabilidade

R> pdf_w <- function(par, x, var = NULL){
+   alpha <- par[1]
+   beta <- par[2]
+   if (is.list(var)) eval(parse(text = paste(var[[1]], " <- ",
+   unlist(var[[2]]), sep = "")))
+   dweibull(x, shape = alpha, scale = beta)
+ }

R> # Gerador de números pseudo-aleatórios da distribuicao Weibull

R> rw <- function(n = 1L, alpha, beta){
+   rweibull(n = n, shape = alpha, scale = beta)
+ }

R> # Conjunto de dados do teste

R> set.seed(0) # Fixando uma semente no gerador de números
R> data <- rw(n = 100L, alpha = 1, beta = 1)

R> # Fazendo o uso da função lrt()
```

```
R> lrt(f = pdf_w, data = data, kicks = c(1, 1), par0 = list("beta",1))  
[1] 0.9705401
```

## 5.2 Função 'est\_q '

A distribuição qui-quadrado inf, abordada no Capítulo 4, não possui uma função quantílica em forma fechada, isto é, não é possível obter a função quantílica na forma analítica. Mas, precisamos saber calcular o quantil da distribuição qui-quadrado inf porque este será utilizado para determinar a região crítica do TRV aperfeiçoado. Diante desta necessidade criamos a função `est_q()`.

De um modo geral, a função `est_q()` é responsável por realizar o cálculo do quantil da distribuição qui-quadrado inf, dado uma probabilidade  $\gamma$ . A forma geral para tal função é fornecida abaixo com a descrição de cada um de seus argumentos.

```
est_q(fn, alpha = 0.05, bilateral = FALSE, step = 1e-04, c, k)
```

em que

- **fn**: Recebe a função densidade de probabilidade da distribuição qui-quadrado inf;
- **alpha**: Nível nominal adotado;
- **bilateral**: Se TRUE, retorna os quantis para um teste bilateral. O padrão considera `bilateral = FALSE`;
- **step**: Tamanho do passo da integral numérica responsável pela obtenção dos quantis da qui-quadrado inf. O padrão considera `step = 1e-04`;
- **c**: Parâmetro da distribuição qui-quadrado inf;
- **k**: Parâmetro da distribuição qui-quadrado inf.

Uma observação importante é que, dado o nível nominal  $\gamma$ , se `bilateral = FALSE` a função `est_q()` retornará o valor do quantil  $q$  com base na distribuição qui-quadrado inf com os parâmetros  $k$  e  $c$ , fixados, tal que

$$P(X > q) = \gamma.$$

Por outro lado, dado um nível nominal  $\gamma$ , se `bilateral = TRUE` a função retornará uma lista com dois quantis,  $q_1$  e  $q_2$ , com base na distribuição qui-quadrado inf com

os parâmetros  $k$  e  $c$ , fixados, tal que

$$P(X < q_1) = \frac{\gamma}{2} \quad \text{e} \quad P(X > q_2) = \frac{\gamma}{2}.$$

em que  $q_1$  é o quantil  $\frac{\gamma}{2}$  e  $q_2$  é o quantil  $(1 - \frac{\gamma}{2})$ .

**Exemplo:** Seja  $X$  uma variável aleatória com distribuição qui-quadrado inf. Para exemplificar, fixemos os parâmetros da distribuição em  $k = 1$  e  $c = 1$ . Suponhamos que o interesse seja de calcular o quantil  $1 - \gamma$ , dado  $\gamma = 0.05$  tal que

$$P(X > q) = 0.05. \quad (5.1)$$

Note que, neste caso, o argumento `bilateral = FALSE`. Primeiramente, implementamos a função de densidade de probabilidade da qui-quadrado inf que será passada como argumento para a função `est_q()`. Em seguida fazemos o uso da função `est_q()` onde temos como retorno o valor do quantil  $q = 5.0019$  de acordo com (5.1).

```
R> # Função densidade de probabilidade da qui-quadrado inf

R> fdp_chisq_inf <- function(par, x) {
+   k <- par[1]
+   c <- par[2]
+   dchisq(x = x, df = k) * (1 - (1 - pchisq(q = x, df = k)) ^ c +
+     c * pchisq(q = x, df = k) * (1 - pchisq(q = x, df = k)) ^ (c - 1))
+ }

R> # Fazendo o uso da função est_q()

R> est_q(fn = fdp_chisq_inf, alpha = 0.05, bilateral = FALSE, c = 1,
+   k = 1)
[1] 5.0019
```

**Exemplo:** Seja  $X$  uma variável aleatória com distribuição qui-quadrado inf. Para exemplificar, fixemos os parâmetros da distribuição em  $k = 1$  e  $c = 1$ . Suponhamos, agora, que o interesse seja de calcular os quantis  $q_1$  e  $q_2$  dado  $\gamma = 0.05$  tal que

$$P(X < q_1) = \frac{\gamma}{2} = 0.025 \quad \text{e} \quad P(X > q_2) = \frac{\gamma}{2} = 0.025. \quad (5.2)$$

em que  $q_1$  é o quantil 0.025 e  $q_2$  é o quantil 0.975, ambos com base na distribuição qui-quadrado inf com  $k = 1$  e  $c = 1$ .

Como podemos ver no script abaixo, os quantis da distribuição qui-quadrado inf que verificam as equações (5.2) e que estão de acordo com as condições do exemplo, são  $q_1 = 0.0398$  e  $q_2 = 6.2274$ .

```
R> est_q(fn = fdp_chisq_inf, alpha = 0.05, bilateral = T, c = 1, k = 1)
$q1
[1] 0.0398

$q2
[1] 6.2274
```

### 5.3 Função 'mc'

Nesta pesquisa trabalharemos com vários cenários de teste de hipóteses, baseados na estatística  $RV$ , em que cada um será replicado um número grande de vezes (réplicas de Monte Carlo) para podermos observar a taxa de rejeição da hipótese nula com base no quantil da qui-quadrado e a taxa de rejeição da hipótese nula baseado no quantil qui-quadrado inf. Pensando nisso, criamos a função `mc()`.

Esta função realiza  $N$  réplicas de Monte Carlo para o teste da razão de verossimilhanças. Dado um nível nominal  $\gamma$ , em cada uma das  $N$  iterações, será retornado 1 (um) se o valor da estatística de teste  $RV$  está acima do quantil da distribuição qui-quadrado e 0 (zero), caso contrário. Da mesma forma, será retornado 1 (um) se o valor da estatística de teste  $RV$  está acima do quantil da distribuição qui-quadrado inf e 0 (zero), caso contrário. Assim, temos o controle das taxas de rejeição da hipótese nula com base nos dois quantis.

Como dito em seções anteriores, a função `mc()` faz o uso da função `lrt()` e `est_q()`. Apresentamos a seguir, a forma geral para tal função, em que descrevemos cada um de seus argumentos.

```
mc(N = 1L, n = 50L, sig = 0.05, f, q, kicks, par0, ncores = 1L, p,
bilateral = FALSE, step = 0.001, ...),
```

em que

- **N**: Número de réplicas de Monte Carlo a ser considerada;
- **n**: Tamanho da amostra a ser considerada;
- **sig**: Nível de significância adotado.
- **f**: Função densidade de probabilidade considerada no teste. Essa função deverá ser implementada conforme o exemplo abaixo;

- **q**: Função responsável pela geração de observações de uma variável aleatória com função densidade passada para **f**;
- **kicks**: Vetor com os chutes iniciais utilizados para a otimização;
- **par0**: Uma lista contendo como primeiro elemento um vetor com o nomes dos parâmetros que serão fixadas sob hipótese nula e um vetor com os valores fixados para cada um dos respectivos parâmetros;
- **ncores**: Número de núcleos a ser considerado. Por padrão, **ncores** = 1L;
- **p**: Valor utilizado para controlar o parâmetro  $c$  da qui-quadrado inf. Este é o valor de  $\rho$  dado em (4.3). Por padrão **p** = 0.5 ;
- **bilateral**: Se TRUE, retorna os quantis para um teste bilateral. O padrão considera **bilateral** = FALSE;
- **step**: Tamanho do passo da integral numérica responsável pela obtenção dos quantis da qui-quadrado inf. O padrão considera **step** = 1e-3;
- **...**: Lista de argumentos que serão passados para a função passada à **q**. Isto é, os parâmetros para geração da amostra verdadeira, que serão passados para a função **q**.

A estatística  $\lambda^*(\mathbf{x})$  com base em uma amostra  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$  de tamanho  $n$  é calculada da mesma forma como é calculada no Teorema 2.2. Sabemos que, dado um nível nominal  $\gamma$ , é razoável rejeitar  $H_0$  se  $RV > q$  em que a constante  $q$  é obtida de modo que

$$P(RV > q) = \gamma.$$

A partir do Teorema 2.2, para  $n$  grande,  $RV$  segue distribuição qui-quadrado com  $k$  graus de liberdade, em que  $k$  é o número de parâmetros testados em  $H_0$ . Desta forma, podemos obter o quantil  $q$ , que determina a região crítica do teste, através da distribuição qui-quadrado com  $k$  graus de liberdade.

Como falado em seções anteriores, dado um nível nominal  $\gamma$ , propomos neste trabalho, rejeitar  $H_0$  se  $RV > q_{\inf}$  em que  $q_{\inf}$  é obtido de modo que

$$P(RV > q_{\inf}) = \gamma.$$

Lembrando que, agora, consideramos que  $RV$  segue distribuição qui-quadrado inf com parâmetros  $k$  e  $c = 0.5 \log(n)$ .

De um modo geral, a função **mc()** é responsável por replicar o teste da razão de verossimilhanças um número  $N$  de vezes. Inicialmente, após definir todos os

argumentos da função `mc()` são calculados os quantís  $q$  e  $q_{\text{inf}}$  com base no valor nominal  $\gamma$  dado ao teste. Para cada uma das  $N$  iterações temos o seguinte algoritmo:

- é gerado uma amostra  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$  de tamanho  $n$  com base na função geradora de números pseudo-aleatórios passada para o argumento `q`, em que foram fixados os valores dos parâmetros sob  $H_0$  e o restante dos parâmetros também fixados de acordo com o interesse do usuário;
- com base na amostra gerada é calculada a estatística  $RV$ ;
- se  $RV > q$  então a hipótese nula é rejeitada e será retornado 1, caso contrario retornará 0. O mesmo é feito com o quantil  $q_{\text{inf}}$ .

A função `mc()` tem como retorno uma `tibble` (um tipo de tabela de forma compacta), um vetor contendo a taxa de rejeição da hipótese nula associado ao TRV, e um vetor com a taxa de rejeição da hipótese nula associado ao TRV\*.

A `tibble` possui  $N$  linhas por 3 colunas. Na primeira coluna estão os valores da estatística  $RV$  calculada nas  $N$  iterações. A segunda e terceira coluna são compostas com os números 0 e 1. Sendo que na segunda, 1 representa que  $H_0$  foi rejeitado baseada no quantil da qui-quadrado e na terceira, 1 representa que  $H_0$  foi rejeitado com base no quantil da qui-quadrado inf.

**Exemplo:** Seja  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$   $n$  observações da variável aleatória  $X \sim \text{Exp-Weibull}(\alpha, \beta, \lambda)$  (Weibull Exponencializada), com f.d.p dada na equação (4.4) em que  $\alpha$ ,  $\beta$  e  $\lambda$  são parâmetros positivos e  $x > 0$ .

Considere que queremos testar  $H_0 : \alpha = 0.6$  e  $\lambda = 1$  *versus*  $H_1 : \alpha \neq 0.6$  ou  $\lambda \neq 1$ . Fixamos o parâmetro de perturbação em  $\beta = 1.4$ . Consideremos como nível nominal do teste  $\gamma = 0.05$  e o tamanho da amostra como  $n = 20$ . Suponhamos que o interesse seja de replicar este teste  $N = 10000$  vezes para observarmos a taxa de rejeição da hipótese nula  $H_0$  com base tanto no quantil da qui-quadrado quanto no quantil da qui-quadrado inf.

Neste exemplo, temos que o espaço paramétrico e o espaço paramétrico sob  $H_0$  são dados, respectivamente, por

$$\Theta = \{(\alpha, \beta, \lambda); \alpha > 0, \beta > 0, \lambda > 0\}, \quad \Theta_0 = \{(\alpha, \beta, \lambda); \alpha = 0.6, \beta > 0, \lambda = 1\}.$$

Neste exemplo, temos que a estatística da razão de verossimilhanças é dada por

$$\lambda^*(\mathbf{x}) = \frac{L(\tilde{\theta}; \mathbf{x})}{L(\hat{\theta}; \mathbf{x})},$$

em que  $\tilde{\theta} = (0.6, \tilde{\beta}, 1)$  e  $\hat{\theta} = (\hat{\alpha}, \hat{\beta}, \hat{\lambda})$  são os EMV restrito e irrestrito, respectivamente. Perceba que  $k = 2$ . A implementação para realizar o uso da função `mc()` é

dado no script abaixo.

```
R> # Função densidade de probabilidade

R> fdp_exp_weibull2 <- function(par, x, var = NULL){
+   alpha = par[1]
+   beta = par[2]
+   lambda = par[3]
+
+   if (is.list(var)) eval(parse(text = paste(var[[1]], " <- ",
+     unlist(var[[2]]), sep = "")))
+
+   ((beta*lambda)/alpha)*((x/alpha)^(beta-1))*exp(-(x/alpha)^beta)*
+   (1-exp(-(x/alpha)^beta))^(lambda-1)
+ }

R> # Função geradora de números pseudo-aleatórios

> r_exp_weibull2 <- function(n, alpha, beta, lambda){
+   u <- runif(n, 0, 1)
+   alpha*(-log(1-u^(1/lambda)))^(1/beta)
+ }

R> # Fazendo o uso da função mc()

R> RNGkind("L'Ecuyer-CMRG")
R> set.seed(1L)      #fixando uma semente no gerador
>
R> tictoc::tic()
>
R> mc(
+   N = 10000L,
+   n = 20,
+   sig = 0.05,
+   f = fdp_exp_weibull2,
+   q = r_exp_weibull2,
+   kicks = c(1, 1, 1),
+   par0 = list(c("alpha","lambda"), c(0.6,1)),
+   ncores = 4L,
```

```

+   alpha = 0.6,
+   beta = 1.4,
+   lambda = 1,
+   p = 0.6,
+   bilateral = F
+ )
|=====| 100%, Elapsed 01:05
>

R> # Retorno da função mc()

$result
# A tibble: 10,000 x 3
  lambda      sucess_chisq  sucess_inf
<dbl>         <dbl>         <dbl>
1  4.32             0             0
2  3.84             0             0
3  0.679            0             0
4  0.410            0             0
5  1.25             0             0
6  0.926            0             0
7  0.465            0             0
8  0.218            0             0
9  9.95             1             1
10 2.57             0             0
# ... with 9,990 more rows

$prop_chisq
[1] 0.0638

$prop_inf
[1] 0.0537
>

R> parallel::mc.reset.stream()
>

```

Observe que ao realizar o uso da função `mc()` para o exemplo em questão, obtemos que a taxa de rejeição da hipótese nula com base no quantil da qui-quadrado foi de 0.0638 e que a taxa de rejeição nula com base no quantil da qui-quadrado



inf foi de 0.0537. Daí, segue que para este cenário de teste de hipótese, a taxa de rejeição de  $H_0$  com base no quantil da qui-quadrado inf ficou mais próxima do nível nominal do teste, que neste caso é de 0.05.

# Capítulo 6

## Simulações e Exemplos Numéricos

### 6.1 Simulações

Nesta seção, objetivamos avaliar o desempenho do TRV aperfeiçoado (TRV\*) e compará-lo com o TRV clássico. Para isso, realizamos simulações de Monte Carlo e utilizamos como critério de comparação a taxa de rejeição empírica da hipótese nula.

Para realizarmos as simulações faremos o uso da função `mc()`, a qual é detalhada na Seção 5.3. Escolhemos três distribuições de probabilidade, a saber: Weibull Exponencializada, Gama e Weibull inversa de dois parâmetros. Em cada um dos cenários consideramos 20000 réplicas de Monte Carlo. Os tamanhos amostrais considerados no estudo foram  $n = 10, 15, 20, 25, 30, 40, 50, 100, 1000$  e os níveis nominais foram  $\gamma = 1\%$ ,  $\gamma = 5\%$  e  $\gamma = 10\%$ . Para cada distribuição, realizamos inferência sobre um e dois parâmetros.

Com relação as distribuições mencionadas anteriormente, já conhecemos a f.d.p da distribuição Weibull exponencializada na equação (4.4). Sendo assim, mencionaremos a função densidade de probabilidade das outras duas distribuições aqui não mencionadas.

Seja  $X$  uma variável aleatória com distribuição gama, onde denotamos por  $X \sim Gama(\alpha, \beta)$ . Sua f.d.p é dada por

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

em que  $\alpha$  e  $\beta$  são parâmetros positivos com  $x > 0$ .

Seja  $X$  uma variável aleatória que segue distribuição Weibull inversa de dois parâmetros, em que denotamos por  $X \sim Weibull\_inversa2(\alpha, \beta)$ . Sua f.d.p é dada por

$$g(x; \alpha, \beta) = \frac{\alpha}{\beta} \left( \frac{-x}{\beta} \right)^{\alpha-1} \exp \left\{ - \left( \frac{-x}{\beta} \right)^{\alpha} \right\}$$

em que  $\alpha$  e  $\beta$  são parâmetros positivos com  $x < 0$ .

### 6.1.1 Testes envolvendo a Dist. Weibull exponencializada

Para o primeiro cenário, consideramos a distribuição Weibull exponencializada, em que realizamos inferência sobre o parâmetro  $\lambda$ . Os outros parâmetros desta distribuição, os parâmetros de perturbação, são fixados em  $\alpha = 0.6$  e  $\beta = 1.4$ . Seja  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$   $n$  observações da variável aleatória  $X \sim \text{Exp} - \text{Weibull}(\alpha, \beta, \lambda)$ . Consideramos o seguinte teste

$$H_0 : \lambda = 1 \quad \text{versus} \quad H_1 : \lambda \neq 1.$$

Na Tabela 6.1 apresentamos as taxas de rejeição nulas para este cenário. Avaliando essa tabela, observamos que para os testes que envolviam amostras de tamanhos maiores que vinte, o TRV é liberal, ou seja, as taxas de rejeições estimadas estão acima do nível nominal especificado.

Ainda analisando a Tabela 6.1, percebemos que em algumas situações TRV\* se mostra liberal, e em outras, foi conservativo, isto é, as taxas de rejeições estimadas estão abaixo do nível nominal especificado. No entanto, em geral, os testes baseados no quantil da distribuição qui-quadrado inf apresentaram os melhores desempenhos quando comparados aos testes baseados no quantil da qui-quadrado, principalmente para amostras de tamanho pequeno a moderado, pois possuem taxas de rejeições estimadas mais próximas do nível nominal especificado. Por exemplo, para  $n = 15$ , ao nível de  $\gamma = 10\%$ , as taxas de rejeição correspondente ao TRV e ao TRV\* foram 13.035% e 10.03%, respectivamente. Concluimos também que as taxas de rejeição correspondentes a ambos os testes se aproximam do nível nominal especificado conforme o tamanho da amostra aumenta. O que indica que esses testes são assintoticamente equivalentes.

Vale ressaltar que para  $n = 10$  as taxas de rejeições nulas ficaram muito distantes dos níveis nominais especificados em alguns casos. Por exemplo, ao nível nominal de  $\gamma = 1\%$  as taxas de rejeição com base no TRV e TRV\* foram 0.210 e 0.045, respectivamente. Note também, que os tempos de simulações para  $n = 10$  são bem superiores do que os tempos dos outros tamanhos amostrais próximos deste. Um dos motivos desses resultados é que, neste caso, consideramos uma amostra de tamanho muito pequeno para uma distribuição de três parâmetros. Além disso, o algoritmo implementado para essas simulações realiza o descarte das amostras em que não há convergência. Logo, de acordo os tempos obtidos percebe-se que houve um grande

Tabela 6.1: Teste envolvendo a distribuição Weibull Exponencializada. Taxas de rejeição nulas (%), inferências sobre  $\lambda$ .

| n    | $\gamma$ (%) | TRV (%)      | TRV* (%)      | Tempo (seg.) |
|------|--------------|--------------|---------------|--------------|
| 10   | 1.0          | <b>0.210</b> | 0.045         | 101.114      |
|      | 5.0          | 7.575        | <b>3.725</b>  | 121.751      |
|      | 10.0         | 16.230       | <b>10.065</b> | 121.841      |
| 15   | 1.0          | <b>0.745</b> | 0.550         | 81.983       |
|      | 5.0          | 7.080        | <b>5.430</b>  | 99.792       |
|      | 10.0         | 13.035       | <b>10.030</b> | 99.615       |
| 20   | 1.0          | <b>0.975</b> | 0.875         | 74.816       |
|      | 5.0          | 6.540        | <b>5.480</b>  | 91.954       |
|      | 10.0         | 11.940       | <b>9.660</b>  | 88.761       |
| 25   | 1.0          | 1.085        | <b>0.990</b>  | 70.000       |
|      | 5.0          | 6.065        | <b>5.425</b>  | 86.911       |
|      | 10.0         | 11.595       | <b>9.840</b>  | 85.048       |
| 30   | 1.0          | 1.090        | <b>1.045</b>  | 69.352       |
|      | 5.0          | 5.960        | <b>5.370</b>  | 86.001       |
|      | 10.0         | 11.740       | <b>10.220</b> | 83.078       |
| 40   | 1.0          | 1.080        | <b>1.045</b>  | 67.592       |
|      | 5.0          | 5.780        | <b>5.300</b>  | 84.196       |
|      | 10.0         | 11.675       | <b>10.395</b> | 82.336       |
| 50   | 1.0          | 1.265        | <b>1.240</b>  | 68.612       |
|      | 5.0          | 6.035        | <b>5.710</b>  | 83.625       |
|      | 10.0         | 11.64        | <b>10.615</b> | 81.934       |
| 100  | 1.0          | 1.195        | 1.195         | 79.089       |
|      | 5.0          | 5.805        | <b>5.735</b>  | 95.276       |
|      | 10.0         | 11.030       | <b>10.540</b> | 97.590       |
| 1000 | 1.0          | 1.010        | 1.010         | 283.406      |
|      | 5.0          | 5.210        | 5.210         | 339.497      |
|      | 10.0         | 10.415       | <b>10.380</b> | 343.018      |

números de não convergências. Desta forma, não recomendamos realizar testes de hipóteses fazendo o uso do TRV e nem do TRV\* quando o tamanho amostral for  $n = 10$  para uma distribuição de três parâmetros, pois isso pode acarretar em inferências imprecisas.

Para o segundo cenário, continuamos fazendo o uso da distribuição Weibull exponencializada, em que realizamos inferências sobre os parâmetros  $\alpha$  e  $\lambda$ . Quanto ao parâmetro de perturbação, fixamos este em  $\beta = 1.4$ . Seja  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$   $n$  observações da variável aleatória  $X \sim Exp - Weibull(\alpha, \beta, \lambda)$ . Consideramos o seguinte teste de hipóteses

$$H_0 : \alpha = 0.6 \text{ e } \lambda = 1 \quad \text{versus} \quad H_1 : \alpha \neq 0.6 \text{ ou } \lambda \neq 1.$$

Tabela 6.2: Teste envolvendo a distribuição Weibull Exponencializada. Taxas de rejeição nulas (%), inferências sobre  $\alpha$  e  $\lambda$ .

| n    | $\gamma$ (%) | TRV (%)      | TRV* (%)      | Tempo (seg.) |
|------|--------------|--------------|---------------|--------------|
| 10   | 1.0          | <b>1.120</b> | 0.830         | 109.409      |
|      | 5.0          | 6.425        | <b>3.885</b>  | 89.278       |
|      | 10.0         | 14.190       | <b>8.240</b>  | 110.882      |
| 15   | 1.0          | <b>1.065</b> | 0.875         | 90.696       |
|      | 5.0          | 6.630        | <b>4.985</b>  | 73.189       |
|      | 10.0         | 13.030       | <b>9.760</b>  | 88.088       |
| 20   | 1.0          | 1.100        | 1.010         | 85.275       |
|      | 5.0          | 6.340        | <b>5.200</b>  | 66.987       |
|      | 10.0         | 12.265       | <b>9.880</b>  | 82.877       |
| 25   | 1.0          | 1.265        | <b>1.195</b>  | 79.322       |
|      | 5.0          | 6.095        | <b>5.410</b>  | 63.584       |
|      | 10.0         | 11.800       | <b>9.800</b>  | 77.353       |
| 30   | 1.0          | 1.100        | <b>1.060</b>  | 77.436       |
|      | 5.0          | 5.870        | <b>5.210</b>  | 63.409       |
|      | 10.0         | 11.325       | <b>9.800</b>  | 75.637       |
| 40   | 1.0          | 1.120        | <b>1.110</b>  | 75.760       |
|      | 5.0          | 5.745        | <b>5.370</b>  | 61.373       |
|      | 10.0         | 10.925       | <b>9.805</b>  | 73.574       |
| 50   | 1.0          | 1.185        | <b>1.170</b>  | 76.534       |
|      | 5.0          | 5.675        | <b>5.410</b>  | 61.889       |
|      | 10.0         | 11.110       | <b>10.180</b> | 74.730       |
| 100  | 1.0          | 1.210        | 1.210         | 88.086       |
|      | 5.0          | 5.540        | <b>5.490</b>  | 73.263       |
|      | 10.0         | 10.910       | <b>10.495</b> | 85.883       |
| 1000 | 1.0          | 1.120        | 1.120         | 317.787      |
|      | 5.0          | 5.185        | <b>5.180</b>  | 260.651      |
|      | 10.0         | 10.200       | <b>10.155</b> | 318.31       |

Na Tabela 6.2 apresentamos as taxas de rejeição nulas para este cenário. Avaliando essa tabela, observamos que para todos os tamanhos de amostra o TRV é liberal. Notamos, também, que o TRV\* apresentou o melhor desempenho do que o TRV, especialmente para amostras de tamanho pequeno a moderado. Por exemplo, para  $n = 15$ , ao nível nominal  $\gamma = 10\%$ , as taxas de rejeição correspondente ao TRV e ao TRV\* foram 13.03% e 9.76%, respectivamente. Concluimos também que as taxas dos testes baseados em ambos os quantis se aproximam do nível nominal especificado conforme o tamanho da amostra aumenta.

A Figura 6.1 mostra a distorção dos testes, ou seja, a diferença entre a taxa de rejeição estimada e o nível nominal correspondente. Aqui, consideramos o nível nominal  $\gamma = 10\%$ . O primeiro gráfico ilustra os resultados apresentados na Tabela 6.1. O segundo gráfico ilustra os resultados apresentados na Tabela 6.2. Essa figura,

evidencia toda análise feita anteriormente sobre a comparação do TRV e TRV\* e mostra claramente que as distorções dos testes baseados no TRV\* são os que mais se aproximam da linha de referência (ordenada em zero), logo esse teste teve o melhor desempenho.

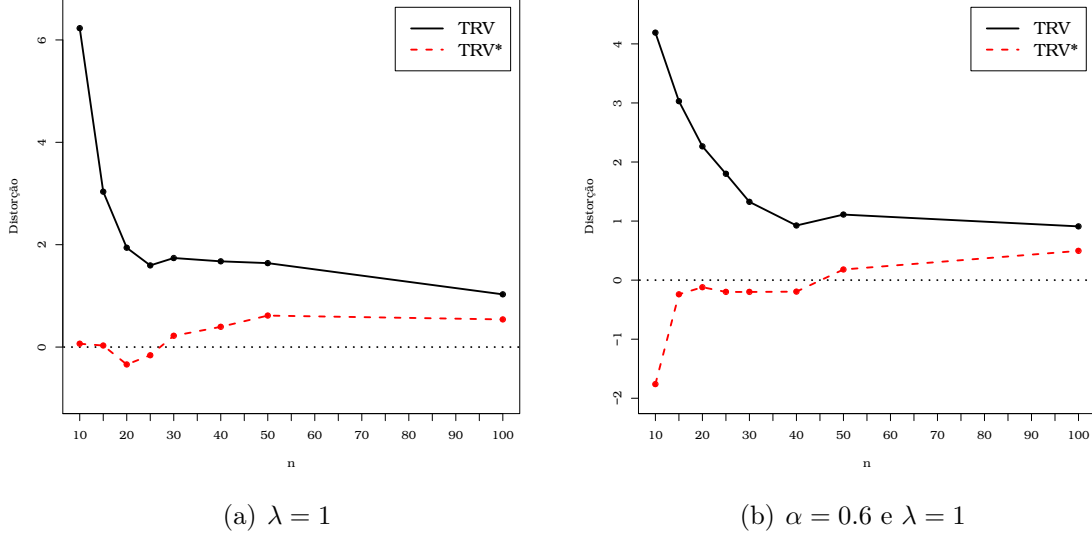


Figura 6.1: Distorções dos tamanhos dos testes envolvendo a distribuição Weibull Exponencializada:  $\gamma = 10\%$

### 6.1.2 Testes envolvendo a Dist. Gama

Para o terceiro cenário, fizemos o uso da distribuição Gama, em que realizamos inferência sobre o parâmetro  $\alpha$ . Fixamos o parâmetro de perturbação em  $\beta = 0.5$ . Seja  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$   $n$  observações da variável aleatória  $X \sim Gama(\alpha, \beta)$ . Consideramos o seguinte teste

$$H_0 : \alpha = 1 \quad \text{versus} \quad H_1 : \alpha \neq 1.$$

A Tabela 6.3 contém as taxas de rejeições nulas para este cenário. Avaliando essa tabela, observamos que para todos os tamanhos de amostra  $n \leq 100$  o TRV é liberal. Notamos que o TRV\* apresentou o melhor desempenho do que o TRV, especialmente para amostras de tamanho pequeno a moderado. Por exemplo, para  $n = 20$ , ao nível de  $\gamma = 5\%$ , as taxas de rejeição correspondente ao TRV e ao TRV\* foram 6.01% e 5.09%, respectivamente.

Para o quarto cenário, permanecemos fazendo o uso da distribuição Gama, em que realizamos inferência sobre  $\alpha$  e  $\beta$ . Seja  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$   $n$  observações da

Tabela 6.3: Teste envolvendo a distribuição Gama. Taxas de rejeição nulas (%), inferências sobre  $\alpha$ .

| n    | $\gamma$ (%) | TRV (%)      | TRV* (%)     | Tempo (seg.) |
|------|--------------|--------------|--------------|--------------|
| 10   | 1.0          | 1.830        | <b>1.385</b> | 29.709       |
|      | 5.0          | 7.395        | <b>5.030</b> | 30.128       |
|      | 10.0         | 13.400       | <b>8.900</b> | 30.090       |
| 15   | 1.0          | 1.555        | <b>1.310</b> | 31.140       |
|      | 5.0          | 6.315        | <b>4.485</b> | 30.601       |
|      | 10.0         | 11.825       | <b>8.960</b> | 31.850       |
| 20   | 1.0          | 1.365        | <b>1.245</b> | 32.743       |
|      | 5.0          | 6.010        | <b>5.090</b> | 31.598       |
|      | 10.0         | 11.345       | <b>9.265</b> | 33.425       |
| 25   | 1.0          | 1.225        | <b>1.170</b> | 32.482       |
|      | 5.0          | 6.005        | <b>5.300</b> | 32.833       |
|      | 10.0         | 11.250       | <b>9.475</b> | 34.585       |
| 30   | 1.0          | 1.240        | <b>1.190</b> | 32.824       |
|      | 5.0          | 5.725        | <b>5.145</b> | 33.669       |
|      | 10.0         | 11.025       | <b>9.570</b> | 34.804       |
| 40   | 1.0          | 1.160        | <b>1.125</b> | 34.573       |
|      | 5.0          | 5.330        | <b>5.030</b> | 34.582       |
|      | 10.0         | 10.675       | <b>9.655</b> | 36.856       |
| 50   | 1.0          | 1.090        | <b>1.075</b> | 36.036       |
|      | 5.0          | 5.265        | <b>4.985</b> | 37.397       |
|      | 10.0         | 10.565       | <b>9.690</b> | 37.570       |
| 100  | 1.0          | 1.055        | <b>1.045</b> | 45.670       |
|      | 5.0          | 5.095        | <b>5.005</b> | 45.997       |
|      | 10.0         | 10.165       | <b>9.865</b> | 47.626       |
| 1000 | 1.0          | 0.985        | <b>0.985</b> | 187.898      |
|      | 5.0          | <b>4.945</b> | 4.940        | 185.168      |
|      | 10.0         | <b>9.885</b> | 9.845        | 187.401      |

variável aleatória  $X \sim Gama(\alpha, \beta)$ . Consideramos o seguinte teste

$$H_0 : \alpha = 2 \text{ e } \beta = 0.5 \quad \text{versus} \quad H_1 : \alpha \neq 2 \text{ ou } \beta \neq 0.5.$$

A Tabela 6.4 é referente às taxas de rejeições nulas para este cenário. Analisando essa tabela, observamos que para todos os tamanhos de amostra  $n \leq 50$  o TRV é liberal. Notamos que o TRV\* apresentou o melhor desempenho do que o TRV, especialmente para amostras de tamanho pequeno a moderado. Por exemplo, para  $n = 10$ , ao nível de  $\gamma = 5\%$ , as taxas de rejeição correspondente ao TRV e ao TRV\* foram 6.755% e 4.35%, respectivamente.

A Figura 6.2 mostra a distorção dos testes. Aqui, consideramos o nível nominal  $\gamma = 10\%$ . O primeiro gráfico ilustra os resultados apresentados na Tabela 6.3. O segundo gráfico ilustra os resultados apresentados na Tabela 6.4. Essa figura,

Tabela 6.4: Teste envolvendo a distribuição Gama. Taxas de rejeição nulas (%), inferências sobre  $\alpha$  e  $\beta$ .

| n    | $\gamma$ (%) | TRV (%)       | TRV* (%)      | Tempo (seg.) |
|------|--------------|---------------|---------------|--------------|
| 10   | 1.0          | 1.595         | <b>1.120</b>  | 14.315       |
|      | 5.0          | 6.755         | <b>4.350</b>  | 11.255       |
|      | 10.0         | 12.570        | <b>8.225</b>  | 10.935       |
| 15   | 1.0          | 1.250         | <b>1.000</b>  | 14.960       |
|      | 5.0          | 5.945         | <b>4.675</b>  | 12.210       |
|      | 10.0         | 11.500        | <b>8.660</b>  | 11.306       |
| 20   | 1.0          | 1.190         | <b>1.060</b>  | 16.256       |
|      | 5.0          | 5.570         | <b>4.650</b>  | 12.162       |
|      | 10.0         | <b>11.115</b> | 8.870         | 11.965       |
| 25   | 1.0          | 1.160         | <b>1.050</b>  | 16.330       |
|      | 5.0          | 5.605         | <b>4.860</b>  | 12.934       |
|      | 10.0         | 10.955        | <b>9.160</b>  | 12.606       |
| 30   | 1.0          | 1.110         | <b>1.050</b>  | 16.945       |
|      | 5.0          | 5.370         | <b>4.820</b>  | 13.343       |
|      | 10.0         | 10.760        | <b>9.340</b>  | 13.429       |
| 40   | 1.0          | 1.100         | <b>1.080</b>  | 18.315       |
|      | 5.0          | 5.240         | <b>4.910</b>  | 14.164       |
|      | 10.0         | 10.860        | <b>9.780</b>  | 14.297       |
| 50   | 1.0          | 1.105         | 1.105         | 20.100       |
|      | 5.0          | 5.440         | <b>5.130</b>  | 15.334       |
|      | 10.0         | 10.600        | <b>9.690</b>  | 16.008       |
| 100  | 1.0          | 0.990         | 0.990         | 27.116       |
|      | 5.0          | 5.125         | <b>4.990</b>  | 21.408       |
|      | 10.0         | 10.315        | <b>9.890</b>  | 20.916       |
| 1000 | 1.0          | 1.095         | 1.095         | 149.007      |
|      | 5.0          | 5.105         | <b>5.095</b>  | 119.523      |
|      | 10.0         | 10.205        | <b>10.140</b> | 122.832      |

evidencia toda análise feita anteriormente sobre a comparação do TRV e TRV\* e mostra notoriamente que as distorções do testes baseados no TRV\* são os que mais se aproximam da linha de referência, logo esse teste teve o melhor desempenho.

### 6.1.3 Testes envolvendo a Dist. Weibull inversa de dois parâmetros

Para o quinto cenário, fizemos o uso da distribuição Weibull inversa de dois parâmetros, em que realizamos inferências sobre  $\beta$ . Fixamos o parâmetro de perturbação em  $\alpha = 2.5$ . Seja  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$   $n$  observações da variável aleatória  $X \sim \text{Weibull\_inversa}(\alpha, \beta)$ . Consideramos o seguinte teste

$$H_0 : \beta = 1 \quad \text{versus} \quad H_1 : \beta \neq 1.$$



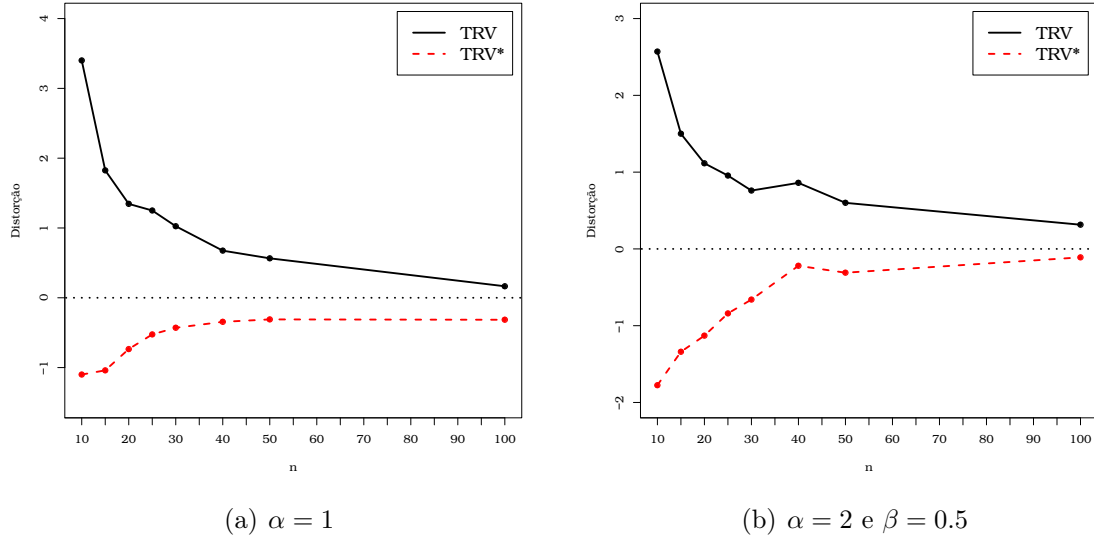


Figura 6.2: Distorções dos tamanhos dos testes envolvendo a distribuição Gama:  $\gamma = 10\%$

A Tabela 6.5 é referente às taxas de rejeições nulas para este cenário. Analisando essa tabela, observamos que para todos os tamanhos amostrais  $n \leq 100$  o TRV é liberal. Percebemos que o TRV\* apresentou o melhor desempenho do que o TRV, especialmente para amostras de tamanho pequeno a moderado. Por exemplo, para  $n = 15$ , ao nível de  $\gamma = 5\%$ , as taxas de rejeição correspondente ao TRV e ao TRV\* foram 6.17% e 4.985%, respectivamente.

Para o sexto e ultimo cenário, continuamos fazendo o uso da distribuição Weibull inversa de dois parâmetros, em que realizamos inferência sobre  $\alpha$  e  $\beta$ . Seja  $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top$   $n$  observações da variável aleatória  $X \sim Weibull\_inversa2(\alpha, \beta)$ . Consideramos o seguinte teste

$$H_0 : \alpha = 1 \text{ e } \beta = 1 \quad \text{versus} \quad H_1 : \alpha \neq 1 \text{ ou } \beta \neq 1.$$

A Tabela 6.6 é referente às taxas de rejeições nulas para este cenário. Analisando essa tabela, observamos que para todos os tamanhos amostrais o TRV é liberal. Percebemos que o TRV\* apresentou o melhor desempenho do que o TRV, especialmente para amostras de tamanho pequeno a moderado. Por exemplo, para  $n = 15$ , ao nível de  $\gamma = 5\%$ , as taxas de rejeição correspondente ao TRV e ao TRV\* foram 6.175% e 4.86%, respectivamente.

A Figura 6.3 mostra a distorção dos testes. Aqui, consideramos o nível nominal  $\gamma = 10\%$ . O primeiro gráfico ilustra os resultados apresentados na Tabela 6.5. O segundo gráfico ilustra os resultados apresentados na Tabela 6.6. Essa figura, evidencia toda análise feita anteriormente sobre a comparação do TRV e TRV\* e

Tabela 6.5: Teste envolvendo a distribuição Weibull Inversa. Taxas de rejeição nulas (%), inferências sobre  $\beta$ .

| n    | $\gamma$ (%) | TRV (%) | TRV* (%)     | Tempo (seg.) |
|------|--------------|---------|--------------|--------------|
| 10   | 1.0          | 1.750   | <b>1.220</b> | 37.694       |
|      | 5.0          | 7.140   | <b>4.700</b> | 29.014       |
|      | 10.0         | 12.800  | <b>8.510</b> | 34.391       |
| 15   | 1.0          | 1.560   | <b>1.270</b> | 37.763       |
|      | 5.0          | 6.170   | <b>4.985</b> | 29.615       |
|      | 10.0         | 11.715  | <b>8.720</b> | 35.500       |
| 20   | 1.0          | 1.385   | <b>1.310</b> | 39.143       |
|      | 5.0          | 6.070   | <b>5.150</b> | 30.695       |
|      | 10.0         | 11.370  | <b>9.275</b> | 36.435       |
| 25   | 1.0          | 1.340   | <b>1.245</b> | 39.487       |
|      | 5.0          | 5.925   | <b>5.225</b> | 29.825       |
|      | 10.0         | 11.185  | <b>9.380</b> | 37.076       |
| 30   | 1.0          | 1.145   | <b>1.090</b> | 38.341       |
|      | 5.0          | 5.645   | <b>5.180</b> | 30.714       |
|      | 10.0         | 10.830  | <b>9.430</b> | 37.252       |
| 40   | 1.0          | 1.125   | <b>1.110</b> | 39.392       |
|      | 5.0          | 5.385   | <b>5.100</b> | 32.207       |
|      | 10.0         | 10.440  | <b>9.365</b> | 38.170       |
| 50   | 1.0          | 1.0550  | <b>1.045</b> | 40.789       |
|      | 5.0          | 5.340   | <b>5.150</b> | 32.702       |
|      | 10.0         | 10.655  | <b>9.765</b> | 39.725       |
| 100  | 1.0          | 1.085   | <b>1.080</b> | 46.665       |
|      | 5.0          | 5.135   | <b>5.045</b> | 37.024       |
|      | 10.0         | 10.370  | <b>9.870</b> | 45.255       |
| 1000 | 1.0          | 0.995   | 0.995        | 136.515      |
|      | 5.0          | 4.880   | 4.880        | 114.221      |
|      | 10.0         | 9.810   | <b>9.850</b> | 130.955      |

mostra notoriamente que as distorções dos testes baseados no TRV\* são os que mais se aproxima da linha de referência, logo, novamente esse teste teve o melhor desempenho.

Com base em todas as tabelas podemos concluir que, em geral, quando aumentamos o número de parâmetros de perturbação as taxas de rejeição nulas também aumentam, o que mostra que a quantidade de parâmetros de perturbação tem impacto no desempenho do TRV e TRV\*.

## 6.2 Exemplos Numéricos

Nesta seção, mostraremos dois exemplos numéricos, em que utilizamos duas amostras completas. No primeiro, assumimos que as observações são independentes e

Tabela 6.6: Teste envolvendo a distribuição Weibull Inversa. Taxas de rejeição nulas (%), inferências sobre  $\alpha$  e  $\beta$ .

| n    | $\gamma$ (%) | TRV (%) | TRV* (%)      | Tempo (seg.) |
|------|--------------|---------|---------------|--------------|
| 10   | 1.0          | 1.705   | <b>1.210</b>  | 12.043       |
|      | 5.0          | 6.785   | <b>4.600</b>  | 11.892       |
|      | 10.0         | 12.315  | <b>8.130</b>  | 11.529       |
| 15   | 1.0          | 1.485   | <b>1.235</b>  | 12.351       |
|      | 5.0          | 6.175   | <b>4.860</b>  | 11.694       |
|      | 10.0         | 11.710  | <b>8.735</b>  | 11.179       |
| 20   | 1.0          | 1.280   | <b>1.185</b>  | 12.465       |
|      | 5.0          | 5.670   | <b>4.820</b>  | 12.333       |
|      | 10.0         | 11.125  | <b>8.935</b>  | 11.085       |
| 25   | 1.0          | 1.210   | <b>1.175</b>  | 12.532       |
|      | 5.0          | 5.835   | <b>5.090</b>  | 12.109       |
|      | 10.0         | 11.075  | <b>9.330</b>  | 11.714       |
| 30   | 1.0          | 1.120   | <b>1.085</b>  | 12.823       |
|      | 5.0          | 5.405   | <b>4.860</b>  | 12.379       |
|      | 10.0         | 10.765  | <b>9.320</b>  | 11.714       |
| 40   | 1.0          | 1.210   | <b>1.185</b>  | 13.282       |
|      | 5.0          | 5.450   | <b>5.040</b>  | 13.095       |
|      | 10.0         | 10.375  | <b>9.295</b>  | 12.033       |
| 50   | 1.0          | 1.075   | <b>1.060</b>  | 13.412       |
|      | 5.0          | 5.250   | <b>4.990</b>  | 13.225       |
|      | 10.0         | 10.460  | <b>9.550</b>  | 12.180       |
| 100  | 1.0          | 1.110   | 1.110         | 15.573       |
|      | 5.0          | 5.070   | <b>4.995</b>  | 14.944       |
|      | 10.0         | 10.220  | 9.780         | 14.284       |
| 1000 | 1.0          | 1.110   | 1.110         | 54.415       |
|      | 5.0          | 5.175   | <b>5.170</b>  | 53.693       |
|      | 10.0         | 10.130  | <b>10.085</b> | 51.744       |

provenientes da distribuição  $Gama(\alpha, \beta)$ . No segundo, assumimos que as observações são independentes e provenientes da distribuição  $Exp - Weibull(\alpha, \beta, \lambda)$ .

Realizamos um teste de Kolmogorov - Smirnov (KS) para os dois conjuntos de dados visando validar o que assumimos anteriormente. Para isso, fizemos o uso do pacote na linguagem R, **AdequacyModel** (DINIZ MARINHO *et al.*, 2016). Para o primeiro conjunto de dados testamos as seguintes hipóteses  $H_0$  : os dados seguem distribuição Gama *versus*  $H_1$  :  $H_0$  é falso. Aqui obtemos que  $p - valor = 0.8484$  (probabilidade de rejeitar  $H_0$ , dado que  $H_0$  verdadeira, com base nos dados amostrais). Logo, aos níveis nominais usuais (1%, 5% e 10%),  $H_0$  não é rejeitada, pois o  $p - valor$  excede tais níveis. Assim, há evidências fortes que os dados seguem distribuição Gama.

Para o segundo conjunto de dados testamos as seguintes hipóteses  $H_0$  : os dados

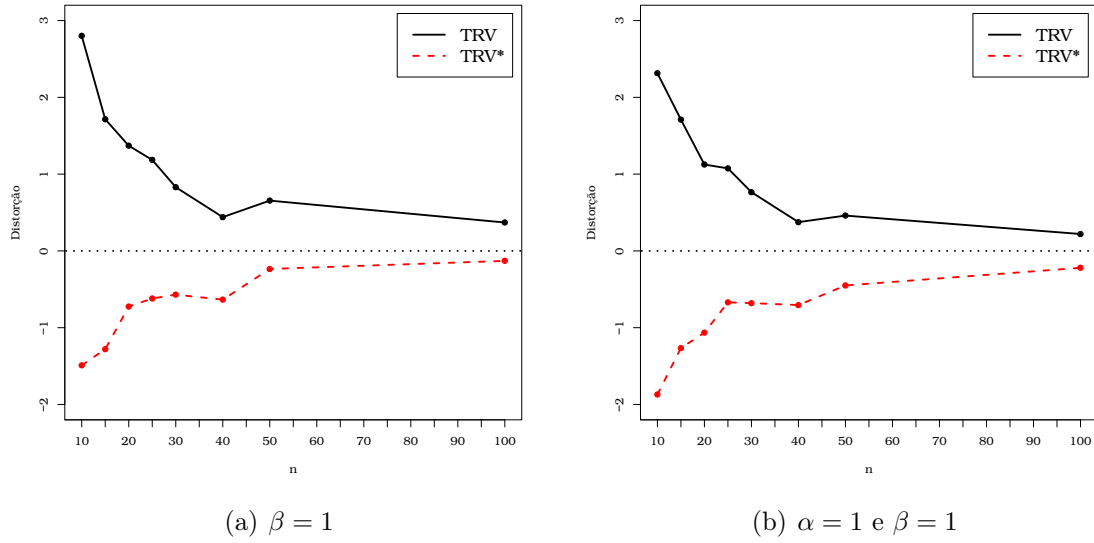


Figura 6.3: Distorções dos tamanhos dos testes envolvendo a distribuição Weibull Inversa de dois parâmetros:  $\gamma = 10\%$

seguem distribuição Weibull Exponencializada *versus*  $H_1 : H_0$  é falso. Aqui obtemos que  $p - \text{valor} = 0.7615$ . Logo, aos níveis nominais usuais,  $H_0$  não é rejeitada, pois o  $p - \text{valor}$  excede tais níveis. Assim, há evidências fortes que os dados seguem distribuição Weibull Exponencializada.

No primeiro exemplo, o interesse é testar a hipótese nula  $H_0 : \alpha = 1$  *versus*  $H_1 : \alpha \neq 1$ , ou seja, estamos testando o modelo *Exponencial*( $\beta$ ). Para este exemplo, vamos considerar o conjunto de dados Dados I, contido na Tabela 6.7, o qual consiste em 20 observações referentes a vinte itens testados até a falha (MURTHY *et al.*, 2004, p. 82). Aqui, obtemos que  $RV = 6.6363$ . Temos que o quantil 99% da qui-quadrado e da qui-quadrado inf, são  $q = 6.6349$  e  $q_{\text{inf}} = 6.7978$ , respectivamente. Logo, ao nível nominal de 1% o TRV rejeita a hipótese nula, enquanto TRV\* não rejeita a hipótese nula.

Tabela 6.7: Dados I: Vinte itens testados até a falha.

|       |       |       |       |      |       |      |      |       |       |
|-------|-------|-------|-------|------|-------|------|------|-------|-------|
| 11.24 | 11.50 | 30.42 | 38.14 | 1.92 | 8.86  | 9.17 | 2.99 | 12.74 | 7.75  |
| 10.20 | 16.58 | 22.48 | 5.73  | 5.52 | 18.92 | 9.60 | 9.37 | 5.85  | 13.36 |

No segundo exemplo, o interesse é testar a hipótese nula  $H_0 : \beta = 1$  e  $\lambda = 1$  *versus*  $H_1 : \beta \neq 1$  ou  $\lambda \neq 1$ , isto é, estamos testando o modelo *Exponencial*( $\frac{1}{\alpha}$ ). Consideramos agora, o conjunto de dados Dados II, contido na Tabela 6.8, o qual consiste em 30 observações referentes aos fluxos anuais do rio Weldon em Mill Grove, Missouri (1930-1959) (VUJICA, 1972). Esse mesmo conjunto de dados foi analisado por Bryson (1974). Aqui, obtemos que  $RV = 4.6635$ . Temos que o quantil 90% da

qui-quadrado e da qui-quadrado inf são  $q = 4.6052$  e  $q_{\text{inf}} = 4.9086$ , respectivamente. Logo, ao nível nominal de 10% o TRV rejeita a hipótese nula, enquanto TRV\* não rejeita a hipótese nula.

Tabela 6.8: Dados II: Fluxos anuais do rio Weldon em Mill Grove, Missouri (1930-59).

|       |       |       |       |       |       |       |       |       |       |
|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|
| 108.0 | 472.0 | 143.0 | 441.0 | 244.0 | 132.0 | 53.6  | 96.5  | 93.7  | 386.0 |
| 400.0 | 44.0  | 585.0 | 217.0 | 398.0 | 567.0 | 245.0 | 72.5  | 98.1  | 42.7  |
| 298.0 | 122.0 | 114.0 | 135.0 | 40.6  | 208.0 | 248.0 | 151.0 | 659.0 | 635.0 |

Considerando os exemplos numéricos apresentados, podemos perceber que os testes TRV e TRV\* levam a conclusões diferentes. Com base nos resultados das simulações feitas sobre os modelos Gama e Weibull Exponencializada, observamos que considerar o TRV para tomar decisões em amostras pequenas a moderadas nesses modelos pode levar a inferências imprecisas, enquanto que o TRV\* fornece melhores resultados em amostras das mesmas características. Dessa forma, recomenda-se fazer as conclusões com base no TRV\*.

# Capítulo 7

## Conclusões

Neste trabalho, inicialmente, estudamos as famílias de distribuições generalizadas sup e inf (TABLADA, 2017), onde obtivemos novas propriedades que acrescentam a mesma. Tais propriedades deram suporte a nossa proposta em que propomos a distribuição qui-quadrado inf como uma distribuição corretora da distribuição qui-quadrado para a obtenção do quantil que determina a região crítica do teste da razão de verossimilhanças.

Desenvolvemos o pacote **LikRatioTest** para a linguagem de programação R, o qual foi utilizado para realizar todos as simulações desta pesquisa. Tal pacote é Open Source e encontra-se disponível no GitHub para instalação no R e poderá servir para outros pesquisadores abordarem cenários diversos.

Realizamos um estudo de simulação de Monte Carlo fazendo o uso de tal pacote, impondo vários cenários, em que trabalhamos com três distribuições e fizemos inferências (teste de hipóteses) em um e dois parâmetros em cada uma destas. Para cada teste de hipóteses avaliamos o desempenho do TRV aperfeiçoado (TRV\*) e comparamos com o TRV clássico. Além disso, mostramos através de exemplos numéricos, envolvendo conjuntos de dados reais, a utilidade da nossa proposta, em que para os níveis nominais utilizados o TRV rejeita a hipótese nula, enquanto que o TRV\* não rejeita. Desta forma, fazer o uso do TRV aperfeiçoado pode nos levar a realizar inferências mais precisas sobre os parâmetros de interesse.

O estudo apontou que, em todos os cenários, as taxas de rejeição nulas dos testes fazendo o uso do TRV\* tiveram os melhores desempenhos, principalmente para amostras de tamanho pequeno a moderado, pois ficaram mais próximas dos níveis nominais especificados. Vale ressaltar também que as taxas dos testes baseados no TRV e no TRV\* se aproximam do nível nominal especificado conforme o tamanho da amostra aumenta. Daí, segue que nossa proposta, de acordo com nossas simulações, teve um bom desempenho tanto em amostras de tamanho pequeno quanto as de tamanho grande.

Portanto, com base em nossas simulações, nossa proposta que consiste em aper-

feioar o teste da razão de verossimilhanças quando o teste for baseado em amostra de tamanho pequeno ou moderado, foi validada em todos os cenários considerados. Além disso, pudemos perceber que nossa proposta é válida, não só para tamanhos de amostras pequenos a moderada, como também para as de tamanho grande, visto que o  $TRV^*$  tende a coincidir com o TRV clássico a medida que o tamanho da amostra aumenta.

# Referências Bibliográficas

- ALZAATREH, A., LEE, C., FAMOYE, F., 2013, “A new method for generating families of continuous distributions”, *METRON*, v. 71, pp. 63–79.
- BARRY, R. J., 1981, *Probabilidade: um curso em nível intermediário*. Rio de Janeiro: Instituto de Matemática Pura e Aplicada (Projeto Euclides).
- BARTLETT, M. S., 1937, “Properties of sufficiency and statistical tests”, *Proceedings of the Royal Society of London. Series A-Mathematical and Physical Sciences*, v. 160, n. 901, pp. 268–282.
- BRYSON, M. C., 1974, “Heavy-tailed distributions: properties and tests”, *Technometrics*, v. 16, n. 1, pp. 61–68.
- CORDEIRO, G. M., CRIBARI-NETO, F., 2014, *An introduction to Bartlett correction and bias reduction*. Springer.
- CORDEIRO, G. M., FERRARI, S. L. D. P., 1991, “A modified score test statistic having chi-squared distribution to order  $n-1$ ”, *Biometrika*, v. 78, n. 3, pp. 573–582.
- DEKKING, F. M., KRAAIKAMP, C., LOPUHAÄ, H. P., et al., 2005, *A Modern Introduction to Probability and Statistics: Understanding why and how*. Springer Science & Business Media.
- DINIZ MARINHO, P. R., BOURGUIGNON, M., BARROS DIAS, C. R., 2016, *AdequacyModel: Adequacy of Probabilistic Models and General Purpose Optimization*. Disponível em: <<https://CRAN.R-project.org/package=AdequacyModel>>. R package version 2.0.0.
- LEHMANN, E. L., ROMANO, J. P., 2006, *Testing statistical hypotheses*. Springer Science & Business Media.
- MAGALHÃES, M. N., 2006, *Probabilidade e variáveis aleatórias*. São Paulo: Edusp.



- MURTHY, D. P., XIE, M., JIANG, R., 2004, *Weibull models*, v. 505. John Wiley & Sons.
- R CORE TEAM, 2019, *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria. Disponível em: <<https://www.R-project.org/>>.
- ROUSSAS, G. G., 2003, *An introduction to probability and statistical inference*. Elsevier.
- SEVERINI, T. A., 2000, *Likelihood methods in statistics*. Oxford University Press.
- TABLADA, C. J., 2017, *Generalized probability distributions for lifetime applications*. Tese de Doutorado, Universidade Federal de Pernambuco.
- VUJICA, Y., 1972, “Probability and Statistics in Hydrology”, *Water Resources Publication*.
- WASSERMAN, L., 2013, *All of statistics: a concise course in statistical inference*. Springer Science & Business Media.
- WICKHAM, H., HESTER, J., CHANG, W., 2019, *devtools: Tools to Make Developing R Packages Easier*. Disponível em: <<https://CRAN.R-project.org/package=devtools>>. R package version 2.2.1.
- WILKS, S. S., 1938, “The large-sample distribution of the likelihood ratio for testing composite hypotheses”, *The annals of mathematical statistics*, v. 9, n. 1, pp. 60–62.