

MAYRA MARQUES BANDEIRA

**ESTIMAÇÃO DA CAPACIDADE DE PROCESSOS VIA TESTES
DE PERMUTAÇÃO E *BOOTSTRAP* EM APLICAÇÕES
INDUSTRIAIS**

Dissertação apresentada à Universidade Federal de Viçosa, como parte das exigências do Programa de Pós-Graduação em Estatística Aplicada e Biometria, para obtenção do título de *Magister Scientiae*.

**VIÇOSA
MINAS GERAIS - BRASIL
2015**

**Ficha catalográfica preparada pela Biblioteca Central da
Universidade Federal de Viçosa - Câmpus Viçosa**

T

B214e
2015
Bandeira, Mayra Marques, 19-
Estimação da capacidade de processos via testes de
permutação e bootstrap em aplicações industriais / Mayra
Marques Bandeira. - Viçosa, MG, 2015.
xi, 53f. : il. ; 29 cm.

Orientador : Fernando Luiz Pereira de Oliveira.
Dissertação (mestrado) - Universidade Federal de
Viçosa.
Referências bibliográficas: f.49-53.

1. Controle de qualidade - Aplicações industriais.
2. Estatísticas - Métodos gráficos. 3. Amostragem
(Estatística). 4. Controle de processos. I. Universidade
Federal de Viçosa. Departamento de Estatística. Programa
de Pós-graduação em Estatística Aplicada e Biometria.
II. Título.

CDD 22. ed. 658.562

MAYRA MARQUES BANDEIRA

**ESTIMAÇÃO DA CAPACIDADE DE PROCESSOS VIA TESTES
DE PERMUTAÇÃO E *BOOTSTRAP* EM APLICAÇÕES
INDUSTRIAIS**

Dissertação apresentada à Universidade Federal de Viçosa, como parte das exigências do Programa de Pós-Graduação em Estatística Aplicada e Biometria, para obtenção do título de *Magister Scientiae*.

APROVADA: 31 de Julho de 2015

José Ivo Ribeiro Júnior
(Coorientador)

Moyses Nascimento

Frederico Rodrigues Borges da Cruz

Fernando Luiz Pereira de Oliveira
(Orientador)

*À minha mãe, Miriam Louise,
por ser a base de tudo.*

DEDICO

Demore o tempo que for para
decidir o que você quer da vida, e
depois que decidir não recue ante
nenhum pretexto, porque o mundo
tentará te dissuadir.

Friedrich Nietzsche

Agradecimentos

Este trabalho não seria possível sem a ajuda de muitas pessoas, às quais agradeço imensamente.

Agradeço em primeiro lugar a Deus, pela força e sabedoria.

À minha mãe, que esteve sempre ao meu lado me apoiando durante toda esta caminhada. Por sempre acreditar que eu seria capaz de chegar até aqui.

Ao meu irmão Daniel e ao meu tio Airton pelo carinho e todo apoio.

Aos meus queridos companheiros do mestrado Matheus e Filipe pela convivência e amizade.

Às amigas de Viçosa, Juliana e Micherlania pela motivação e paciência todos os dias.

Aos amigos Lázaro e Rafael por me ajudarem sempre que precisei.

Ao meu orientador Fernando Oliveira, pela paciência, compreensão e pelo aprendizado durante esses anos.

Aos meus coorientadores José Ivo e Lupércio pela ajuda e contribuição.

Aos professores do Departamento de Estatística da UFV pelo convívio e pelos conhecimentos que me proporcionaram.

À todos aqueles, amigos e familiares, que torceram por mim.

Finalmente, agradeço à CAPES pelo apoio financeiro indispensável para a realização deste trabalho.

Sumário

Lista de Figuras	vii
Lista de Tabelas	ix
Resumo	x
Abstract	xi
Introdução	1
1 Referencial Teórico	4
1.1 Índices de Capacidade	4
1.2 Índices de Capacidade para Processos Univariados	5
1.2.1 Índice de Capacidade C_p	5
1.2.2 Índice de Capacidade C_{pk}	7
1.2.3 Índices de Capacidade C_{pm} e C_{pkm}	8
1.2.4 Índices de capacidade C_{pi} e C_{ps}	10
1.3 Índices de Capacidade Multivariados	11
1.3.1 Índices de Capacidade Média Geométrica	12
1.3.2 Índices de Capacidade de Niverthi e Dey	13
1.3.3 Índices de capacidade multivariados de Mingoti e Glória .	13
1.3.4 Extensões dos índices Niverthi e Dey e Mingoti e Glória .	14
1.3.5 Índice de capacidade proposto por Veevers	15
1.3.6 Índice de capacidade proposto por Taam, Subbaiah e Liddy	16
1.3.7 Índices propostos por Pan e Li	16

1.4	Índices de capacidade para processos autocorrelacionados	17
1.4.1	Índices Propostos por Pan e Huang	18
1.4.2	Índices Propostos por Mingoti e Oliveira	18
2	Métodos de Reamostragem	19
2.1	Teste de Permutação	19
2.2	Teste <i>bootstrap</i>	20
2.2.1	Intervalos de Confiança <i>bootstrap</i>	23
3	Análise de Dados	25
3.1	Banco de Dados	25
3.2	Análise dos Dados da VALE	26
3.3	Análise dos Dados PDPL	31
4	Conclusão	48
4.1	Considerações finais	48
	Referências Bibliográficas	49

Lista de Figuras

3.1	Boxplot do teor de sílica presente no concentrado após retirada de alguns pontos discrepantes	26
3.2	Histograma do teor de sílica presente no concentrado após retirada de alguns pontos discrepantes	27
3.3	Histograma dos índices de capacidade para 1.000 reamostragens pelo método da permutação	28
3.4	Histograma dos índices de capacidade para 1.000 reamostragens <i>bootstrap</i> com $n = 100, 200, 300, 500$ e 1000	30
3.5	Boxplot da variável A	31
3.6	Histograma da variável A	32
3.7	Boxplot da variável B	32
3.8	Histograma da variável B	32
3.9	Boxplot da variável C	33
3.10	Histograma da variável C	33
3.11	Histograma dos índices de capacidade $\widehat{C}_{pgeom}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para as variáveis A^*, B e C	35
3.12	Histograma dos índices de capacidade $\widehat{C}_{pgeom}$, para as variáveis A, B e C	36
3.13	Histograma dos índices de capacidade $\widehat{C}_{pkgeom}$, para a variável A	37
3.14	Histograma dos índices de capacidade $\widehat{C}_{pkgeom}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , paraa variável B	38
3.15	Histograma dos índices de capacidade $\widehat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20	39
3.16	Histograma dos índices de capacidade $\widehat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20	40
3.17	Histograma dos índices de capacidade $\widehat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20	40

3.18	Histograma dos índices de capacidade $\widehat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável A.	41
3.19	Histograma dos índices de capacidade $\widehat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável B.	42
3.20	Histograma dos índices de capacidade $\widehat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável C.	42
3.21	Histograma dos índices de capacidade $\widehat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável A.	43
3.22	Histograma dos índices de capacidade $\widehat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável B.	44
3.23	Histograma dos índices de capacidade $\widehat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável C.	44
3.24	Histograma dos índices de capacidade $\widehat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável A.	45
3.25	Histograma dos índices de capacidade $\widehat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável B.	46
3.26	Histograma dos índices de capacidade $\widehat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável C.	46

Lista de Tabelas

1.1	Classificação do processo com respeito a sua capacidade	8
3.1	Análise descritiva do teor de sílica no concentrado sem alguns pontos discrepantes	26
3.2	Proporção de itens fora do limite de especificação	26
3.3	Reamostragem pelo método da permutação	29
3.4	Reamostragem pelo método <i>bootstrap</i>	29
3.5	Análise descritiva da variável A	31
3.6	Análise descritiva da variável B	31
3.7	Análise descritiva da variável C	33
3.8	Limites de especificações	34
3.9	Proporção de itens fora dos limites de especificações	34
3.10	Reamostragem pelo método da permutação	35
3.11	Reamostragem pelo método <i>bootstrap</i>	36
3.12	Reamostragem pelo método da permutação	38
3.13	Reamostragem pelo método <i>bootstrap</i>	38
3.14	Reamostragem pelo método da permutação	41
3.15	Reamostragem pelo método <i>bootstrap</i>	43
3.16	Reamostragem pelo método da permutação	45
3.17	Reamostragem pelo método <i>bootstrap</i>	45
3.18	Proporção de itens fora dos limites especificados	47

Resumo

BANDEIRA, Mayra Marques, M.Sc. Universidade Federal de Viçosa, Julho de 2015. **Estimação da capacidade de processos via testes de permutação e *bootstrap* em aplicações industriais**. Orientador: Fernando Luiz Pereira de Oliveira. Coorientadores: José Ivo Ribeiro Junior e Lupércio França Bessegato.

A competitividade do mercado tem levado as empresas e indústrias a elevar a qualidade dos produtos e serviços fazendo com que ela se torne uma das preocupações em diversas áreas. O monitoramento e a avaliação da capacidade do processo torna-se indispensável para as empresas. Os índices de capacidade são medidas que traduzem a capacidade de um processo de atender as especificações estabelecidas. Alguns índices de capacidade disponíveis na literatura são apresentados neste trabalho. Nosso objetivo é obter as estimações pontuais e intervalares, de alguns índices de capacidade auxiliadas através de técnicas de reamostragem. Dentre estas técnicas de reamostragem, os testes de permutação e *bootstrap* são utilizados de forma a reorganizar o conjunto de dados sem alterar as informações existentes e assim obter reamostras para auxiliarmos nas estimativas via índices de capacidade. Nesta dissertação, utilizando dados reais, estimamos alguns índices de capacidade em situações em que suposições usuais para obter estas estimativas não são válidas.

Abstract

BANDEIRA, Mayra Marques, M.Sc. Universidade Federal de Viçosa, July, 2015. **Estimation of process capability via permutation and *bootstrap* tests in industrial applications.** Adviser: Fernando Luiz Pereira de Oliveira. Co-advisers: José Ivo Ribeiro Junior and Lupércio França Bessegato.

The competitiveness of the market have led companies and industries to raise the quality of products and services. As a result, the monitoring and evaluation of the process capability has become an indispensable tool for any company. The process capability index is measured for understanding the ability of a process to produce output within determined specifications. Some of the process capability indices which are available in literature are presented in this dissertation. Our goal is to get the point estimates and interval, some capability indices aided by techniques resampling. Among these techniques, permutation and *bootstrap* tests are used to reorganize the data set without changing the existing information and to obtain estimates for the parameters of interest. In this dissertation, using real data, we estimate some process capability index in situations in which the usual assumptions for these estimates are not valid.

Introdução

Montgomery (2004) define que o controle da qualidade compõe-se de ações que garantam qualidade dos meios, que controlam e medem características de um item, processo ou instalação, de acordo com as especificações estabelecidas. Visto que qualidade significa atender às especificações, segundo Ribeiro (2013), todos os itens de um mesmo produto ou serviço produzido por um mesmo processo ter que apresentar valores da variável resposta localizados dentro dos limites especificados. Segundo Mingoti *et al.* (2011), a qualidade está na capacidade de o produto em atender a determinadas necessidades do cliente, ou seja, de atender às especificações exigidas pelo cliente (consumidor) do produto/serviço. Desse modo, é necessário implantar mecanismos de monitoramento do processo e medidas que quantifiquem sua capacidade em atender às exigências dos clientes a que seus produtos se destinam. Tais medidas podem ser implementadas através do controle estatístico do processo, conhecido como CEP.

Montgomery (2004) define o CEP como a aplicação de um conjunto de técnicas estatísticas para determinar se o resultado do processo segue as especificações estabelecidas para um produto ou serviço. Segundo Souza (2010), o CEP é utilizado no monitoramento e detecção de alterações nos processos.

O CEP, de acordo com Paese *et al.* (2001), permite o monitoramento das características de interesse, proporcionando sua manutenção dentro de limites preestabelecidos e indicando quando ações de correção e melhoria devem ser realizadas. Do mesmo modo, permite a redução da variabilidade nas características da qualidade do produto.

Num ambiente competitivo, segundo Paese *et al.* (2001), o CEP abre-nos caminho para melhorias contínuas, garantindo um processo estável, previsível e com uma identidade e capacidade definidas. Através do CEP pode se obter menor variabilidade e, conseqüentemente, melhores níveis de qualidade nos processos. Para esses objetivos, Montgomery (2004) afirma que o CEP é extremamente útil, já que é uma poderosa coleção de ferramentas para a coleta, análise e interpretação de dados, com o intuito de melhorar a qualidade por meio da eliminação de causas especiais, podendo ser utilizado para a maioria dos processos. As causas especiais são causas controláveis, cuja influencia é grande na variação do processo. Há também as causas naturais (causas comuns) que são as causas inerentes ao processo.

Dos diversos motivos para a utilização do CEP, Montgomery (2004) afirma que :

- Permite prever até que ponto o processo manterá as tolerâncias estabelecidas;
- Auxilia os elaboradores/planejadores de um produto na seleção de modificações de um processo;
- Auxilia a estabelecer intervalos entre amostras para monitoramento de um processo;
- Permite especificar exigências de desempenho para um equipamento novo;
- Permite planejar a sequência de processos de produção quando há um efeito interativo de processo sobre as tolerâncias;
- Permite reduzir a variabilidade em um processo.

O CEP também permite realizar ações corretivas antes que não-conformidades ocorram. De acordo com Montgomery (2004), a não-conformidade é definida como defeito ou falha em algum produto ou serviço de acordo com a especificação estabelecida, permitindo a análise de como deveria estar o processo, ou se o processo está fora das especificações e como manter o processo sob controle estatístico.

Entre as ferramentas do CEP, uma muito usual são as cartas de controle, também conhecidas como gráficos de controle, propostas em 1924 por Shewhart. As cartas de controle são amplamente usadas para monitoramento de processos de acordo com Costa *et al.* (2005). Oliveira (2007) afirma que as cartas de controle de Shewhart mostram o que ocorre no processo, visto que os limites de controle são construídos com base na variabilidade existente no processo em termos da característica de interesse.

Os gráficos de controle são técnicas do CEP que, de acordo com Woodall e Montgomery (2014), estão entre as técnicas mais simples e eficientes. Seu objetivo, de acordo com Souza (2010), é verificar a estabilidade do processo. Porém, eles não indicam se este processo está atendendo às especificações estabelecidas de um produto ou clientes. Isto torna importante o uso dos índices de capacidade pois através dos resultados obtidos por eles podemos analisar a capacidade do processo.

Segundo Oliveira (2007), os índices de capacidade mensuram se o processo está atuando conforme as especificações, que podem ser fornecidas pelos clientes ou internamente no processo. Seu objetivo, de acordo com Ramos e Ho (2003), é avaliar a capacidade do processo produzir produtos que atendam às especificações estabelecidas, dado que o processo está sob controle estatístico (apresenta somente causas naturais). De acordo com Ribeiro (2013), não faz sentido a avaliação da capacidade de um processo fora de controle estatístico, pois o processo pode apresentar comportamento não previsível.

Existem na literatura diversos índices de capacidade. Uma das linhas de pesquisa são os índices de capacidade bayesianos que, de acordo com Woodall e Montgomery (2014), não são amplamente utilizados na área de controle

de qualidade, apesar do sucesso que fazem nas outras áreas da estatística. Desvantagens dos métodos bayesianos incluem sua complexidade e a quantidade de computação necessária para realizá-la. Segundo Woodall e Montgomery (2014), poucos métodos bayesianos têm sido propostos para o monitoramento de processos.

Segundo Pan e Li (2014), os índices de capacidade de processo têm sido muito utilizados na indústria para estimar medidas quantitativas de desempenho de processos que levam à melhoria da qualidade. Em diversas situações, para que estas estimativas sejam confiáveis, algumas suposições precisam ser analisadas. Como exemplo, uma forma de obter estimativas da capacidade de processos, quando estas suposições não são atendidas, é através dos métodos de reamostragem.

De acordo com a literatura é usual a utilização de métodos de reamostragem, como os testes de permutação e *bootstrap*, para obter distribuições empíricas. Segundo Costa (2006), com o uso da reamostragem, a distribuição assumida por amostragem é desconsiderada e gera-se uma distribuição empírica. Sendo assim as estimativas são obtidas nestas distribuições empíricas.

Tais procedimentos podem ser utilizados para auxiliar na estimativa da capacidade de processos.

Objetivos

- Proporcionar uma revisão sobre índices de capacidade;
- Obter estimativas intervalares da capacidade de processo para dois bancos de dados reais onde suposições teóricas não são satisfeitas, fornecendo assim informações para ajudar as tomadas de decisões de profissionais em diversas áreas.

Organização da dissertação

Esta dissertação foi organizada da seguinte maneira. No capítulo 1 apresentamos uma revisão de literatura sobre os índices de capacidade. No capítulo 2 apresentamos os métodos de reamostragens utilizados neste trabalho. No capítulo 3 apresentam-se os dados reais e os resultados obtidos das estimativas da capacidade usando os métodos usuais e pelas distribuições empíricas geradas pelas técnicas de reamostragem. Finalmente no capítulo 4 contendo as considerações finais e propostas de trabalhos futuros.

Capítulo 1

Referencial Teórico

1.1 Índices de Capacidade

Os índices de capacidade são medidas adimensionais que traduzem, em números, a capacidade de um processo atender às especificações estabelecidas por seus clientes, ou seja, avalia a capacidade do processo produzir ou não itens de acordo com as especificações. Segundo Bulba e Ho (2004), os índices de capacidade foram introduzidos na década de 70, desde que Juran (1974) apresentou o pioneiro índice de capacidade C_p . A suposição usual é que a característica de interesse seja observável com distribuição normal.

Capacidade de um processo diz respeito à variabilidade de um processo e, para estudar a capacidade, é necessário que o processo esteja sob controle estatístico. Um processo encontra-se sob controle estatístico quando ele apresenta somente causas naturais.

Segundo Hubele *et al.* (2005), um índice de capacidade deve ser calculado utilizando dados a partir de um processo estável, determinando estatísticas das amostras através de gráficos de controle. Uma vez que os gráficos mostram um grau razoável de estabilidade a capacidade do processo pode ser avaliada.

A importância da utilização dos índices de capacidade de acordo com Mingoti *et al.* (2011), é que mesmo um processo sob controle estatístico produz itens fora das especificações. Logo, não é suficiente colocar e manter um processo sob controle. É fundamental avaliar se o processo é capaz de atender às especificações estabelecidas a partir das necessidades dos clientes. É esta avaliação que constitui a análise da capacidade do processo, que é medida através da relação entre a variabilidade natural do processo em relação à variabilidade que é permitida a esse processo, dada pelos limites de especificação.

Diversos índices de capacidade podem ser encontrados na literatura. Segundo Ramos (2003), dois índices são frequentemente utilizados para o caso onde se têm apenas uma característica de interesse: C_p e C_{pk} .

Segundo Gonzalez e Werner (2009), os índices estabelecidos nos estudos

tradicionais de capacidade de processo procuram detectar dois tipos de problemas: o de localização do processo, se o processo atende em média ao valor nominal de especificação; e o de variabilidade, quando o processo apresenta muita dispersão e não atende às especificações que são estabelecidas no projeto. Segundo Montgomery (2004) e Deleryd (1999), são quatro índices de capacidade para dados normalmente distribuídos, C_p , C_{pk} , C_{pm} e C_{pkm} . De acordo com Ribeiro (2013), para verificar a capacidade do processo de acordo com esses índices, é necessário que o processo esteja sob controle estatístico e que a característica de interesse seja independente e normalmente distribuída, cujo monitoramento da variabilidade pode ser obtida a partir dos gráficos de controle.

Atualmente, muitos pesquisadores têm criado ou aperfeiçoado índices de capacidade. De acordo com Pan e Li (2014), nas últimas duas décadas, muitos pesquisadores e profissionais têm se dedicado a desenvolver índices de capacidade apropriado para vários processos de fabricação, por exemplo, em situações onde se têm mais de uma característica de interesse.

De acordo com o resultado e a interpretação dos índices de capacidade, segundo Souza (2010), tem-se um guia estratégico para melhoria da qualidade do processo, possibilitando diminuir custos de produção bem como obter clientes mais satisfeitos.

1.2 Índices de Capacidade para Processos Univariados

Os índices de capacidade para processos podem ser univariados, quando se estuda a capacidade de uma única característica de qualidade. Em processos univariados os índices C_p , C_{pk} , C_{pm} e C_{pkm} são frequentemente utilizados. Esses índices, segundo Oliveira (2007), são utilizados para processos univariados, ou seja, quando uma característica de qualidade está sendo avaliada e supondo que a sua distribuição de probabilidades seja normal. De acordo com Montgomery (2004), no controle estatístico da qualidade univariado, em geral, é usada a distribuição normal para descrever o comportamento de uma característica de qualidade contínua. A função de densidade de probabilidade da normal univariada é a seguinte:

$$f(x) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}, -\infty < x < \infty$$

Em que μ é a média e σ^2 , a variância.

1.2.1 Índice de Capacidade C_p

É a forma mais simples de avaliar a capacidade de um processo, chamado de índice de capacidade potencial do processo. Este índice, definido inicialmente por

Juran (1974), considera que o processo está centrado, segundo Gonzalez e Werner (2009) no valor nominal da especificação, desconsidera a centralização (média) do processo tratando apenas da sua variação. Não é sensível ao deslocamento (causas especiais) dos dados que afetam a média do processo. Segundo Ramos e Ho (2003) o C_p é definido como sendo a razão entre a tolerância de engenharia e a dispersão do processo. De acordo com Junior (2013), este índice relaciona aquilo que deve ser produzido com a variabilidade natural do processo, sendo calculado pela seguinte fórmula:

$$C_p = \frac{LSE - LIE}{6\sigma},$$

em que,

LSE: limite superior de especificação;

LIE: limite inferior de especificação;

σ : desvio-padrão do processo.

Se o desvio padrão (σ) for desconhecido, deve-se utilizar uma estimativa de σ , resultando em uma estimativa do índice C_p . Então:

$$\hat{C}_p = \frac{LSE - LIE}{6\hat{\sigma}}$$

Para ilustrar o cálculo de C_p , foi utilizado um exemplo (Montgomery [33]) de dados de diâmetro de anéis de pistão para motores de automóveis que são produzidos por um processo de forja. Vinte e cinco amostras, cada uma de tamanho cinco, foram extraídas desse processo quando se pensava que o mesmo estava sob controle.

Com a utilização de gráficos de controle $\bar{X}$ e R , as especificações dos anéis de pistão são $LSE = 74,05$ mm e $LIE = 73,95$ mm e $\hat{\sigma} = 0,01$. Assim, temos uma estimativa para C_p :

$$\hat{C}_p = \frac{LSE - LIE}{6\hat{\sigma}} = \frac{74,05 - 73,95}{6 \times 0,01} = 1,667.$$

Quando um processo apresentar somente um limite de especificação (limite de especificação unilateral), Kane (1996) propôs os índices unilaterais, C_{pS} , para o caso do processo apresentar limite superior, e o índice C_{pI} , no caso de se ter o limite inferior. Definidos a seguir:

$$C_{pS} = \frac{LSE - \mu}{3\sigma},$$

$$C_{pI} = \frac{\mu - LIE}{3\sigma}.$$

1.2.2 Índice de Capacidade C_{pk}

É um ajuste do índice C_p , chamado de índice de desempenho. Este índice, inicialmente proposto por Kane (1996), não considera a centralização do processo e é sensível ao deslocamento dos dados em relação a média. De acordo com Gonzalez e Werner (2009), na prática nem sempre o processo está centrado no valor nominal da especificação, então, o uso do índice C_p pode levar a conclusões erradas. O índice C_{pk} é definido como o mínimo dos índices, C_{pI} e C_{pS} , definido pela seguinte equação:

$$C_{pk} = \min(C_{pI}, C_{pS})$$

$$C_{pk} = \min\left(\frac{\mu - LIE}{3\sigma}, \frac{LSE - \mu}{3\sigma}\right)$$

em que, LSE: limite superior de especificação;

LIE: limite inferior de especificação;

μ : média do processo.

Se o desvio-padrão e a média do processo forem desconhecidos, a estimativa para o índice C_{pk} é calculada da seguinte forma:

$$\hat{C}_{pk} = \min\left(\frac{\bar{\mu} - LIE}{3\hat{\sigma}}, \frac{LSE - \bar{\mu}}{3\hat{\sigma}}\right)$$

Em casos de processos em que a média é centralizada na especificação, segundo Ramos e Ho (2003) e a característica de interesse obedece a uma distribuição normal, é válida a igualdade entre C_p e C_{pk} estimados.

Para ilustrar o cálculo de C_{pk} (Montgomery (2004)), foi suposto que um processo está sob controle com $\hat{\mu} = 100$, $\hat{\sigma} = 1,05$ e $n = 5$. As especificações do processo são $LSE = 105$ e $LIE = 85$. Assim, tem-se uma estimativa para C_{pk} :

$$\hat{C}_{pk} = \min\left(\frac{\bar{\mu} - LIE}{3\hat{\sigma}}, \frac{LSE - \bar{\mu}}{3\hat{\sigma}}\right)$$

$$\min\left(\frac{(100 - 85)}{3 \times 1,05}, \frac{(105 - 100)}{3 \times 1,05}\right) = \min(4,76; 1,59) = 1,587.$$

De acordo com Montgomery (2004), se $C_p = C_{pk}$, o processo está centrado no ponto médio das especificações e, quando $C_{pk} < C_p$, o processo está descentrado. Montgomery (2004) também afirma que, o índice C_p mede a capacidade potencial do processo enquanto o índice C_{pk} mede a capacidade efetiva.

Na literatura há diversos índices para análise de capacidade. Porém os índices de capacidade do processo frequentemente utilizados para processos univariados são o C_p (índice de capacidade potencial) e o C_{pk} (índice de capacidade efetiva), classificados de acordo com a Tabela 1.1.

Tabela 1.1: Classificação do processo com respeito a sua capacidade

Classificação	Valor de C_p/C_{pk}	itens não conformes (ppm)	Descrição
Capaz	$\geq 1,33$	$\leq 64\text{ppm}$	Está dentro dos limites especificados, produzindo praticamente todos os produtos com qualidade.
Razoavelmente Capaz	$1 \leq C_{pk} \leq 1,33$	de 64ppm a 1350ppm	Está sujeito a frequentes ocorrências de causas especiais, necessitando ser rigidamente controlado.
Incapaz	≤ 1	$> 1350\text{ppm}$	Está produzindo uma porcentagem considerável de itens fora das especificações, o processo está fora de controle.

FONTE: Adaptado de Costa *et al.* (2008)

1.2.3 Índices de Capacidade C_{pm} e C_{pkm}

O índice C_{pm} , proposto por Chan *et al.* (1988), é outra possibilidade aos índice C_p , sendo, de acordo com Ribeiro (2013), uma alternativa melhor, pois este considera a distância entre a média e o valor-alvo (T), além de considerar a variação do processo. Segundo Souza (2010), não basta produzir peças dentro dos limites, mas também é preciso que o processo esteja o mais próximo possível do alvo. O valor alvo é o valor especificado que o produtor/empresa deseja que tenha seu produto ou processo.

O índice C_{pm} possui uma vantagem em relação ao índice C_p . Segundo Gonçalves e Werner (2009), ele proporciona uma boa ideia da capacidade do processo, tanto para os processos com valores próximos ao valor nominal quanto para os que se apresentam mais distante dele.

Da mesma forma, o índice C_{pkm} proposto por Pearn *et al.* (1992), é outra possibilidade ao índice C_{pk} , sendo, de acordo com Ribeiro (2013), uma alternativa melhor, pois além de considerar a menor distância entre a média e o valor-alvo (T) ele a pondera em relação ao valor-alvo (T).

Gonçalves e Werner (2009) afirmam que ambos são alternativos aos índices C_p e C_{pk} e consideram além da variação do processo, a distância de sua média em

relação ao valor nominal da especificação, como seguem:

$$C_{pm} = \frac{(LSE - LIE)}{6\sqrt{\sigma^2 + (\mu - T)^2}},$$

$$C_{pkm} = \min \left(\frac{(\mu - LIE)}{3\sqrt{\sigma^2 + (\mu - T)^2}}, \frac{(LSE - \mu)}{3\sqrt{\sigma^2 + (\mu - T)^2}} \right)$$

e

$$T = \frac{LSE + LIE}{2},$$

em que,

LSE: limite superior de especificação;

LIE: limite inferior de especificação;

μ : é a média do processo;

σ : desvio-padrão do processo;

T : é o valor-alvo .

A estimação dos índices Cpm e Cpmk são dadas por:

$$\hat{C}_{pm} = \frac{(LSE - LIE)}{6\sqrt{s^2 + (\bar{x} - T)^2}},$$

$$\hat{C}_{pkm} = \min \left(\frac{(\bar{\mu} - LIE)}{3\sqrt{s^2 + (\bar{\mu} - T)^2}}, \frac{(LSE - \bar{\mu})}{3\sqrt{s^2 + (\bar{\mu} - T)^2}} \right)$$

e

$$T = \frac{LSE + LIE}{2},$$

em que,

LSE: limite superior de especificação;

LIE: limite inferior de especificação;

$\bar{\mu}$: é a estimativa da média (μ) do processo;

σ : é a estimativa desvio-padrão (σ) do processo;

T : é o valor-alvo.

Para ilustrar o cálculo dos índice C_{pm} e C_{pkm} , foi utilizado como exemplo um processo sob controle estatístico constituído por $\mu = 70$, $\sigma = 1$, $LIE = 65$, $LSE = 75$, $T = 70$ (Ribeiro (2013)).

$$C_{pm} = \frac{(75 - 65)}{6\sqrt{1^2 + (70 - 70)^2}} = 1,667.$$

$$C_{pkm} = \min \left(\frac{(70 - 65)}{3\sqrt{1^2 + (70 - 70)^2}}, \frac{(75 - 70)}{3\sqrt{1^2 + (70 - 70)^2}} \right) = \\ = \min (3,333; 1,667) = 1,667.$$

1.2.4 Índices de capacidade C_{pi} e C_{ps}

Os índices de capacidade C_{pi} e C_{ps} são utilizados, segundo Ribeiro (2013), quando o processo apresentar somente um dos limites de especificação (LSE ou LIE).

Quando somente o LIE é determinado, podemos utilizar o índice C_{pi} , definido por:

$$C_{pi} = \frac{\mu - LIE}{3\sigma}.$$

Quando somente o LSE é determinado, podemos utilizar o índice C_{ps} definido por:

$$C_{ps} = \frac{LSE - \mu}{3\sigma},$$

em que, μ é a média do processo; σ é o desvio padrão do processo.

A estimação dos índices C_{pi} e C_{ps} são dadas por:

$$\hat{C}_{pi} = \frac{\bar{x} - LIE}{3s_x}, \\ \hat{C}_{ps} = \frac{LSE - \bar{x}}{3s_x},$$

em que,

$\bar{x}$ é a média amostral; s é o desvio-padrão amostral.

O desenvolvimento de novos índices e o surgimento de novas abordagens, segundo Souza (2010), estão relacionados às necessidades das empresas, considerando a vantagem que pode ser obtida através da qualidade dos produtos e processos. Na próxima seção iremos definir alguns índices de capacidade utilizados em situações em que temos mais de uma característica de interesse.

1.3 Índices de Capacidade Multivariados

Na análise de um processo, existem ocasiões em que um grande número de variáveis deve ser analisado simultaneamente incorporando a correlação existente entre as características de interesse. Segundo Mingoti *et al.* (2011), é muito comum que a qualidade de um processo seja determinada por mais de uma característica, sendo então desejável avaliar a capacidade do processo por medidas que levem em consideração todas as características simultaneamente. Para contornar essa situação, de acordo com Oliveira (2007), existe uma tendência moderna para o uso dos métodos estatísticos multivariados.

Segundo Ferreira (2011), com os métodos estatísticos multivariados, as análises, descrições e inferências são realizadas com base nas respostas simultâneas, valendo-se a correlação entre as variáveis. Então a avaliação da capacidade do processo deve ser feita através dos índices multivariados, em que o conjunto de variáveis é representado na forma de vetores. A notação de vetor aleatório é utilizada para cada unidade amostral. Um vetor aleatório é um vetor cujos elementos são variáveis aleatórias, isto é:

$$\tilde{X}^t = (x_1, x_2, \dots, x_p),$$

em que, x_i é variável aleatória.

Uma matriz aleatória é uma matriz cujos elementos são variáveis aleatórias, isto é:

$$X = [x_{ij}],$$

em que, x_{ij} são variáveis aleatórias.

Na análise multivariada, uma amostra aleatória de tamanho n consiste em n observações independentes com mesma distribuição conjunta das p variáveis em cada repetição.

A abordagem usada no caso univariado também pode ser usada no caso multivariado. Supondo que tenhamos p variáveis, dadas por $x_1, x_2, \dots, x_p$, arranjam essas variáveis em um vetor de p componentes $x' = [x_1, x_2, \dots, x_p]$. A função densidade de probabilidade da normal multivariada é a seguinte:

$$f(x) = \frac{1}{(2\pi)^{\frac{p}{2}} |\Sigma|^{\frac{1}{2}}} e^{-\frac{1}{2}(x-\mu)'\Sigma^{-1}(x-\mu)},$$

em que, μ' é o vetor das médias dos $x's$, Σ é matriz de covariâncias de dimensão $p \times p$.

Temos então a distância padronizada ao quadrado de x à média para o caso multivariado definida por:

$$(x - \mu)' \Sigma^{-1} (x - \mu).$$

A matriz de covariância Σ é denotada por:

$$\Sigma_{pp} = \begin{bmatrix} \sigma_{11} & \sigma_{12} & \dots & \sigma_{1p} \\ \sigma_{21} & \sigma_{22} & \dots & \sigma_{2p} \\ \dots & \dots & \dots & \dots \\ \sigma_{p1} & \sigma_{p2} & \dots & \sigma_{pp} \end{bmatrix},$$

no qual σ_{ij} é a covariância entre X_i e X_j , para $i \neq k$ e σ_{ij} é a variância de X_i se $i = k$.

O trabalho original em controle da qualidade multivariado foi feito por Hotelling (1947), que aplicou seus procedimentos a dados de visores de bombardeiros durante a Segunda Guerra Mundial, segundo Montgomery (2004). Posteriormente, surgiram vários artigos sobre procedimentos de controle para várias variáveis relacionadas.

Nos dias de hoje, o controle da qualidade multivariado é muito importante devido à grande necessidade de controle de duas ou mais características simultaneamente, justificando o aumento do uso dos métodos multivariados nos últimos anos. Os métodos multivariados, segundo Woodall e Montgomery (2014), são fundamentais quando se quer monitorar várias variáveis de qualidade e tirar proveito de todas as relações entre elas.

É possível avaliar a capacidade do processo multivariado, de acordo com Mingoti *et al.* (2011), através do cálculo dos índices Cp, Cpk e Cpm para cada uma das variáveis separadamente. Porém, dessa maneira não se considera a correlação existente entre as características de interesse.

1.3.1 Índices de Capacidade Média Geométrica

A média geométrica dos índices de capacidade univariados é uma das formas de obter os índices de capacidade para processos multivariados, definidos por:

$$C_{pgeom} = \left(\prod_{i=1}^p C_p(x_i) \right)^{1/p},$$

$$C_{pkgeom} = \left(\prod_{i=1}^p C_{pk}(x_i) \right)^{1/p}.$$

Porém, caso haja correlação entre as características de interesse, ela não estará sendo levada em consideração.

Há alguns índices multivariados como os de Chen (1994), Veevers (1998),

Niverthi e Dey (2000), Mingoti e Glória (2008), Bernardo e Irony (1996), dentre outros, que são extensões dos índices Cp e Cpk, de acordo com Mingoti *et al.* (2011).

1.3.2 Índices de Capacidade de Niverthi e Dey

Os índices de capacidade multivariados de Niverthi e Dey (2000), são definidos por:

$$C_{pND} = \frac{1}{2m} \Sigma^{-1/2} (LSE - LIE),$$

$$C_{pkND} = \Sigma^{-1/2} \min \left[\left(\frac{LSE - \mu^0}{m}; \frac{(\mu^0 - LIE)}{m} \right) \right],$$

em que,

$\Sigma^{-1/2} \Sigma^{-1/2} = \Sigma^{-1}$ e $\Sigma^{-1/2}$ é matriz que contém os pesos usados na composição das combinações lineares;

μ^0 o vetor de médias centradas e

$m = 3$ corresponde a um nível de confiança igual a 0,997, ou um nível de significância $\alpha = 0,003$.

Então temos que os índices C_{pND} e C_{pkND} são vetores de dimensão $p \times 1$, sendo p o número de variáveis. Uma media global da capacidade do processo pode ser apresentada como o mínimo dos valores obtidos nos vetores C_{pND} e C_{pkND} , segundo Mingoti e Glória (2008).

Um problema atual é que este índice ainda não possui valores de referências para comparação dos resultados obtidos, como os índices univariados possuem.

1.3.3 Índices de capacidade multivariados de Mingoti e Glória

Os índices de capacidade multivariados de Mingoti e Glória (2008) são definidos por:

$$C_p^m = \min \{ C_{pi}^m; i = 1, \dots, p \},$$

$$C_{pi}^m = \left[\frac{LSE_i - LIE_i}{2\sigma_i C_{r\alpha}} \right] = \left[\frac{r_i^2 + r_i^1}{2\sigma_i C_{r\alpha}} \right],$$

em que,

$C_{r\alpha}$ é o ponto crítico usado no teste estatístico de Hayter e Tsui [22] para o vetor de médias populacional;

μ_i^s é o valor nominal da variável aleatória X_i , $i = 1, \dots, p$;

σ_i desvio-padrão da variável aleatória X_i ;

$r_i^1 = \mu_i^s - LIE$;

$r_i^2 = LSE - \mu_i^s$;

LSE: limite superior de especificação;

LIE: limite inferior de especificação.

$$C_{pk}^m = \min \{ C_{pki}^m; i = 1, \dots, p \},$$

$$C_{pki}^m = \min \left(\frac{\mu_i^0 - LIE_i}{\sigma_i C_{r\alpha}}, \frac{LSE_i - \mu_i^0}{\sigma_i C_{r\alpha}} \right), i = 1, \dots, p,$$

em que,

μ_i^0 é o vetor de médias centradas ;

σ_i o desvio-padrão da variável aleatória X_i ;

$LIE = \mu_i^s - r_i^1$ e $LSE = \mu_i^s + r_i^2$.

O processo é considerado capaz quando esses C_{pki}^m ou C_p^m forem maiores ou iguais a 1 com um nível de confiança $(1 - \alpha)100\%$. A vantagem desses índices é que eles consideram possíveis desvios médios em relação aos valores médios especificados.

1.3.4 Extensões dos índices Niverthi e Dey e Mingoti e Glória

Mingoti *et al.* (2011), propuseram dois novos índices fundamentados nas proposições de Niverthi e Dey (2000 e Mingoti e Glória (2008) que são extensões do índice C_{pm} univariado para o caso multivariado via Niverthi e Dey e Mingoti e Glória. A ideia foi construir índices multivariados que levassem em consideração a diferença entre os vetores de médias do processo e de especificação, a estrutura de variabilidade e de correlação das características de qualidade do processo e os limites de especificação de cada variável. O primeiro índice que de acordo com Mingoti *et al.* (2011), segue a formulação de índice de Niverthi e Dey (2000) é definido por:

$$C_{pmA}^M = (\Sigma_{p \times p} + A)^{-1/2} \left[\frac{LSE - LIE}{2m} \right],$$

em que,

$A = (\mu^s - \mu^0)(\mu^s - \mu^0)'$ é uma matriz relacionada com os desvios das médias de especificação e de processo de cada característica de qualidade com dimensão $p \times p$;

$$\mu^s = (\mu_1^s, \mu_2^s, \dots, \mu_p^s) \text{ e } \mu^0 = (\mu_1^0, \mu_2^0, \dots, \mu_p^0).$$

Para processos centrados nas médias de especificação, de acordo com Mingoti *et al.* (2011), o valor de C_{pmA}^M será igual aos dos índices C_{pND} e C_{pkND} . m : é uma constante resultante da distribuição normal univariada padronizada sendo determinada conforme o nível de confiança que o usuário deseja utilizar para medir a capacidade do processo.

O segundo índice que, de acordo com Mingoti *et al.* (2011), segue a formulação do índice de Mingoti e Glória (2008) definido por:

$$C_{pmB}^M = \min \{C_{pmiB}^M; i = 1, 2, \dots, p\},$$

$$C_{pmiB}^M = \frac{LSE - LIE}{2C_{r\alpha}(\sigma_i^2 + (\mu_i^s - \mu_i^0)^2)^{1/2}} = \frac{r_i^1 + r_i^2}{2C_{r\alpha}(\sigma_i^2 + (\mu_i^s - \mu_i^0)^2)^{1/2}}.$$

Segundo Mingoti *et al.* (2011), quando o processo é centrado no vetor de médias o valor de C_{pmB}^M será igual C_p^m e C_{pk}^m .

1.3.5 Índice de capacidade proposto por Veevers

Os índices propostos por Veeves (1998) são definidos pelo produto dos índices univariados. Definidos como:

$$C_{pv} = \prod_{i=1}^n C_p(X_i)^{I_i},$$

em que, n é o número de variáveis e

$$I_i = \begin{cases} 0 & \text{se } C_p(x_i) \geq 1; \\ 1 & \text{se } C_p(x_i) < 1. \end{cases}$$

Então o índice C_{pv} existe se houver pelo menos um $C_p(x_i)$ menor que 1.

Para o caso de todos os índices $C_p(x_i)$ maiores que 1, deve-se utilizar :

$$C_{pv} = \frac{\prod_{i=1}^n C_p(x_i)}{\prod_{i=1}^n C_p(x_i) - \prod_{i=1}^n [C_p(x_i) - 1]}.$$

1.3.6 Índice de capacidade proposto por Taam, Subbaiah e Liddy

Taam *et al.* (1993) propuseram índices de processos multivariados, MCP e MCPM, definidos por:

$$C_{pt} = \frac{(\prod_{i=1}^c r_i) \pi^{c/2} (\Gamma(\frac{c}{2} + 1))^{-1}}{|\Sigma|^{1/2} (\pi \chi_{1-\alpha, c}^2)^{c/2} (\Gamma(\frac{c}{2} + 1))^{-1}},$$

em que, c é o número de características de qualidade;

Σ matriz de covariância;

e Γ é uma função Gamma.

Quando a média processo é distante do valor alvo, o índice de C_{pt} pode ser reescrito da seguinte maneira:

$$C_{pt} = \frac{C_{pt}}{d},$$

no qual $d = (1 + (\mu - T)' \Sigma^{-1} (\mu - T))^{1/2}$,

μ é a média do processo e T é o valor-alvo.

1.3.7 Índices propostos por Pan e Li

A fim de avaliar adequadamente o desempenho do processo de precisão, Pan e Li (2014) propuseram o seguinte índice de capacidade de processos com variâncias heterogêneas:

$$NPC_a = \frac{\sum_{i=1}^p (\mu_i - T_i)^2}{R^2},$$

em que,

μ é a média do processo;

T é o valor alvo e

U é a área da região de especificação.

Uma vez que o valor do índice NPC_a é grande significa que a distância entre o processo e a localização do alvo é grande. Quando o valor do índice é pequeno, significa que a distância entre o processo e a localização do alvo é pequeno. O índice NPC_a pode fornecer a informação relativa de precisão do processo. Para medir o desempenho de precisão do processo, Pan e Li (2010) propuseram:

$$NPC_p = \frac{U^2}{C_p \sum_{i=1}^p \sigma_i^2},$$

em que,

$C_p = (\chi_{p,0.9973}^2)^{p/2}/p$ e $\chi_{p,0.9973}^2$ é o percentual de distribuição qui-quadrado com p graus de liberdade;

σ desvio padrão do processo e

U é a área da região de especificação.

Seus respectivos estimadores são dados por:

$$\widehat{NPC}_a = \frac{\sum_{i=1}^p (\hat{X}_i - t_i)^2}{U^2},$$

em que, $\hat{X}_i$ é a média da amostra de X_i característica de qualidade.

$$\widehat{NPC}_p = \frac{U^2}{C_p \sum_{i=1}^p s_i^2},$$

em que, $s_i^2 = \sum_{j=1}^n (x_j - \hat{x}_i^2)/(n-1)$ é a variancia amostral da i -ésima característica da qualidade.

1.4 Índices de capacidade para processos autocorrelacionados

Devido a autocorrelação, a estimação dos parâmetros podem ser afetadas, segundo Ramos e Ho (2003), sendo assim as estimativas dos índices C_p e C_{pk} estarão viesadas. Segundo Oliveira (2007), a presença de autocorrelação acarreta no problema para o monitoramento de processos através dos gráficos de controle de Shewart, pois eleva as estimativas do desvio padrão do processo. Autocorrelação é a correlação de uma variável, consigo mesma ao longo do tempo.

A correlação e sua intensidade existente nos dados pode ser medida, segundo Ramos e Ho (2003) utilizando o coeficiente de correlação, definido como:

$$\rho_L = \frac{Cov(X_t, X_{t-1})}{\sigma^2(X_t)}, L = 1, 2, \dots, n-1.$$

Mingoti e Oliveira (2011) mostraram a importância de considerar o efeito da autocorrelação no processo para o cálculo dos índices de capacidade multivariados. Segundo estes autores, a presença de autocorrelação é capaz de afetar o resultado fornecido pelos índices de capacidade quando a autocorrelação não está sendo considerada.

1.4.1 Índices Propostos por Pan e Huang

Se um processo autocorrelacionado é tratado como um processo independente, isso irá resultar em uma decisão inadequada e consequentemente levará a uma perda de qualidade, o que motivou Pan e Huang (2013) a desenvolverem novos índices de capacidade multivariados para dados autocorrelacionados. Definidos por:

$$MAC_p = \frac{\prod_{i=1}^{r_i}}{(\chi_{1-\alpha, v}^2 |\Gamma(0)|^{1/2}},$$

em que,

$$\Gamma(0) = \Phi\Gamma\Phi' + \Theta\Sigma\Theta' - \Theta\Sigma\Phi' + \Sigma \text{ (definido por Mingoti e Oliveira (2011))},$$

$$r_i = LSE - LIE/2 \quad i = 1, \dots, v,$$

v é o número de características de qualidade.

$$MAC_{pm} = \frac{MAC_p}{D^\Gamma} = \frac{MAC_p}{[1 + (\underline{\mu} - \underline{T})^T \Gamma(0)^{-1} (\underline{\mu} - \underline{T})]^{1/2}},$$

em que, D^Γ é o novo fator de correção com matriz autocovariância, $\underline{\mu}$ é a média do processo, e $\underline{T}$ é o vetor alvo.

1.4.2 Índices Propostos por Mingoti e Oliveira

Mingoti e Oliveira (2011) propuseram uma a extensão do índice de Veevers (1998) para processos não-centrados, definidos por:

$$C_{pkv} = \prod_{i=1}^n C_p k(X_i)^{I_i},$$

em que, n é o número de variáveis e

$$I_i = \begin{cases} 0 & \text{se } C_p k(x_i) \geq 1; \\ 1 & \text{se } C_p k(x_i) < 1. \end{cases}$$

O índice de capacidade C_{pkv} , segundo Mingoti e Oliveira (2011), é definido de modo que o seu valor seja sempre menor do $\prod_{i=1}^n C_p k(X_i)$.

Mingoti e Oliveira (2011) propuseram correções para alguns índices de capacidade de processos multivariados autocorrelacionados sugerindo uma alternativa que é realizar o cálculo dos índices usando a variância e covariância, ou seja, usar as informações contidas $\Gamma(0)$, em vez da matriz Σ_{pxp} , a fim de calcular os índices de capacidade.

Capítulo 2

Métodos de Reamostragem

Tais métodos são maneiras de reorganizar um conjunto de dados definido sem alterar as informações existentes.

Reorganizar o conjunto de dados observados, segundo Gentle (2009), pode dar uma indicação de como o conjunto de dados é incomum no que diz respeito a uma dada hipótese nula. Existem muitos procedimentos úteis para a análise dos dados que envolvem particionamento a amostra original, como o teste de permutação e o *bootstrap*. Usando os subgrupos da amostra completa, somos capazes de obter estimativas ou estatísticas de teste sem depender dos pressupostos que levaram a escolha destes. Segundo Costa (2006), com o uso da reamostragem obtém-se uma distribuição empírica desconsiderando a distribuição normal assumida de uma estatística. A fim de que a aplicação das técnicas de reamostragem resultem em valores confiáveis, segundo Domingues (2014), é necessário que a partir de uma amostra mestre sejam feitas milhares ou até centenas de milhares de reamostragens de tamanhos iguais. Na literatura é recomendado o uso de 1000 reamostragens.

2.1 Teste de Permutação

O teste de permutação, de acordo com Efron e Tibshirani (1993), foi introduzido por R.A. Fisher em 1930, como um argumento teórico de suporte ao teste t de Student, tratando-se da reamostragem mais simples, no caso não paramétrico, segundo Davison e Hinkley (1997).

De acordo com Davison e Hinkley (1997) um teste de permutação é um teste comparativo, onde a estatística de teste envolve algum tipo de comparação entre as funções de distribuições de probabilidade, sendo aplicado principalmente para o caso de duas populações, segundo Silveira (2008).

A ideia básica é, para cada permutação, calcular uma estatística de teste. Segundo Ferreira (2013), para utilizar o teste devemos primeiramente computar uma estatística com os dados originais.

A vantagem de se utilizar os métodos de permutação, de acordo com Polansky (2006), é que os níveis de significância dos testes de permutação são exatas, independentemente da distribuição dos dados de processo.

No teste de permutação é comparado o valor da estatística observada com as saídas apresentadas através de um algoritmo computacional, que consiste em criar, a partir dos dados originais, reamostras por simples permutações.

Pontes e Corrente (2001) definem os seguintes passos para a realização do método de permutação.

Primeiro passo: Calcula-se uma estatística de teste para os dados experimentais. Segundo passo: Os dados são permutados repetidamente e a estatística de teste é calculada para cada uma das permutações.

Os dados permutados constituem o conjunto referência para determinar a significância. A proporção de dados permutados no conjunto referência que têm a estatística de teste com valores maiores ou iguais (ou menores ou iguais) ao valor dos resultados obtidos experimentalmente fornece o p-valor (significância). Cada conjunto de dados permutados no conjunto de referência representa os resultados que teriam sido obtidos para uma particular atribuição no caso da hipótese nula ser verdadeira.

O número total de permutações é definido por:

$$N_{Perm} = \frac{N!}{n_1!, n_2, \dots, n_k!},$$

em que, $n_1!, n_2, \dots, n_k$ repetições e $N = \sum_{i=1}^k n_i$.

De acordo com Efron e Tibshirani (1993) a modernidade dos computadores faz com que os testes de permutação sejam práticos para usar em uma base de rotina, sendo a idéia básica, simples e livre de suposições matemáticas.

Uma de suas vantagens, é a fácil implementação de testes complexos. Segundo Ferreira (2013), a versatilidade é a principal vantagem, pois podem ser aplicados em diversas situações.

2.2 Teste *bootstrap*

O método *bootstrap* é um método de reamostragem que, segundo Efron e Tibshirani (1993), é definido como um método de simulação de dados baseado em inferência estatística, que pode ser usado para produzir inferências. Este consiste na reamostragem com reposição da amostra original, de acordo com Ferreira (2013), sendo um procedimento computacional intensivo de reamostragem, baseado na técnica da substituição, que possibilita estimar a distribuição amostral de estatísticas de interesse, segundo Filho *et al.* (2002).

Sua idéia, segundo Ramos e Ho (2003), consiste em que, por detrás de todo procedimento estatístico, existe sempre um processo físico que gera os dados em análise e, portanto, permite que se trabalhe diretamente com este modelo físico, através de simulação, ao invés da forma teórica da estatística indutiva clássica. De acordo com Maciel *et al.* (2013), o método *bootstrap* é utilizado para correção de vies de estimadores, construção de intervalos de confiança, testes de hipóteses, estimação do erro-padrão de um estimador, entre outros.

De acordo com Gentle (2009), os métodos de reamostragem envolvem o uso de várias amostras, tomadas a partir de uma única amostra que foi feita a partir da população de interesse, ou seja, a idéia é a partir de uma amostra original obter várias amostras e então extrair informações destas que sejam de interesse para a realização da inferência.

Para estimar parâmetros estatísticos desconhecendo a distribuição dos valores observados, uma maneira é utilizar a técnica da reamostragem aleatória, na qual cada observação de uma população tem a mesma probabilidade de ser incluída num conjunto reamostrado segundo Filho *et al.* (2002).

De acordo com Gentle (2009), a inferência é baseada em uma reamostragem utilizando a distribuição de amostragem condicional de uma nova amostra tirada a partir de uma dada amostra. Funções estatísticas sobre a amostra dada, um conjunto finito, pode ser facilmente avaliada. Métodos de reamostragem, portanto, pode ser útil, mesmo quando muito pouco se sabe sobre a distribuição.

Para Ramos e Ho (2003) esta é uma técnica geral que permite avaliar a incerteza na estimação de parâmetros, porém substitui a análise matemática pela amostragem (com reposição) dos dados amostrais originais. Na realidade, este método não deixa de ser uma forma diferente de simulação Monte Carlo, definido por métodos estatísticos que consistem em amostragens aleatórias para obter resultados numéricos, porém o método *bootstrap* utiliza a própria distribuição empírica dos dados.

De acordo com Gentle (2009), a ideia básica é que, porque a amostra observada contém toda a informação disponível sobre a população subjacente, a amostra observada pode ser considerada a população, assim, a distribuição de qualquer teste estatístico relevante pode ser simulada usando amostras aleatórias da população que consiste na amostra original.

As vantagens do método são suas aplicações em situações na qual os métodos tradicionais não poderiam ser aplicados, segundo Ferreira (2013), como no caso de distribuições inadequadas, distribuição de amostragem da quantidade amostral de interesse desconhecida entre outras. Sendo desvantagem a sua realização em amostras pequenas, pois garantias não podem ser fornecidas, segundo Ferreira (2013).

De acordo com Davison e Hinkley (1997), para utilização do método *bootstrap* é necessário a realização de um número grande de reamostragens e o cálculo de diversas estatísticas para cada uma destas reamostragens.

A atual disponibilidade de recursos de informática e o avanço dos softwares,

segundo Filho *et al.* (2002), permitiu que fossem desenvolvidas poderosas metodologias de reamostragem. O que agiliza os cálculos e a realização das reamostragens, contribuindo para o aumento do uso dessa técnica.

Na literatura existem muitos trabalhos que fazem uso de metodologias *bootstrap*. Em geral, o método *bootstrap* é utilizado para correção de vies de estimadores, construção de intervalos de confiança, testes de hipóteses, estimação do erro-padrão de um estimador, entre outros.

As metodologias *bootstrap* apresentam dois enfoques, sendo eles o paramétrico e o não-paramétrico. Segundo Maciel *et al.* (2013) o bootstrap paramétrico refere-se ao caso em que a reamostragem é feita com base em uma distribuição conhecida ou estabelecida e no bootstrap não-paramétrico não há o conhecimento da distribuição verdadeira dos dados.

Carrasco *et al.* (2012) também afirma que existem basicamente dois tipos de *bootstrap*: o paramétrico, no qual os estimadores de máxima verossimilhança (EMV) são obtidos através do modelo ajustado, isto é, geramos dados do modelo ajustado com os valores dos parâmetros fixados nos EMV obtidos da amostra original; e o *bootstrap* não paramétrico, onde os EMV são baseados em reamostras com reposição obtidas da amostra original.

O que as difere, segundo Domingues (2014) é a forma de obtenção da amostra, pois no caso paramétrico, a amostra será composta realizando-se a amostragem diretamente dessa distribuição e os parâmetros desconhecidos serão substituídos por estimativas paramétricas. Já no caso não paramétrico a amostra será composta retirando uma amostra com reposição de tamanho n , da amostra original.

A forma não paramétrica, segundo Cymrot e Rizzo (2006), é a mais utilizada e uma vantagem é poder aplicá-la de maneira geral, pois esta requer que menos suposições sejam feitas, tornando a técnica muito útil.

Ramos e Ho (2003) definem os seguintes passos para a realização do método de reamostragem:

Primeiro passo: deve-se considerar uma amostra constituída de n elementos da população em estudo $(x_i, i = 1, 2, 3, \dots, n)$.

Segundo passo: através de sorteio com reposição, retira-se novas amostras (chamadas de reamostras) de r valores ($r < n$) a partir da amostra original $(x_{ij}^*, i = 1, 2, 3, \dots, r; j = 1, 2, \dots, B)$.

Terceiro passo : calcula-se uma estimativa pontual do parâmetro de interesse com base nas observações reamostradas no passo 2.

Quarto passo : repetem-se os passos 1 a 3 B vezes.

Quinto passo : calcula-se a média amostral e o desvio padrão amostral das estimativas.

Filho *et al.* (2002), a reamostragem com reposição dos dados realizada um grande número de vezes, sendo estimados, para cada reamostragem, os

parâmetros de interesse torna-se possível obter uma distribuição empírica para estas estimativas *bootstrap*, bem como novas estimativas dos parâmetros de interesse, como a média, variância, intervalos de confiança, dentre outros. Essa técnica é chamada *bootstrap* não paramétrico, uma vez que as estimativas *bootstrap* são baseadas apenas nos dados, não sendo esperada distribuição de probabilidade para esses dados.

O teste bootstrap e o teste de permutação são eficientes para analisar testes de hipóteses segundo Silveira (2008). Ambos consistem numa estatística de teste tendo como vantagem a não imposição das suposições sobre a distribuição dos dados.

Segundo Efron e Tibshirani (1993), os testes *bootstrap* e os testes de permutação apresentam resultados similares. No teste de permutação a reamostragem é realizada sem reposição e no teste *bootstrap* a reamostragem é realizada com reposição.

2.2.1 Intervalos de Confiança *bootstrap*

A avaliação da incerteza sobre os valores dos parâmetros obtidos é feita através de intervalos de confiança, segundo Cunha e Colosimo (2003), intervalos aproximados de $100(1 - \alpha)\%$ de confiança para o parâmetro de interesse β podem ser obtidos pelo método *bootstrap*.

O intervalo de confiança *bootstrap* é utilizado em situações que outras técnicas não pode ser aplicadas, segundo Costa (2006), especialmente nas ocasiões em que o número de amostras é pequeno.

Um intervalo de confiança será definido por limites, tal que para qualquer um α ,

$$P(\theta < \hat{\theta}_\alpha) = \alpha.$$

O coeficiente de confiança de um intervalo é definido por $\gamma = 1 - (\alpha_1 + \alpha_2)$, e α_1 e α_2 são as probabilidades de erro pertencentes a cauda esquerda e direita, respectivamente. Para algumas aplicações é necessário apenas um limite, um limite de confiança inferior $\hat{\theta}_\alpha$ ou limite de confiança superior $\hat{\theta}_{1-\alpha}$, ambos com coeficiente de confiança $1 - \alpha$.

Segundo Davison e Hinklen (1997), a região de confiança única pode não nos fornecer um resumo adequado da incerteza sobre β , por isso, na prática, deve-se dar regiões para três ou quatro níveis de confiança entre 0,50 e 0,99, juntamente com a estimativa pontual para θ . Uma vantagem, segundo Davison e Hinklen (1997), é que qualquer assimetria na incerteza sobre o parâmetro será bastante clara.

O desejável é obter intervalos de confiança menores possíveis e nível de confiança maior possível, ou seja, próximo de 100%.

O intervalo $100(1 - \alpha)\%$ de confiança *bootstrap*-t para uma determinada estatística, de acordo com Cymrot e Rizzo (2006), é dado por,

$$IC_{boot.t} = [estatstica \pm t \times \hat{S}_{boot}],$$

em que, t é encontrado utilizando-se $(n - 1)$ graus de liberdade, sendo n o tamanho da amostra mestre e $\hat{S}_{boot}$, o desvio padrão bootstrap, obtido da seguinte maneira:

$$\hat{S}_{boot} = \sqrt{\frac{\sum_{b=1}^B (\hat{\theta}(b) - \hat{\theta}(\cdot))^2}{B - 1}}$$

O intervalo $100(1 - \alpha)\%$ de confiança *bootstrap* percentil pode ser obtido de duas maneiras:

A primeira, segundo Efron (1986), é encontrar os percentílico da média, ou seja, encontrar $100(1 - \alpha/2)\%$ e $100(\alpha/2)\%$ da média das reamostras da estatística do parâmetro que pretende estimar.

A segunda maneira, de acordo Montgomery e Runger (2003), é através dos percentis das diferenças dos valores das estatísticas das reamostras em relação ao valor médio da estatística nas reamostras. Então, a diferença entre esses valores, para cada reamostra i , é obtida da seguinte forma:

$$dif = \hat{\theta}(b) - \hat{\theta}(\cdot),$$

$$IC_{boot.percentil} = [\hat{\theta} - P_{97,5dif}; \hat{\theta} - P_{2,5dif}].$$

Há também outros intervalos de confiança baseados na técnica *bootstrap* como os intervalos de confiança percentílico BCa, intervalos de confiança BCPB e intervalos de confiança BCa, mais detalhes em Cunha e Colosino (2003) e Efron e Tibshirani (1993).

A partir dos métodos de reamostragem descritos, iremos obter intervalos de confiança *bootstrap* percentílico para obter intervalos de confiança para a média dos índices de capacidade estimados.

Capítulo 3

Análise de Dados

Dentre os objetivos desta dissertação encontra-se a obtenção de estimativas de alguns índices de capacidade, auxiliada pelos métodos de reamostragem *bootstrap* e de permutação. Acredita-se que esta abordagem, de análise de dados, pode fornecer um auxílio para a interpretação da capacidade de processo em setores industriais. Pelo uso dos procedimentos descritos nos capítulos anteriores, os dois bancos de dados apresentados a seguir serão objetos de análise.

3.1 Banco de Dados

Banco 1: O primeiro banco de dados é composto por 4.229 dados referentes ao processo de flotação da usina de tratamento de minério de ferro IB3, do Complexo Minerador de Mariana, MG, pertencente à Companhia Vale do Rio Doce (VALE), maior mineradora do Brasil e segunda maior mineradora mundial, referente ao ano de 2012. O objetivo é tornar igual a zero o teor de sílica presente no produto final (minério de ferro), o que é justificado pelo fato da sílica ser um dos principais contaminantes do minério de ferro. A variável-resposta (Y) é o teor de sílica no concentrado obtido de um processo de flotação com nove variáveis de entrada controláveis, sem limite inferior de especificação (LIE) e com limite superior de especificação (LSE) igual a 1,10%.

Banco 2: O segundo banco de dados é constituído por 31 dados referentes à produção de leite de fazendeiros da região de Viçosa no ano de 2014, pertencentes ao Programa de Desenvolvimento da Pecuária Leiteira da Região de Viçosa (PDPL), referente ao ano de 2014. O PDPL visa promover a evolução acelerada da produção econômica de leite. Tem como objetivo a maximização do retorno econômico dos produtores de leite, com sustentabilidade econômica, financeira, social, cultural e ambiental. As variáveis analisadas são denominadas por A , B e C . A variável A apresenta LIE igual a 531,57 e LSE igual a 4.907,93; B apresenta LIE igual a 15,04 e LSE igual 28,83; C tem LIE é igual a 0,86 e LSE igual 1,13.

3.2 Análise dos Dados da VALE

Análise Descritiva

Devido ao grande número de valores discrepantes do banco de dados, retiramos os valores acima de 1,67. Portanto, utilizamos 3.680 dados para esta análise. Conforme visto na Tabela 3.2, a média do teor de sílica considerando os dados menores que 1,67 é igual 0,8048 com desvio padrão igual a 0,2679. O coeficiente de variação (CV) é igual 33,29%, o que indica que os dados ainda são heterogêneos, conforme confirmado pelas Figuras 3.1 e 3.2.

Tabela 3.1: Análise descritiva do teor de sílica no concentrado sem alguns pontos discrepantes

Teor de sílica no concentrado (Y)	
Valor Mínimo:	0,350
Mediana:	0,740
Média:	0,805
Valor Máximo:	1,670
Desvio padrão:	0,268

Tabela 3.2: Proporção de itens fora do limite de especificação

Proporção	
Teor de Sílica	0,151

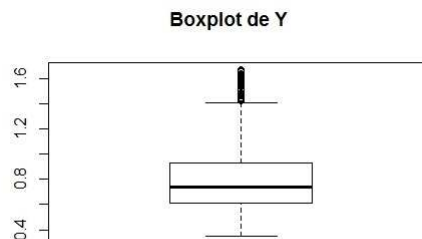


Figura 3.1: Boxplot do teor de sílica presente no concentrado após retirada de alguns pontos discrepantes

Mesmo após a retirada dos valores maiores que 1,67 ainda há valores discrepantes (*outliers*), o que explica o valor alto obtido pelo coeficiente de variação.

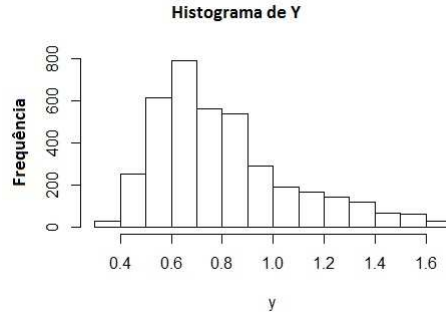


Figura 3.2: Histograma do teor de sílica presente no concentrado após retirada de alguns pontos discrepantes

Teste de Normalidade

Para verificar se os dados seguem uma distribuição normal utilizamos o teste Shapiro Wilk. As hipóteses testadas são:

$$\begin{cases} H_0 : & \text{os dados seguem distribuição normal;} \\ H_a : & \text{os dados não seguem distribuição normal.} \end{cases}$$

A estatística de teste (W) é dada pela seguinte equação.

$$W = \frac{b^2}{\sum_{i=1}^n (x_i - \bar{x})^2},$$

em que x_i são os valores da amostra, ordenados de forma crescente, e b é determinado por:

$$b = \begin{cases} \sum_{i=1}^{n/2} a_{n-i+1} \times (x_{(n-i+1)} - x_i), & \text{se } n \text{ é par,} \\ \sum_{i=1}^{(n+1)/2} a_{n-i+1} \times (x_{(n-i+1)} - x_i), & \text{se } n \text{ é ímpar.} \end{cases}$$

Obtivemos $W = 0,919$ e $p\text{-valor} < 2,2 \times 10^{-16}$. Considerando $\alpha = 0,05$, rejeitamos a hipótese nula ($p\text{-valor} < \alpha$), logo podemos acreditar que os dados não seguem distribuição normal.

Índice de Capacidade C_{ps}

Como o processo apresenta somente limite superior de especificação (LSE), utilizamos apenas o índice C_{ps} .

$$C_{ps} = \frac{1,100 - 0,805}{3 \times 0,268} = 0,367.$$

O processo é considerado capaz quando o valor de $\hat{C}_{ps}$ é maior que 1,25. Neste caso, nosso processo não é considerado capaz de atender as especificações estabelecidas mesmo após a retirada dos valores acima de 1,67.

Reamostragem pelo Método da Permutação

Realizamos 1.000 reamostragens pelo método da permutação para cinco tamanhos (n) diferentes de amostra para os dados menores ou iguais a 1,67. A Figura 3.3 apresenta os histogramas que permitem visualizar as distribuições empíricas para as 1.000 reamostragens pelo método da permutação, para os dados menores que 1,67.

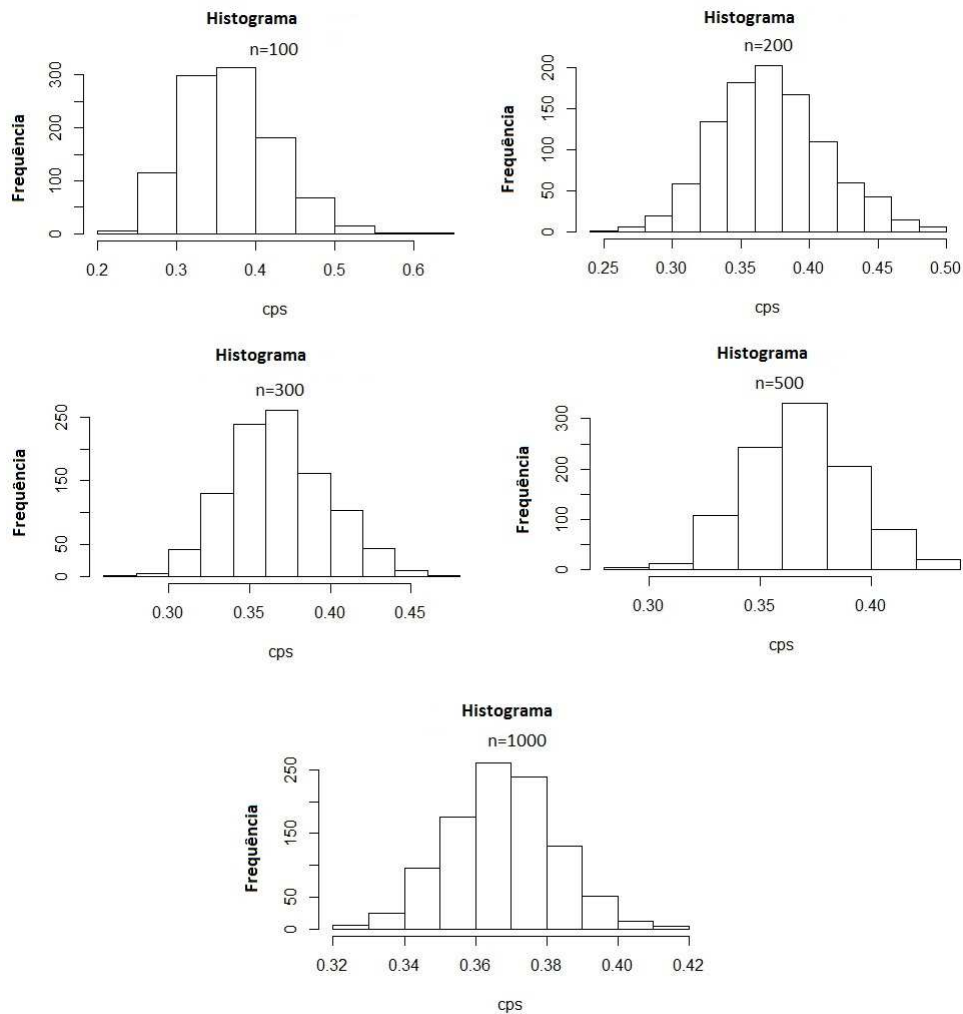


Figura 3.3: Histograma dos índices de capacidade para 1.000 reamostragens pelo método da permutação

A Tabela 3.3 apresenta os valores obtidos para a média dos $\widehat{C}_{ps}$, de acordo com o tamanho da amostra, bem como o erro quadrático médio associado (EQM) a elas, todas realizadas com reamostragem *bootstrap* de tamanho 1.000. O EQM é pequeno, o que significa que o estimador está com boa precisão. Quando aumentamos o tamanho da amostra menor é o valor do EQM. Os intervalos de confiança percentil de 95% para os índices de capacidade do processo foram construídos a partir da metodologia *bootstrap*, utilizando-se 1.000 reamostragens de tamanhos $n = 100, 200, 300, 500$ e 1000. Pelos intervalos de confiança *bootstrap* percentílicos construídos, pode ser observado que o processo é incapaz de atender as especificações.

Tabela 3.3: Reamostragem pelo método da permutação

n	$\widehat{C}_{ps}$		
	Média	EQM	$IC_{\text{percentílico}}$
100	0,367	0,003	(0,363 ; 0,370)
200	0,372	0,002	(0,369 ; 0,375)
300	0,368	0,001	(0,366 ; 0,370)
500	0,368	0,001	(0,367 ; 0,370)
1.000	0,368	0,001	(0,367 ; 0,368)

Reamostragem pelo Método *Bootstrap*

Realizamos reamostragens pelo método *bootstrap*, para diferentes tamanhos de amostra. Os histogramas que permitem visualizar as distribuições empíricas podem ser vistos na Figura 3.4.

A Tabela 3.4 apresenta os valores obtidos para a média dos C_{ps} , de acordo com o tamanho da amostra e o erro quadrático médio associado (EQM) a elas, todos obtidos para reamostragem *bootstrap* de tamanho 1.000. O EQM é pequeno o que significa que o estimador está com boa precisão. Os intervalos de confiança percentílico nos mostram que com 95% de confiança o processo não é capaz.

Tabela 3.4: Reamostragem pelo método *bootstrap*

n	$\widehat{C}_{ps}$		
	Média	EQM	$IC_{\text{percentílico}}$
100	0,371	0,003	(0,367 ; 0,374)
200	0,369	0,002	(0,366 ; 0,371)
300	0,369	0,001	(0,367 ; 0,371)
500	0,368	0,001	(0,366 ; 0,369)
1.000	0,367	0,000	(0,366 ; 0,368)

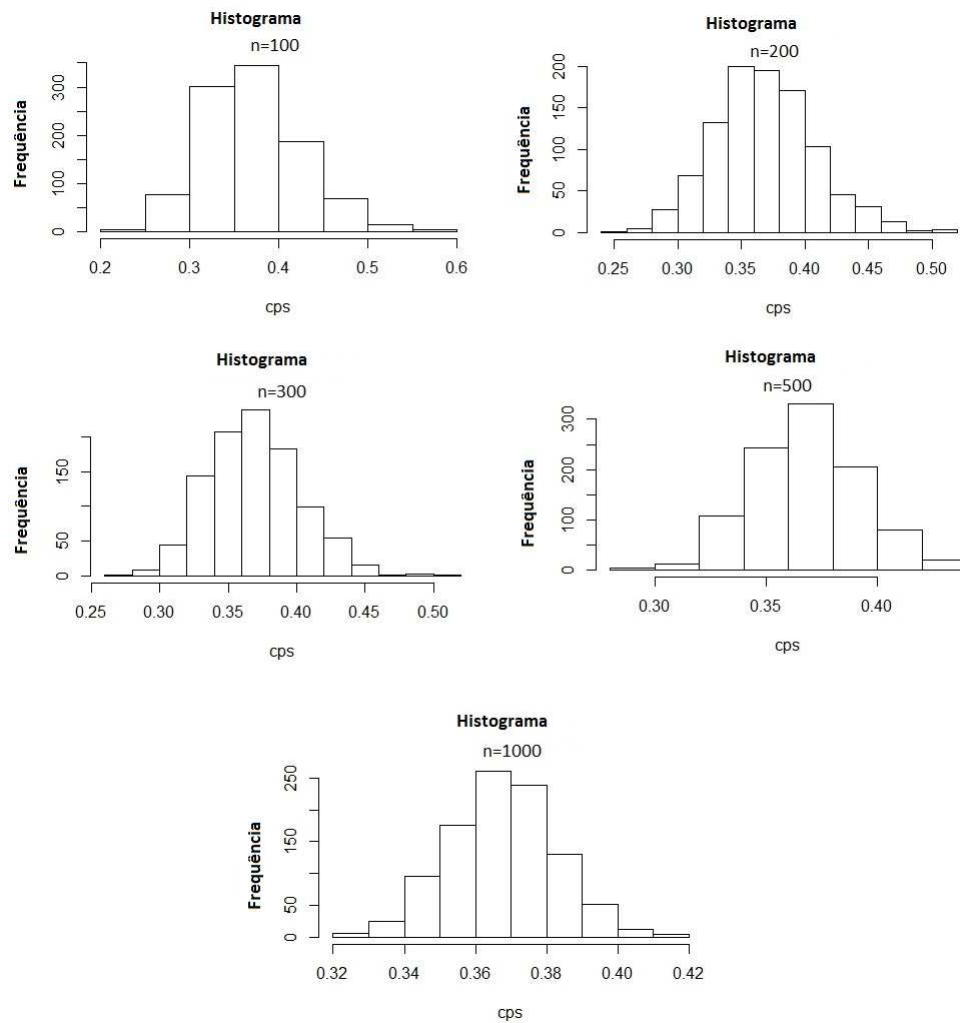


Figura 3.4: Histograma dos índices de capacidade para 1.000 reamostragens *bootstrap* com $n = 100, 200, 300, 500$ e 1000

3.3 Análise dos Dados PDPL

Análise Descritiva

Conforme visto na Tabela 3.5, a média da variável A é igual 758,500 com desvio padrão igual a 579,223. O coeficiente de variação (CV) é igual 76,364% , o que indica heterogeneidade dos dados e é justificado pela grande diferença entre os valores mínimo e máximo no processo (Figuras 3.5 e 3.6).

Tabela 3.5: Análise descritiva da variável A

	Variável A
Valor Mínimo:	178,000
Mediana:	532,800
Média:	758,500
Valor Máximo:	2305,000
Desvio padrão:	579,223

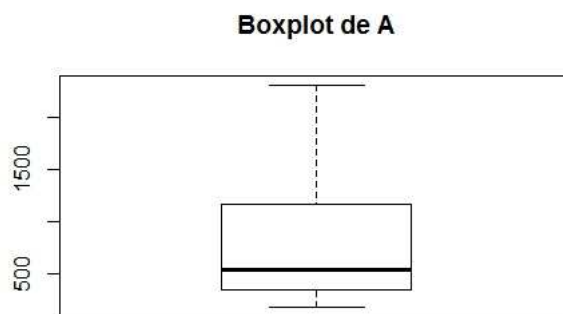
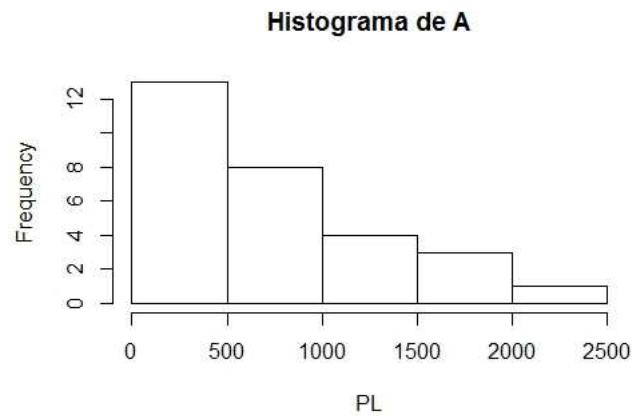
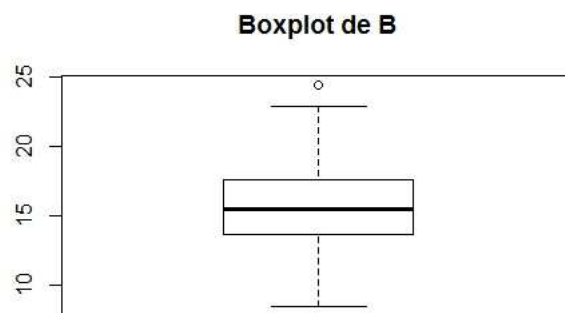
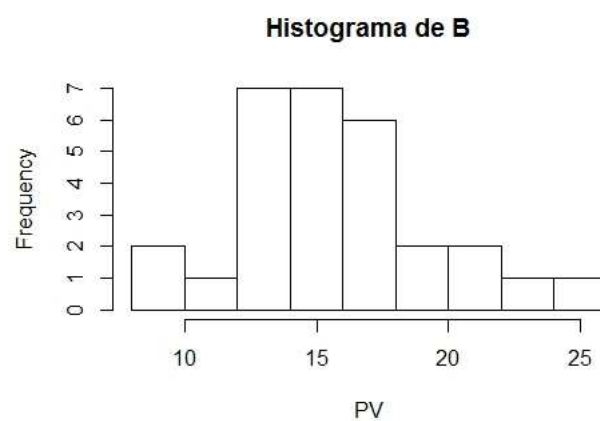


Figura 3.5: Boxplot da variável A

De acordo com a Tabela 3.6, a média da variável B é igual 15,770 com desvio padrão igual a 3,718, o que produz um coeficiente de variação (CV) igual a 23,576%. O CV indica homogeneidade dos dados, conforme visto nas Figuras 3.7 e 3.8.

Tabela 3.6: Análise descritiva da variável B

	Variável B
Valor Mínimo:	8,440
Mediana:	15,460
Média:	15,770
Valor Máximo:	24,480
Desvio padrão:	3,718

Figura 3.6: Histograma da variável A Figura 3.7: Boxplot da variável B Figura 3.8: Histograma da variável B

Conforme visto na Tabela 3.7, a média da variável C é igual 1,084 com desvio padrão igual a 0,113, resultando em um coeficiente de variação (CV) igual 10,424%. A homogeneidade dos dados fica aparente pelo CV e pelas Figuras 3.9

e 3.10.

Tabela 3.7: Análise descritiva da variável C

	Variável C
Valor Mínimo:	0,860
Mediana:	1,110
Média:	1,084
Valor Máximo:	1,290
Desvio padrão:	0,113

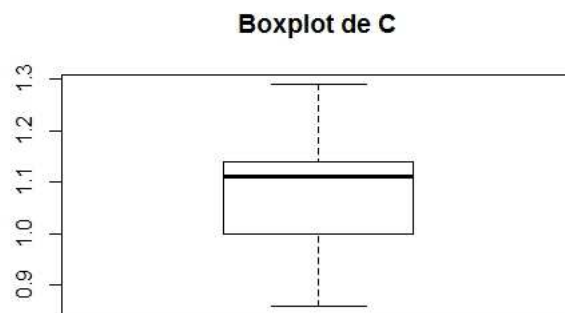


Figura 3.9: Boxplot da variável C

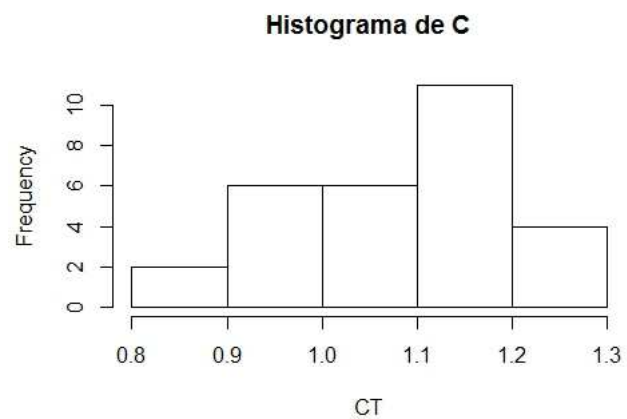


Figura 3.10: Histograma da variável C

Os limite de especificação definidos para as variáveis A , B e C podem ser vistos na Tabela 3.9.

Tabela 3.8: Limites de especificações

Variável	LIE	LSE
A	531,57	4907,93
B	15,040	28,830
C	0,860	1,130

Tabela 3.9: Proporção de itens fora dos limites de especificações

Variável	Proporção
A	0,460
B	0,586
C	0,655

Teste de Normalidade

Para verificar se os dados seguem uma distribuição normal utilizamos o teste Shapiro Wilk, considerando $\alpha = 0,05$. Para a variável A obtivemos $W = 0,845$ e p -valor $= 0,001$. Como p -valor menor que α , rejeitamos H_0 e podemos dizer que os dados da variável A *não* seguem uma distribuição normal. Para a variável B temos $W = 0,9731$ e p -valor $= 0,647$. Como o p -valor é maior que α , não rejeitamos H_0 e podemos dizer que os dados da variável B *seguem* uma distribuição normal. Finalmente, para a variável C temos $W = 0,968$ e p -valor $= 0,514$. Como o p -valor é maior que α , não rejeitamos H_0 e podemos dizer que os dados da variável C *seguem* uma distribuição normal.

Realizamos a transformação raiz quadrada nos dados da variável A . Verificamos a normalidade dos dados transformados (A^*), utilizando o teste Shapiro Wilk, considerando $\alpha = 0,01$. Para a variável A^* obtivemos $W = 0,913$ e p -valor $= 0,020$. Como p -valor maior que α , não rejeitamos H_0 e podemos dizer que os dados da variável A^* *seguem* uma distribuição normal.

Cálculo do Índice de Capacidade C_p Geométrico

Os índices de capacidade geométrico, C_{pgeom} , definidos abaixo, foram obtidos para as variáveis A^* , B e C , conforme apresentado no histograma da Figura 3.11 a 4.18.

$$C_{pgeom} = \left(\prod_{i=1}^p C_p(x_i) \right)^{1/p} = 0,581.$$

São apresentados na Tabela 3.10 os valores obtidos para a média dos $\hat{C}_{pgeom}$ de acordo com o tamanho da amostra e o erro quadrático médio associado (EQM) a elas, todas realizadas com reamostragem pelo método da permutação de

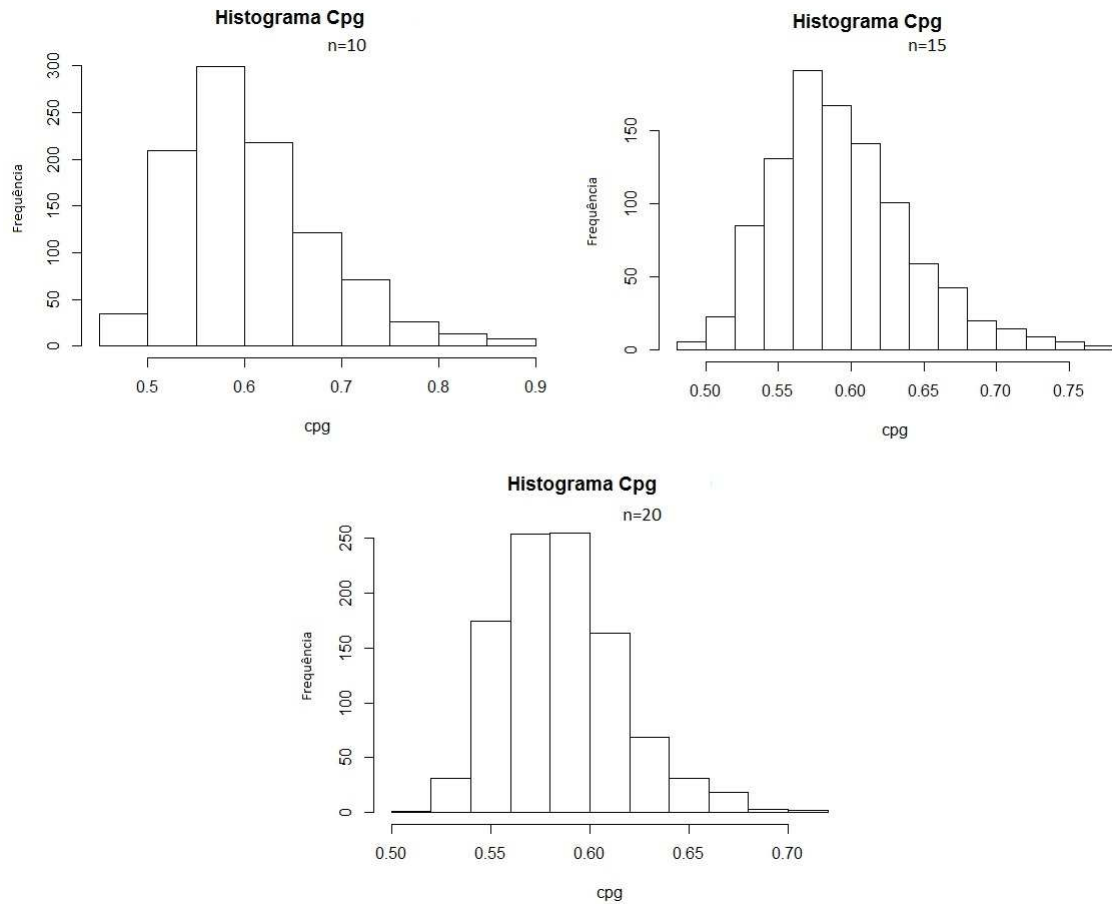


Figura 3.11: Histograma dos índices de capacidade $\hat{C}_{pgeom}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para as variáveis A*,B e C.

tamanho 1.000. O EQM diminui quando aumentamos o tamanho (n) da amostra. Os intervalos de confiança percentil de 95% para os índice de capacidade do processo foram construídos a partir da metodologia *bootstrap*, utilizando-se 1.000 reamostragens de tamanhos $n = 10, 15$ e 20 . Os intervalos de confiança *bootstrap* percentílicos nos mostram com 95% de confiança que o processo não é capaz.

Tabela 3.10: Reamostragem pelo método da permutação

n	$\hat{C}_{pgeom}$		
	Média	EQM	$IC_{\text{percentílico}}$
10	0,606	0,002	(0,601 ; 0,611)
15	0,593	0,002	(0,590 ; 0,596)
20	0,586	0,002	(0,584 ; 0,588)

Na Figura 3.12 temos apresentados os histogramas que permitem visualizar as distribuições empíricas para as 1.000 reamostragens pelo método da permutação, calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , respectivamente.

Na Tabela 3.11, são apresentados os valores obtidos para a média dos $\hat{C}_{pgeom}$ de acordo com o tamanho da amostra e o erro quadrático médio associado (EQM)

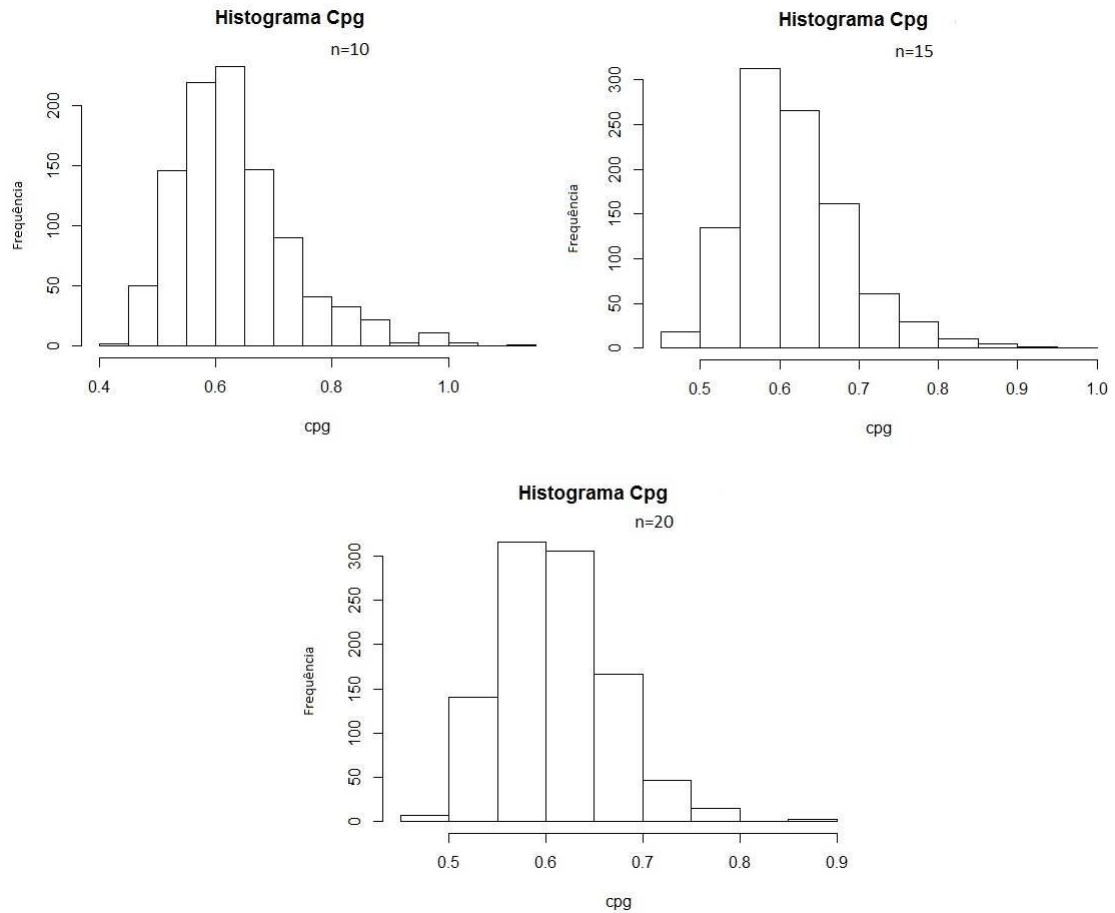


Figura 3.12: Histograma dos índices de capacidade $\hat{C}_{pgeom}$, para as variáveis A,B e C.

a elas, todas realizadas com reamostragem *bootstrap* de tamanho 1.000. O EQM diminui quando aumentamos o tamanho (n) da amostra. Os intervalos de confiança percentílicos nos mostram que com 95% de confiança o processo não é capaz.

Tabela 3.11: Reamostragem pelo método *bootstrap*

n	$\hat{C}_{pgeom}$		
	Média	EQM	$IC_{percentil}$
10	0,633	0,003	(0,584 ; 0,588)
15	0,616	0,001	(0,612 ; 0,620)
20	0,610	$6,641e^{-5}$	(0,607 ; 0,614)

Cálculo do Índice de Capacidade C_{pk} Geométrico

Os índices de capacidade geométricos $\hat{C}_{pkgeom}$, Figura 3.13, para as variáveis A , B e C são dados por:

$$C_{pkgeom} = \left(\prod_{i=1}^p C_{pk}(x_i) \right)^{1/p} = 0,028.$$

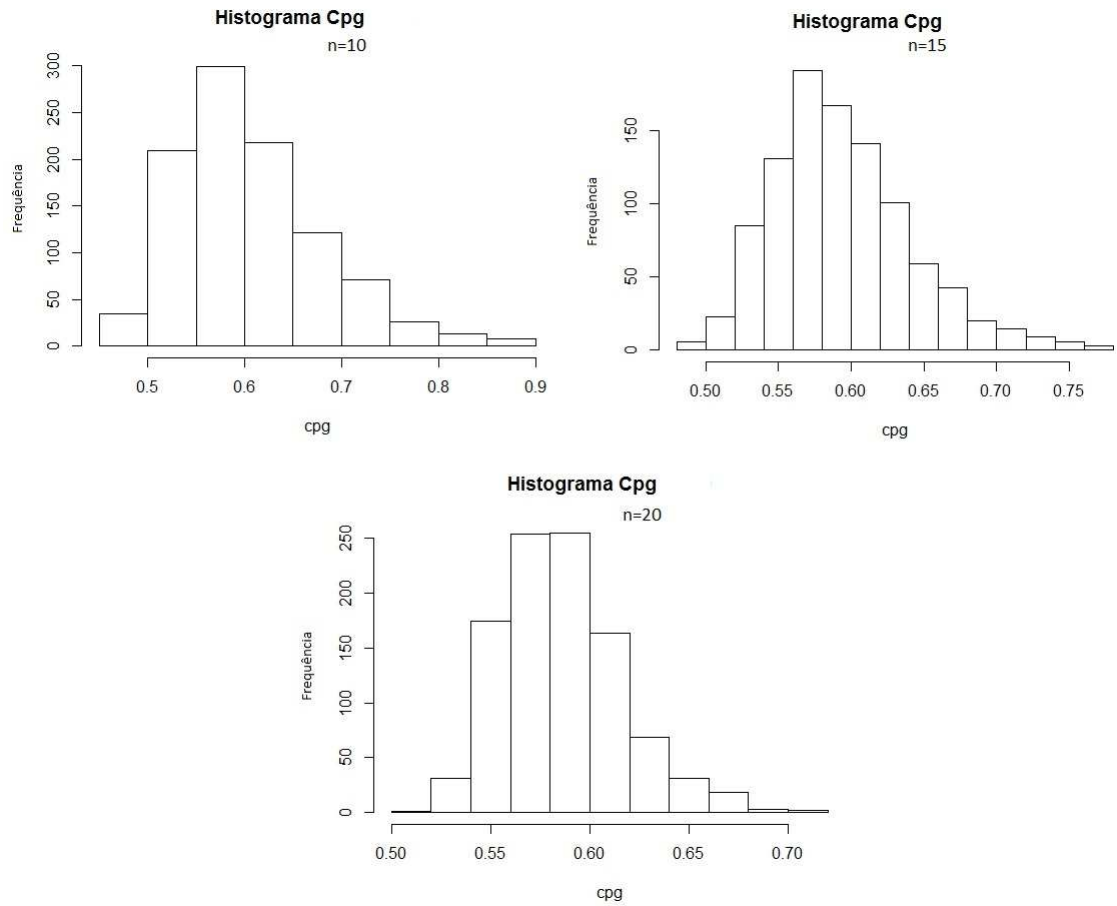


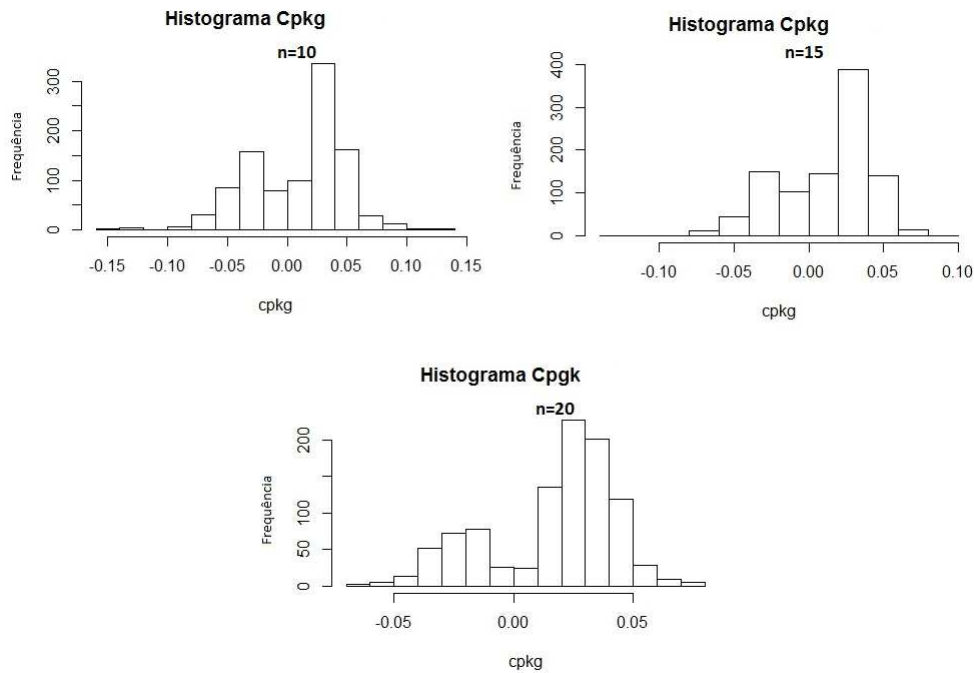
Figura 3.13: Histograma dos índices de capacidade $\hat{C}_{pkgeom}$, para a variável A .

Na Tabela 3.12, são apresentados os valores obtidos para a média dos $\hat{C}_{pkgeom}$ de acordo com o tamanho da amostra e o erro quadrático médio associado (EQM) a elas, todos obtidos com reamostragem pelo método da permutação de tamanho 1.000. O EQM é bem próximo para os diferentes tamanhos (n) de amostra. Os intervalos de confiança percentílicos nos mostra que com 95% de confiança o processo não é capaz. O histograma é apresentado na Figura 3.14.

São apresentados na Tabela 3.13 os valores obtidos para a média dos $\hat{C}_{pkgeom}$, pelo método *bootstrap*, de acordo com o tamanho da amostra e o erro quadrático médio associado (EQM) a elas, realizadas por reamostragem *bootstrap* de tamanho 1000. O EQM é bem próximo para os diferentes tamanhos (n) de amostra. Os

Tabela 3.12: Reamostragem pelo método da permutação

n	$\hat{C}_{pkgeom}$		
	Média	EQM	$IC_{percentil}$
10	0,013	0,001	(0,011 ; 0,015)
15	0,019	$4,634E^{-5}$	(0,018 ; 0,021)
20	0,024	$2,421E^{-5}$	(0,023 ; 0,025)

Figura 3.14: Histograma dos índices de capacidade $\hat{C}_{pkgeom}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , paraa variável B

intervalos de confiança percentílicos nos mostra que o processo não é capaz, com 95% de confiança.

Tabela 3.13: Reamostragem pelo método *bootstrap*

n	$\hat{C}_{pkgeom}$		
	Média	EQM	$IC_{percentil}$
10	0,010	0,002	(0,007 ; 0,012)
15	0,012	0,001	(0,010 ; 0,014)
20	0,017	0,001	(0,015 ; 0,018)

Cálculo do Índice de Capacidade C_p de Niverthi-Dey

Para todos os dados, temos o índice de capacidade $\hat{C}_{pND}$ definido abaixo:

$$\hat{C}_{pND} = \begin{bmatrix} 0,771 \\ 0,502 \\ 0,417 \end{bmatrix}.$$

Como estamos analisando três variáveis (A , B e C), o índice de capacidade $\hat{C}_{pND}$ é um vetor de dimensão 3×1 . Não há valores referência para este índice, sendo usual considerar o valor mínimo entres, chamado de índice global. Neste caso o índice global é igual 0,417, o que indica que o processo não está atendendo às especificações.

Realizamos 1.000 reamostragens pelo método de permutação para três tamanhos (n) diferentes de amostra para as variáveis A , B e C . A Figura 3.15 apresenta os histogramas que permitem visualizar as distribuições empíricas para as 1.000 reamostragens pelo método da permutação para a média dos índices $\hat{C}_{pND}$ para a variável A , a Figura 3.16, para a variável B , e a Figura 3.17, para a variável C .

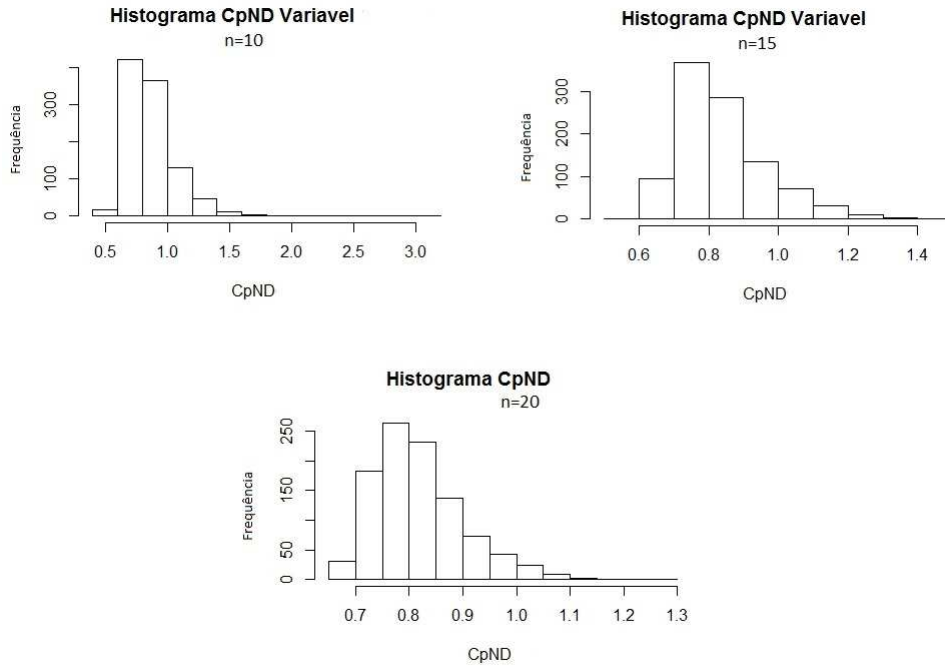


Figura 3.15: Histograma dos índices de capacidade $\hat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20

Temos na Tabela 3.14 os valores obtidos para a média dos $\hat{C}_{pND}$ de acordo com o tamanho da amostra e o erro quadrático médio associado (EQM) a elas, para reamostragens pela permutação de tamanho 1.000. O EQM não diminui muito quando aumentamos o tamanho (n) da amostra, sendo próximo os valores das médias dos $\hat{C}_{pND}$ obtidos para amostras de tamanho 15 e 20. Os intervalos de confiança percentílico nos mostram que com 95% de confiança o processo não é capaz.

Realizamos 1000 reamostragens pelo método bootstrap para três tamanhos

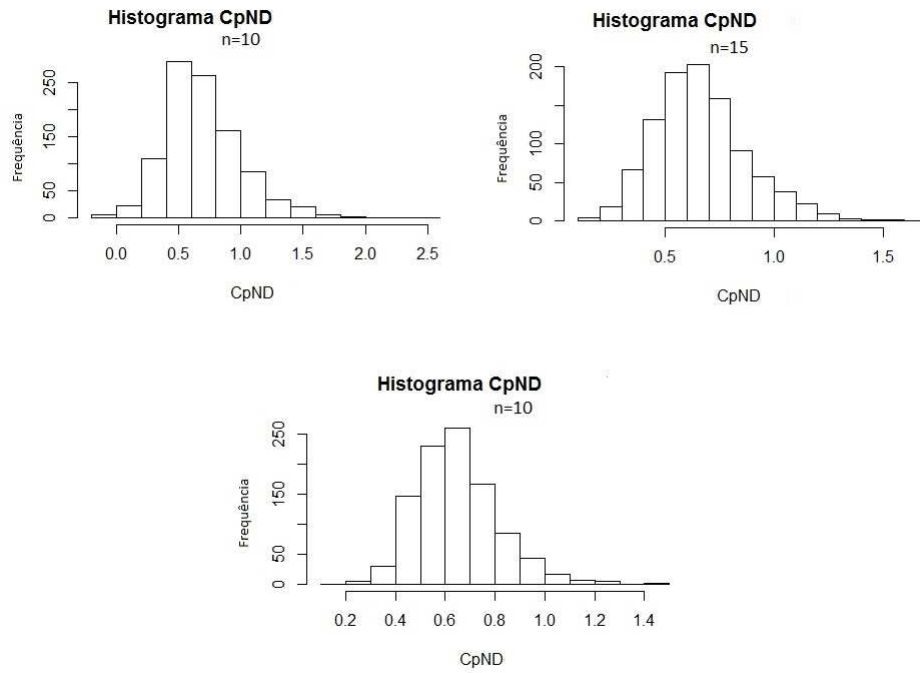


Figura 3.16: Histograma dos índices de capacidade $\hat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20

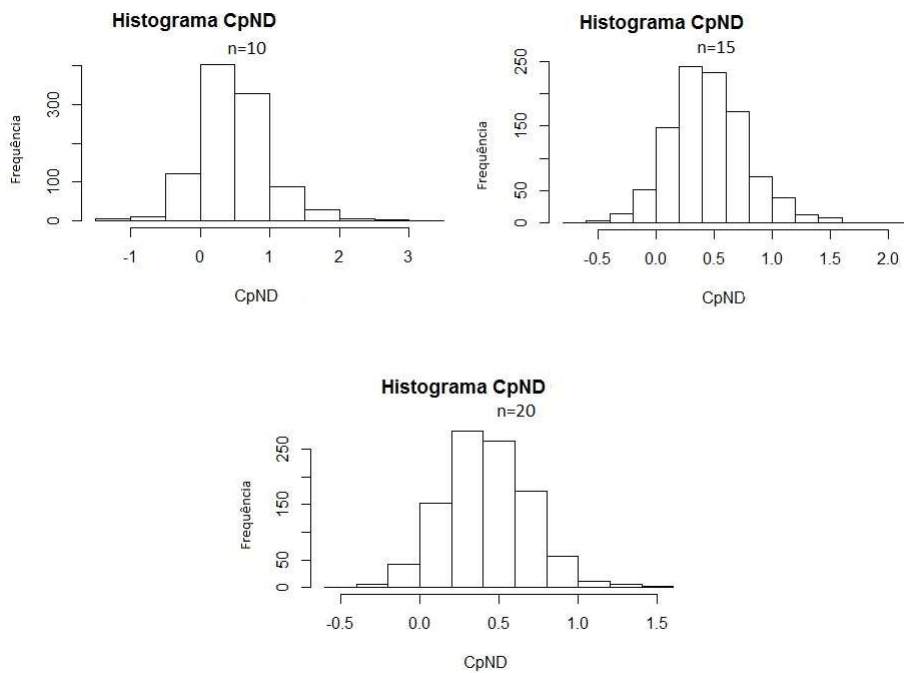


Figura 3.17: Histograma dos índices de capacidade $\hat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20

(n) diferentes de amostra para a média dos $\hat{C}_{pND}$ das variáveis A , B e C . A Figura 3.18 apresenta os histogramas que permitem visualizar as distribuições

Tabela 3.14: Reamostragem pelo método da permutação

n	$\hat{C}_{pND}$		
	Media	EQM	$IC_{percentilico}$
10	0,491	0,092	(0,458 ; 0,523)
15	0,452	0,201	(0,431 ; 0,475)
20	0,423	0,103	(0,406 ; 0,440)

empíricas para as 1.000 reamostragens pelo método *bootstrap* para a média dos índices $\hat{C}_{pND}$ para variável A , a Figura 3.19, para a variável B , e a Figura 3.20, para a variável C .

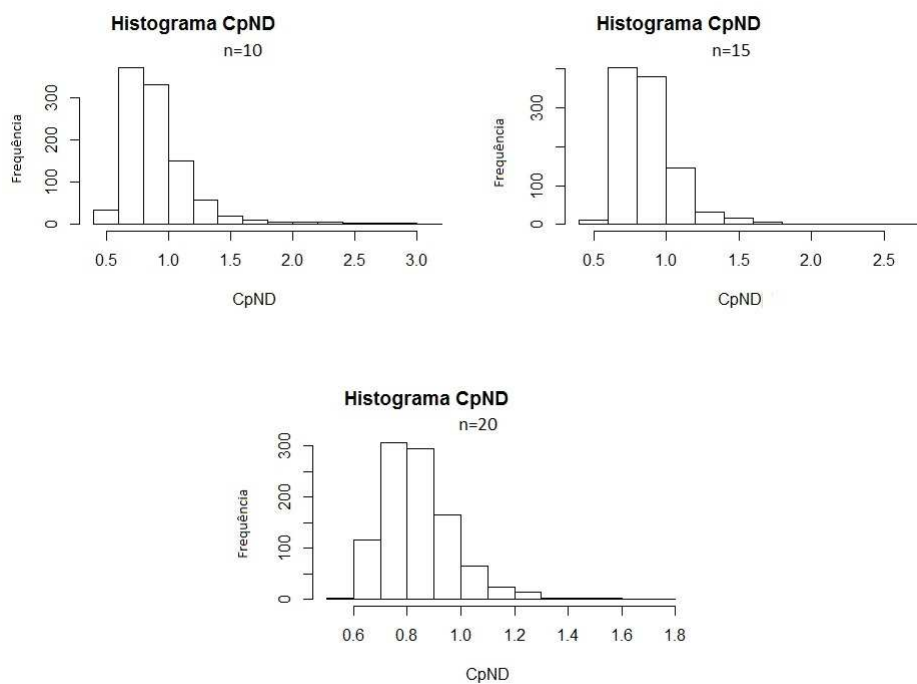


Figura 3.18: Histograma dos índices de capacidade $\hat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável A .

Temos na Tabela 3.15 os valores obtidos para a média dos $\hat{C}_{pND}$ de acordo com o tamanho da amostra e o erro quadrático médio associado (EQM) a elas. Todas realizadas com reamostragem pela permutação de tamanho 1.000. O EQM diminui não diminui muito quando aumentamos o tamanho (n) da amostra. Sendo próximo os valores das médias dos $\hat{C}_{pND}$ obtidos para todas as amostras. Os intervalos de confiança percentil nos mostram que com 95% de confiança o processo não é capaz.

Cálculo do Índice de Capacidade C_{pk} de Niverthi-Dey

Para todos os dados, temos o índice de capacidade $\hat{C}_{pkND}$ definido abaixo:

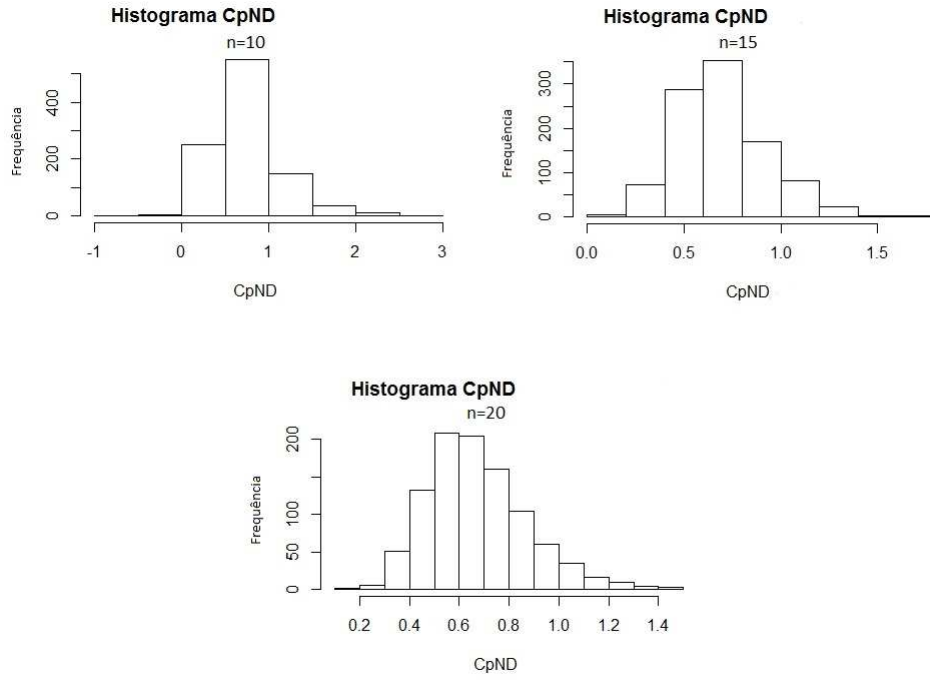


Figura 3.19: Histograma dos índices de capacidade $\hat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável B.

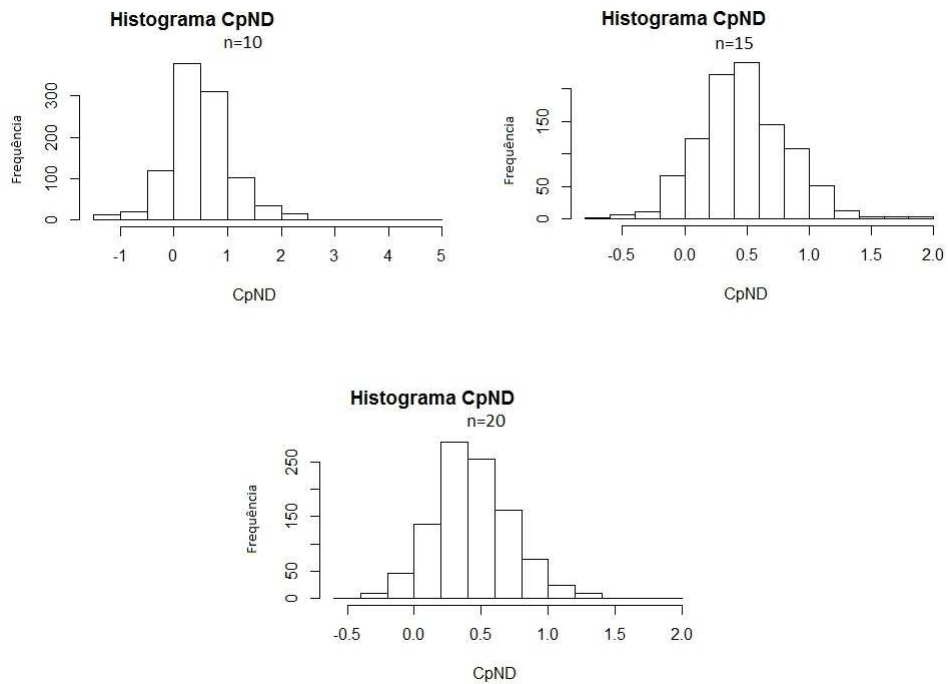


Figura 3.20: Histograma dos índices de capacidade $\hat{C}_{pND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável C.

$$\hat{C}_{pkND} = \begin{bmatrix} -0,239 \\ -0,152 \\ 0,728 \end{bmatrix}.$$

Tabela 3.15: Reamostragem pelo método *bootstrap*

n	$\hat{C}_{pND}$		
	Media	EQM	$IC_{percentilico}$
10	0,506		(0,469 ; 0,548)
15	0,473		(0,451 ; 0,497)
20	0,437		(0,419 ; 0,456)

O índice global, definido como o valor mínimo dos índices presentes no vetor, é igual a $-0,239$, o que também indica que o processo não está capaz de atender às especificações.

Realizamos 1.000 reamostragens pelo método de permutação para três tamanhos (n) diferentes de amostra para a média dos $\hat{C}_{pkND}$ das variáveis A , B e C . A Figura 3.21 apresenta os histogramas que permitem visualizar as distribuições empíricas para as 1.000 reamostragens pelo método da permutação para a média dos índices $\hat{C}_{pkND}$, para a variável A , a Figura 3.22, para a variável B , e a Figura 3.23, para a variável C .

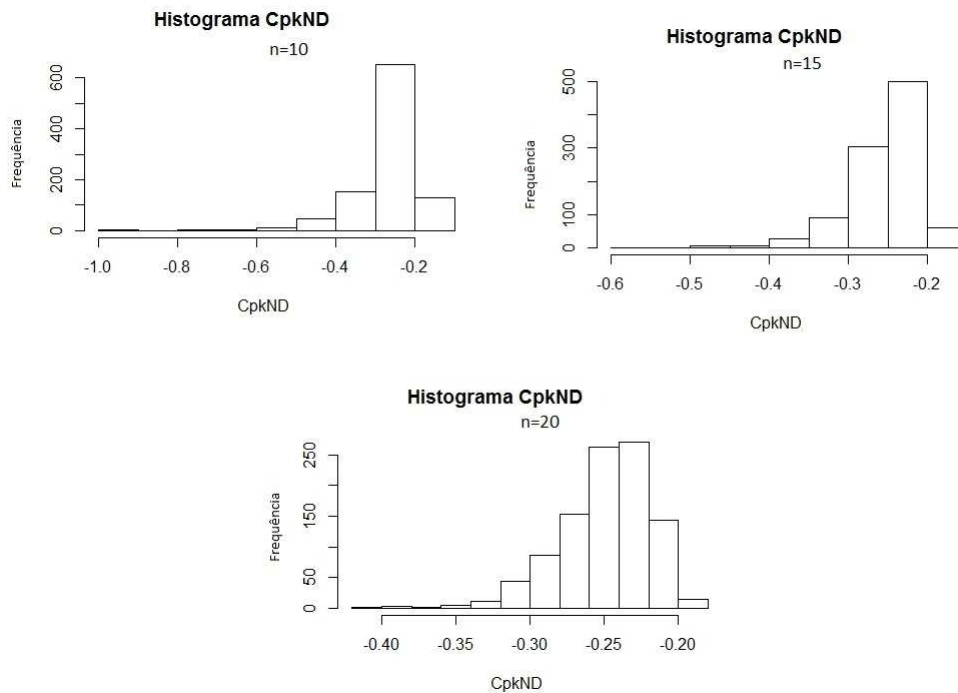


Figura 3.21: Histograma dos índices de capacidade $\hat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável A .

Temos na Tabela 3.16 os valores obtidos para a média dos $\hat{C}_{pkND}$ de acordo com o tamanho da amostra e o erro quadrático médio associado (EQM) a elas, realizadas com reamostragem pela permutação de tamanho 1.000. O EQM não diminui muito quando aumentamos o tamanho (n) da amostra e são próximos os valores das médias dos $\hat{C}_{pkND}$ obtidos para todas as amostras. Os intervalos de confiança percentílico nos mostram que com 95% de confiança o processo não é capaz.

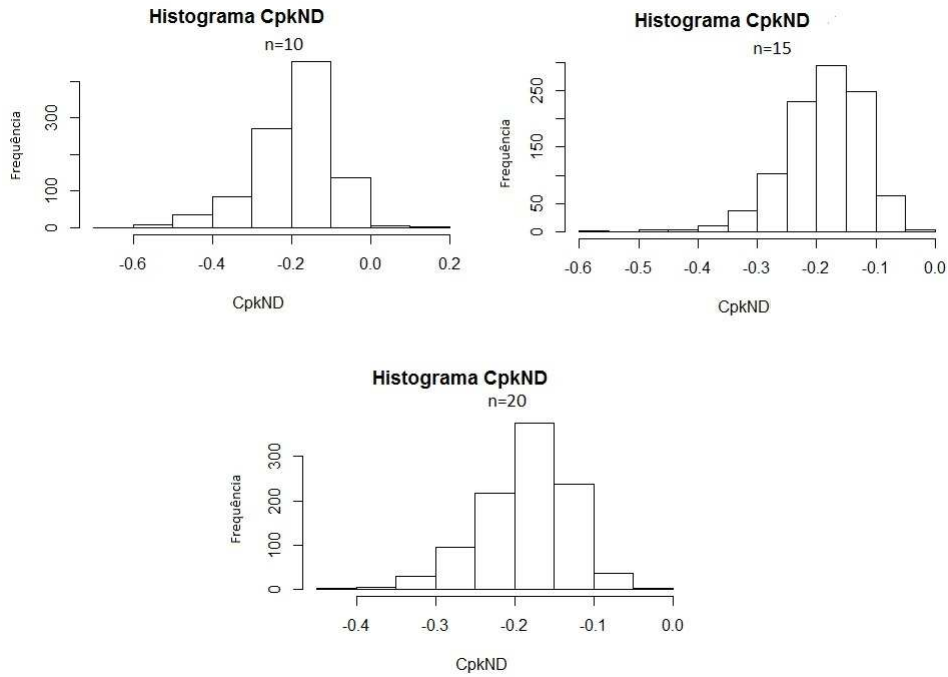


Figura 3.22: Histograma dos índices de capacidade $\hat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável B.

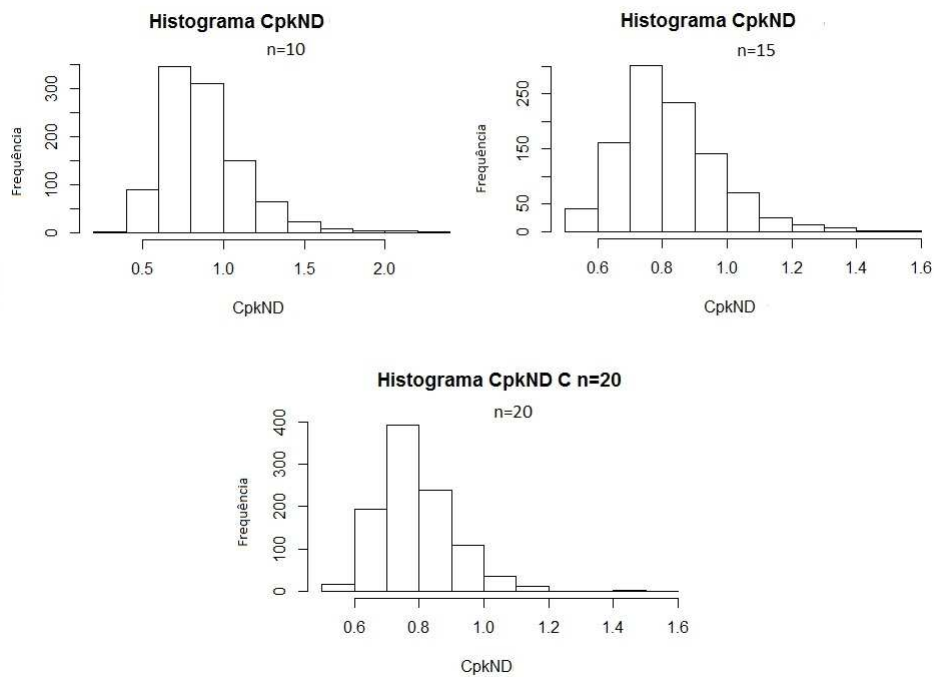
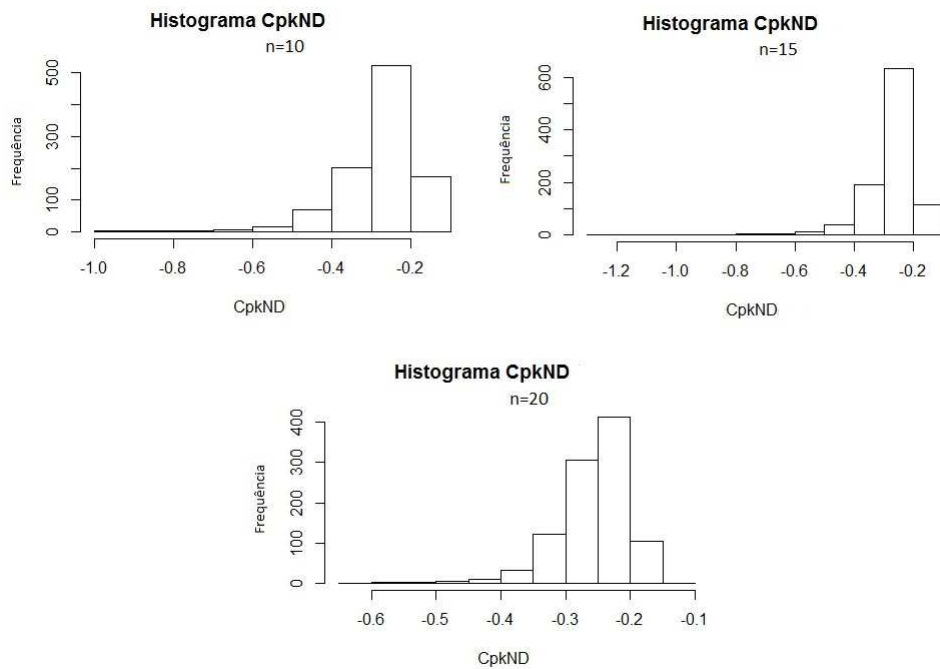


Figura 3.23: Histograma dos índices de capacidade $\hat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável C.

Realizamos também reamostragens pelo método *bootstrap* para a média dos $\hat{C}_{pkND}$ das variáveis A, B e C. A Figura 3.24 apresenta os histogramas para a variável A, a Figura 3.25, para a variável B, e a Figura 3.26, para a variável C.

Tabela 3.16: Reamostragem pelo método da permutação

n	$\hat{C}_{pkND}$		
	Media	EQM	$IC_{percentilico}$
10	-0,268	0,003	(-0,274 ; -0,263)
15	-0,254	0,002	(-0,257 ; -0,251)
20	-0,249	0,001	(-0,252 ; -0,248)

Figura 3.24: Histograma dos índices de capacidade $\hat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável A.

A Tabela 3.18 apresenta os valores obtidos para a média dos $\hat{C}_{pkND}$ e o erro quadrático médio associado (EQM), em função do tamanho da amostra. Todas foram realizadas com reamostragem pela permutação de tamanho 1.000. O EQM não diminui muito quando aumentamos o tamanho (n) da amostra e são próximos os valores das médias dos $\hat{C}_{pkND}$ obtidos para todas as amostras. Os intervalos de confiança percentílicos nos mostram que o processo não é capaz com 95% de confiança.

Tabela 3.17: Reamostragem pelo método *bootstrap*

n	$\hat{C}_{pkND}$		
	Media	EQM	$IC_{percentilico}$
10	-0,279	0,002	(-0,285 ; -0,273)
15	-0,271	0,001	(-0,276 ; -0,266)
20	-0,258	0,001	(-0,262, -0,255)

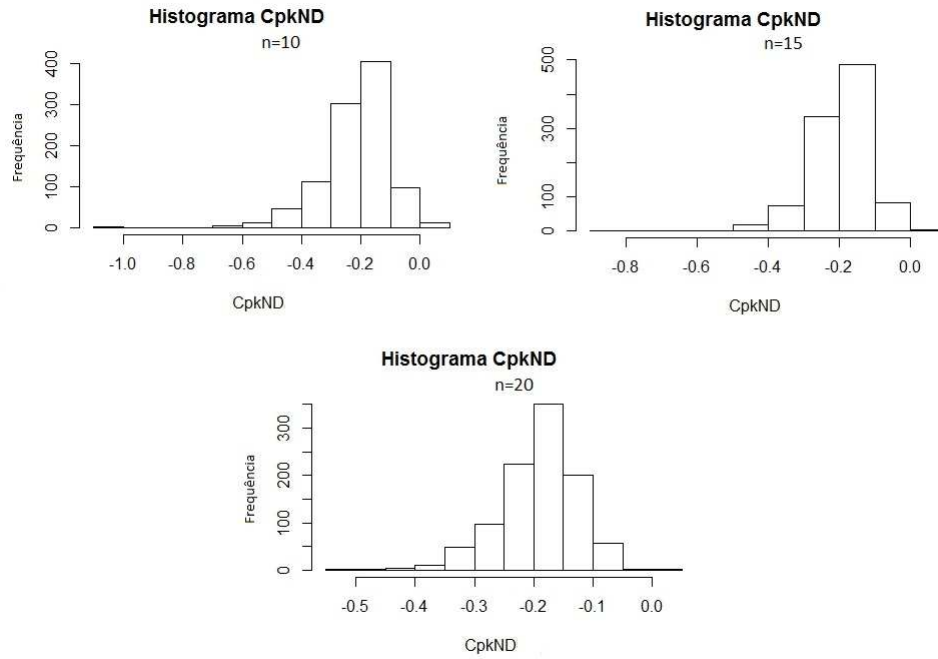


Figura 3.25: Histograma dos índices de capacidade $\hat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável B.

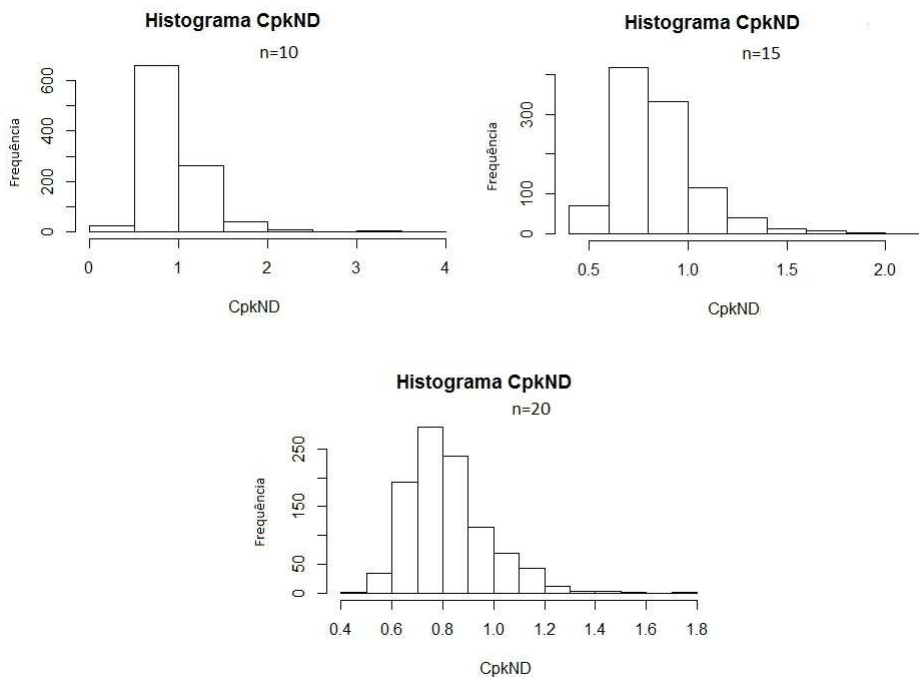


Figura 3.26: Histograma dos índices de capacidade $\hat{C}_{pkND}$ calculados para 1.000 reamostragens com $n = 10, 15$ e 20 , para a variável C.

Para as variáveis A , B e C calculamos as proporções de itens fora da especificação para 100 reamostragens de tamanho 10 e o EQM associado a elas. Com esta informação o profissional tem uma referência da % de não conformidade

com os valores estimados da capacidade neste conjunto de dados.

Tabela 3.18: Proporção de itens fora dos limites especificados

Variável	Media	EQM
A	59,2%	0,017
B	59,1%	$3,025e^{-5}$
C	64,6 %	0,015

Capítulo 4

Conclusão

4.1 Considerações finais

Após a realização das reamostragens pelos métodos da permutação e *bootstrap* podemos concluir que os dois processos não são capazes de atender às especificações estabelecidas. A retirada de alguns pontos discrepantes no banco de dados do processo de flotação da usina de tratamento de minério de ferro resulta em uma melhora do processo, porém isto não o torna capaz de atender as especificações estabelecidas. Isso indica que medidas devem ser realizadas para identificar o que está afetando o processo. Na produção de leite, a variação do tamanho de produtores pode ter ocasionado a não capacidade de atender as especificações. Observamos que quanto maior o tamanho da amostra (n), mais simétrica é a distribuição empírica, permitindo a realização de inferências. Há uma sensível redução do erro quadrático médio (EQM) quando aumentamos o tamanho (n) da amostra. Os índices de capacidade geométricos apresentaram valores muito próximos de zero, enquanto valores negativos foram obtidos pelos índices de Niverthi Dey. Dessa forma constatamos o quanto o índice de capacidade de Niverthi Dey penaliza o processo, conforme verificado em outros trabalhos. Em vista disso, sugerimos um estudo que determine valores referência para o índice de Niverthi Dey. Os intervalos de confiança percentil obtidos variam pouco para cada índice estimado. Pela natureza dos dados reais, os valores calculados dos índices podem dar uma ideia sobre a capacidade dos processos, muito embora não impliquem que serão observadas as mesmas proporções de não-conformidade que seriam esperadas para dados que seguissem uma distribuição normal. Investigações futuras deverão ser conduzidas, envolvendo estas e outras questões correlatas.

Referências Bibliográficas

- [1] BERNARDO, J. M.; IRONY, T. X. **A general multivariate bayesian process capability index.** *Statistician*, v. 45, n. 3, p. 487-502, 1996.
- [2] BULBA, E. A., HO. L. L. **Índices de capacidade de relações funcionais lineares e não-lineares**, *Revista Produção*, v. 14 n. 1 2004.
- [3] CARARI, M. L.; LIMA, P. C.; FERREIRA, D. F.; CIRILLO, M. A. **Estimação de diferenças entre duas proporções binomiais via bootstrap**, *Rev. Bras. Biom.*, São Paulo, v.28, n.3, p.112-134, 2010.
- [4] CARRASCO, C.G.; TUTIA, NAKANO, M.H.; E.Y. **Intervalos de Confiança para os Parâmetros do Modelo Geométrico em Inflação de Zeros Intervalos de Confiança para os Parâmetros do Modelo Geométrico em Inflação de Zeros**, *TEMA (São Carlos)* vol.13 no.3 São Carlos set./dez. 2012.
- [5] CHAN, L.K., CHENG, S.W.; SPIRING, F.A. **A new measure of process capability: Cpm**, *Journal of Quality Technology*, v. 20, p. 162-175, 1988.
- [6] CHEN, H. **A multivariate process capability index over a rectangular solid tolerance zone**, *Statistica Sinica*, v. 4, n. 2, p. 749-758, 1994.
- [7] CLARO, F. A. E., COSTA, A. F. B.; MACHADO, M. A. G. **Gráficos de controle de EWMA e de X-barra para monitoramento de processos autocorrelacionados**, *Prod.* vol.17 no.3 São Paulo Sept./Dec. 2007.
- [8] COSTA, A. F. B., EPPRECHT, E. K., CARPINETTI, L. C. R. **Controle Estatístico da Qualidade**, 2.ed. São Paulo: Editora Atlas, 2005.
- [9] COSTA, G. G. O. **Um procedimento inferencial para análise fatorial utilizando as técnicas bootstrap jackknife: construção de intervalos de confiança e testes de hipóteses**, Tese de Doutorado.Pontifícia Universidade Católica do Rio de Janeiro, 2006.
- [10] CUNHA, W. J., COLOSIMO, E. A. **Intervalos De Confiança Bootstrap para Modelos de Regressão com Erros de Medida**, *Rev. Mat. Estat.*, São Paulo, v.21, n.2, p.25-41, 2003.

- [11] CYMROT, R.; RIZZO, A. L. T. **Aplicação da técnica de reamostragem bootstrap na estimação da probabilidade dos alunos serem usuários de transporte público**, Santos, SP: Environmental and Health Word Congress, p. 292- 296, 2006.
- [12] DAVISON , A.C.; HINKLEY, D.V. **Bootstrap methods and their application**, New York: Ed. Cambridge, 1997.
- [13] DELERYD, M. **The effect of Skewness on estimates of some process capability indices**, International Journal of Applied Quality Management, v. 2, n. 2, p. 153-186, 1999.
- [14] DOMINGUES, K. M. **Estimação de Intervalos de Confiança via Reamostragem bootstrap usando Simulações no Software R**, (Monografia), Ouro Preto, Departamento de Estatística, 2014.
- [15] EFRON, B.; TIBSHIRANI, R. J. **An introduction to the bootstrap**, New York: Chapman & Hall , p. 436, 1993.
- [16] EFRON, B.; TIBSHIRANI, R. J. **Bootstrap methods for standard errors, confidence intervals, and other measures of statistical accuracy**, Statistical Science , v. 1, n. 1, p. 55-77, 1986.
- [17] FERREIRA, D. F. **Estatística Multivariada**, 2 ed. Editora ULFA, Lavras, 2011.
- [18] FERREIRA, D.F. **Estatística computacional em Java**, 1 ed., Editora UFLA, Lavras, 2013.
- [19] FILHO, M. J. F.; LIBARDI, P. L.; LIER, Q. J.; CORRENTE J. E. **Método convencional e bootstrap para estimar o número de observações na determinação dos parâmetros da função k**, Rev. Bras. Ciênc. Solo v. 26, n. 4, Viçosa out/dez 2002.
- [20] GENTLE, J.E. **Computational Statistics**, New York: Ed. Springer, 2009.
- [21] GONÇALEZ, P. U., WERNER, L. **Comparação dos índices de capacidade do processo para distribuições não-normais**, Gestão e Produção, v. 16, n. 1, p. 121-132, jan.-mar. 2009.
- [22] HAYTER, A. J.; TSUI, K-L. **Identification and quantification in multivariate quality control problems**, Journal of Quality Technology, v. 26, n. 3, p. 197-208, 1994.
- [23] HUBELE, N. F.; BERRADO, A.; GEL, E. S. . **A Wald Test for Comparing Multiple Capability Indices**, Journal of Quality Technology 37, p. 304 - 307, 2005.
- [24] JUNIOR, J. I. R. **Métodos estatísticos aplicados ao controle da qualidade**, Viçosa, MG: Ed. UFV, 2013.

- [25] JURAN, J. M. **Quality Control Handbook**, New York, Mc Graw-Hill, 1974.
- [26] KANE, V.E. **Process capability indices**, Journal of Quality Technology, v. 18, n. 1, p. 41-52, 1986.
- [27] LI, Y., LIN, C. **Multivariate Cp value**, Chinese Journal of applied Probability and Statistics, n.12, p. 132-138, 1996.
- [28] MINGOTI, S. A.; CONCEIÇÃO, M. M. C. **Uma proposta do índice Cpm multivariado e implementação computacional no software minitab for Windows dos índices de capacidade multivariados de Chen modificados, propostos por Mingoti e Glória.** (Dissertação de Mestrado), Belo Horizonte, Departamento de estatística da UFMG, 2004.
- [29] MINGOTI, S. A.; GLÓRIA, F. A. A. **Comparing mingoti and glória's and niverthi and dey's multivariate capability indexes**, Produção, v. 18, n. 3, p. 598-608, 2008.
- [30] MINGOTI, S. A.; OLIVEIRA, F. L. P.; CONCEIÇÃO, M. M. C. **Índices de capacidade para processos multivariados independentes: extensões dos índices de Niverthi e Dey e Mingoti e Glória**, Revista Produção, v. 21, n. 1, p. 94-105, jan./mar. 2011.
- [31] MINGOTI, S. A.; OLIVEIRA, F. L. P. **On Capability Indices for Multivariate Autocorrelated Process**, Brazilian Journal of Operations & Production Management, v. 8, n. 1,p. 133-152, 2011.
- [32] Mohamadi, M.; FOUMANI, M.; ABBASI, B. **Process Capability Analysis in the Presence of Autocorrelation**, Journal of Optimization in Industrial Engineering, p.15-20, agosto 2011.
- [33] MONTGOMERY, D. C. **Introdução ao Controle Estatístico da Qualidade**, São Paulo: Ed. LTC, 2004.
- [34] MONTGOMERY, D. C.; RUNGER, G. C. **Estatística Aplicada e Probabilidade para Engenheiros**, Rio de Janeiro: Ed. LTC, 2003.
- [35] NIVERTHI, M.; DEY, D. K. **Multivariate process capability: a bayesian perspective. Communications in Statistics-Simulation and Computation**, v. 29, n. 2, p. 667-687, 2000.
- [36] OLIVEIRA, F. L. P. **Índices de capacidade para processos multivariados autocorrelacionados**, (Dissertação de Mestrado), Belo Horizonte, Departamento de estatística da UFMG, 2007.
- [37] PAESE, C., CATEN, C. T., RIBEIRO, J. L. D. **Aplicação da análise de variância na implantação do CEP** , Revista Produção, v. 11 n. 1, Novembro de 2001.
- [38] PAN. J. N.; HUANG, W. K. C. **Developing New Multivariate ProcessCapability Indices for Autocorrelated Data**, Quality and reliability engineering international, 2013.

- [39] PAN, J. N.; LEE, C. Y. **New capability indices for evaluating the performance of multivariate manufacturing processes**, Quality and Reliability Engineering International, v. 26, n. 1, p. 3-15, 2010.
- [40] PAN, J.; LI, C. **New capability indices for measuring the performance of a multidimensional machining process**, Expert Systems with Applications, v. 41, p. 2409-2414, 2014.
- [41] PEARN, W.L., KOTZ, S.; JOHNSON, N.L. **Distributional and inferential properties of process control indices**, Journal of Quality Technology, v. 24, n. 4, 216-231, 1992.
- [42] POLANSKY, A. M. **Permutation Methods for Comparing Process Capabilities**, Journal of Quality Technology, v. 38, n. 3, Julho 2006.
- [43] PONTES, A. C. F., CORRENTE, J. E. **Comparações Múltiplas não-paramétricas para o delineamento com um fator de classificação simples**, Rev. Mat. Estat., São Paulo, v. 19, p. 179-197, 2001.
- [44] RAMOS, A.W., HO, L.L. **Procedimentos inferenciais em índices de capacidade para dados autocorrelacionados via bootstrap**, Revista Produção, v.13, n.3, p. 50-62, 2003.
- [45] RESENDE, M. D. V., DUDA, L.L., GUIMARÃES, P. R. B., FERNANDES, J. S. C. **Análise de Modelos Lineares Mistos via Inferência Bayesiana**, Rev. Mat. Estat., São Paulo, v. 19, p. 41-70, 2001.
- [46] SANTOS, A. G., LACERDA, E. F., NETO, H. C. A., LUNA, W. A., FURLANETTO, E. L. **A importância dos gráficos de controle para monitorar a qualidade dos processos Industriais: estudo de caso numa indústria metalúrgica**, Cadernos do IME- Série Estatística, v. 28, p. 33-46, 2010.
- [47] SILVEIRA, F. V. J. **Testes de permutação e bootstrap em análise estatística de formas: aplicações à zoologia**, (Dissertação de Mestrado), Recife, Departamento de Estatística, 2008.
- [48] SOUZA, F. S. **Índices de Capacidade para Gráficos de Controle baseados em Modelos de Regressão**, (Dissertação de Mestrado), Porto Alegre, Departamento de Estatística, 2010.
- [49] Taam, W., Subbaiah, P., Liddy, J. W. **A note on multivariate capability indices**, Journal of Applied Statistics, v.20, n.3, p. 339-351, 1993.
- [50] VEEVERS, A. **Viability and capability indexes for multiresponse processes**, Journal of Applied Statistics, v. 25, n. 4, p. 545-558, 1998.
- [51] WANG, F.K., HUBELE, N.F., LAWRENCE, F.P., MISKULIN, J.D. and SHAHRIARI, H. **Comparison of Three Multivariate Process Capability Indices**, Journal of Quality Technology, v.32 n.3, p. 263-275, 2000.

- [52] WOODALL, W. H., MONTGOMERY, D.C. **Some Current Directions in the Theory and Application of Statistical Process Monitoring**, Journal of Quality Technology, v. 46, n. 1, Janeiro 2014.
- [53] YEH, A. B., CHEN, H. **A Nonparametric Multivariate Process Capability Index**, 1999.