



UNIVERSIDADE FEDERAL DE PERNAMBUCO
CENTRO DE TECNOLOGIA E GEOCIÊNCIAS
DEPARTAMENTO DE ENGENHARIA CIVIL
PROGRAMA DE PÓS-GRADUAÇÃO EM ENGENHARIA CIVIL

RODRIGO SANTOS DE OLIVEIRA

**ANÁLISE DE CONFIABILIDADE ESTRUTURAL USANDO O MÉTODO MONTE
CARLO SELETIVO**

Recife
2022

RODRIGO SANTOS DE OLIVEIRA

**ANÁLISE DE CONFIABILIDADE ESTRUTURAL USANDO O MÉTODO MONTE
CARLO SELETIVO**

Dissertação apresentada ao Programa de Pós-Graduação em Engenharia Civil da Universidade Federal de Pernambuco, como requisito parcial para obtenção do título de mestre em Engenharia Civil.

Área de concentração: Estruturas.

Orientador: Prof. Dr. Renato de Siqueira Motta.

Coorientadora: Profa. Dra. Silvana Maria Bastos Afonso.

Recife

2022

Catálogo na fonte
Sandra Maria Neri Santiago, CRB-4 / 1267

- O48a Oliveira, Rodrigo Santos de.
 Análise de confiabilidade estrutural usando o Método Monte Carlo Seletivo /
Rodrigo Santos de Oliveira. – 2022.
 81 f.: figs., tabs.
- Orientador: Prof. Dr. Renato de Siqueira Motta.
 Coorientadora: Profa. Dra. Silvana Maria Bastos Afonso.
 Dissertação (Mestrado) – Universidade Federal de Pernambuco. CTG.
Programa de Pós-Graduação em Engenharia Civil, Recife, 2022.
 Inclui referências.
1. Engenharia civil. 2. Monte Carlo. 3. Análise de confiabilidade. 4. Monte
Carlo Seletivo. 5. Análise estrutural. I. Motta, Renato de Siqueira (Orientador).
II. Afonso, Silvana Maria Bastos (Coorientadora). III. Título.
- UFPE
- 624 CDD (22. ed.) BCTG / 2022-45

RODRIGO SANTOS DE OLIVEIRA

**ANÁLISE DE CONFIABILIDADE ESTRUTURAL USANDO O MÉTODO MONTE
CARLO SELETIVO**

Dissertação apresentada ao Programa de Pós-Graduação em Engenharia Civil da Universidade Federal de Pernambuco, Centro de Tecnologia e Geociências, como requisito parcial para a obtenção do título de Mestre em Engenharia Civil. Área de concentração: Estruturas.

Aprovada em: 04/02/2022.

BANCA EXAMINADORA

Participação por videoconferência
Prof. Dr. Renato de Siqueira Motta (Orientador)
Universidade Federal de Pernambuco

Participação por videoconferência
Prof. Dr. Ramiro Brito Willmersdorf (Examinador Interno)
Universidade Federal de Pernambuco

Participação por videoconferência
Prof. Dr. André Jacomel Torii (Examinador Externo)
Universidade Federal da Integração Latino-Americana

Dedico este trabalho ao meu avô, Elysio, que foi um exemplo na minha vida e sempre acreditou que eu alcançaria tudo que meu coração desejasse.

AGRADECIMENTOS

Aos meus pais, Agildo e Valéria, por sempre me darem apoio e amor para perseguir meus sonhos e evoluir como ser humano.

À minha noiva, Mylena, pela motivação incondicional, compreensão e paciência quando necessário.

Aos meus orientadores, professor Renato e professora Silvana, pela dedicação a este trabalho e por representarem um norte para mim como pesquisador.

Aos meus irmãos, Bruno e André, por ajudarem a me manter firme durante o percurso.

Aos meus avós, Elysio e Dapaz, pela fé inabalável e os ensinamentos transmitidos ao longo da minha vida.

Aos meus colegas do mestrado, por estarem sempre disponíveis para me auxiliar.

O presente trabalho foi realizado com apoio da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior - Brasil (CAPES) - Código de Financiamento 001.

Os apoios financeiros do CNPq e da FACEPE também são reconhecidos pelo autor.

“The most important questions of life are, for the most part, really only problems of probability”
(LAPLACE, 1812).

RESUMO

O método Monte Carlo (MC) tem como principais vantagens a robustez e a simplicidade. Em problemas práticos de confiabilidade estrutural, entretanto, o tamanho da amostra pode tornar o MC inviável devido ao alto custo computacional para o cálculo da probabilidade de falha. Diversos métodos baseados no MC foram desenvolvidos buscando diminuir esse problema, como o *Importance Sampling* (IS), o *Subset Simulation* (SuS) e o *Separable Monte Carlo* (SepMC). Neste trabalho, o método Monte Carlo Seletivo (SMC) é proposto, visando problemas cuja função de falha é localmente monotônica na região de interesse, o que acontece na maior parte dos problemas de engenharia estrutural. O SMC reduz o número de avaliações da função de falha (*NFE*) do MC ao calcular apenas uma parte dos pontos, realizando uma busca pelos pontos falhos no sentido de falha de cada variável. Dois problemas de estruturas e 13 exemplos de referência são analisados, onde o *NFE* para um dado coeficiente de variação da probabilidade de falha é contabilizado e comparado com o MC, IS, SuS e SepMC. Além disso, um estudo paramétrico para avaliar o impacto do número de variáveis da função de falha no método proposto é realizado. Os resultados mostraram que a redução no *NFE* pelo SMC pode chegar a mais de 99,9% em relação ao MC e foi maior que a do SuS para todos os problemas, além de maior que a do IS e do SepMC para a maioria dos problemas. Vale salientar que quando a função de falha é monotônica, para uma mesma amostra, o SMC e o MC resultam em um mesmo valor de probabilidade de falha.

Palavras-chave: Monte Carlo; análise de confiabilidade; Monte Carlo Seletivo; análise estrutural.

ABSTRACT

The Monte Carlo method (MC) has as main advantages the robustness and simplicity. In practical problems of structural reliability, however, the sample size can make the MC unfeasible due to the high computational cost for calculating the failure probability. Several methods based on MC have been developed to reduce this issue, such as Importance Sampling (IS), Subset Simulation (SuS) and Separable Monte Carlo (SepMC). In this work, the Selective Monte Carlo method (SMC) is proposed, aiming at problems whose failure function is locally monotonic in the region of interest, which happens in most structural engineering problems. The SMC reduces the number of evaluations of the failure function (*NFE*) of the MC by calculating only a part of the points, performing a search for the failed points in the direction of failure of each variable. Two structural problems and 13 benchmark examples are analyzed, where the *NFE* for a given coefficient of variation of the failure probability is accounted for and compared with the MC, IS, SuS and SepMC. In addition, a parametric study to evaluate the impact of the number of variables of the failure function on the proposed method is carried out. The results showed that the reduction of the *NFE* by SMC can reach more than 99.9% in relation to MC and was greater than that of SuS for all problems, in addition to greater than that of IS and SepMC for most problems. It is worth noting that when the failure function is monotonic, for the same sample, the SMC and MC result in the same failure probability value.

Keywords: Monte Carlo; reliability analysis; Selective Monte Carlo; structural analysis.

LISTA DE FIGURAS

Figura 1 –	Análise da estrutura de forma (a) determinística e (b) probabilística	19
Figura 2 –	Aproximação linear da função de falha no espaço padrão reduzido	27
Figura 3 –	Transformação de $\mathbf{X}$ para $\mathbf{U}$, onde $\mathbf{X}$ são variáveis possivelmente correlacionadas no espaço original e $\mathbf{U}$ sem correlação no espaço isoprobabilístico	30
Figura 4 –	Exemplo do funcionamento do HL-RF	31
Figura 5 –	IS com a função de amostragem centrada no MPP	37
Figura 6 –	Comparação da amostra gerada pelo (a) MC tradicional com a gerada pelo (b) SepMC	39
Figura 7 –	Exemplo de amostra gerada pelo SuS	46
Figura 8 –	Exemplo de fronteira de Pareto	49
Figura 9 –	Primeira iteração do SMC	53
Figura 10 –	Iterações do SMC: (a) iteração 2, (b) iteração 3 e (c) iteração 4	54
Figura 11 –	Problemas 2 e 11 presentes em Santos et al. (2012)	55
Figura 12 –	Número de pontos avaliados em relação ao número de variáveis para uma amostra com $N = 10^6$	58
Figura 13 –	Problema 21 presente em Santos et al. (2012)	64
Figura 14 –	Viga biapoiada do problema 1	65
Figura 15 –	$CoV(\tilde{P}_f)$ versus o número de avaliações da função de falha calculado pelo SMC em comparação com o MC, IS, SuS e SepMC, onde os pontos no topo têm $CoV(\tilde{P}_f) = \infty$	66
Figura 16 –	Convergência da $\tilde{P}_f$ para cada método	67
Figura 17 –	$CoV(\tilde{P}_f)$ do IS à medida que a simulação avança	68
Figura 18 –	CDF de G estimada pelo SuS	68
Figura 19 –	Análise da taxa de pontos avaliados pelo SMC em relação a N	69
Figura 20 –	Modos de falha do pórtico simples do problema 2	70

LISTA DE TABELAS

Tabela 1 –	Métodos baseados em simulação MC encontrados na literatura recente	23
Tabela 2 –	Resultados de cada iteração do SMC para $N = 10^4$	53
Tabela 3 –	Resultados da amostra do exemplo	56
Tabela 4 –	Problemas de referência	59
Tabela 5 –	Resultados dos problemas de referência para $CoV(\tilde{P}_f) = 5\%$	60
Tabela 6 –	Parâmetros estatísticos das variáveis do problema 1	65
Tabela 7 –	Análise dos resultados do problema da viga biapoiada	66
Tabela 8 –	Parâmetros das distribuições das variáveis do problema 2	70
Tabela 9 –	Resultados de probabilidade de falha e avaliações de função do problema 2	71

LISTA DE ABREVIATURAS E SIGLAS

ADIS	<i>Cross-entropy based Adaptive Directional Importance Sampling</i>
AK-IS	<i>Adaptive Kriging-Importance Sampling</i>
AK-MCS	<i>Adaptive Kriging-Monte Carlo</i>
ASVM-MCS	<i>Adaptive Support Vector Machine-Monte Carlo</i>
AWL-MCS	<i>Active Weight Learning-Monte Carlo</i>
ACS	<i>Adaptive Conditional Sampling</i>
CDF	Função de probabilidade cumulativa
CHL-RF	<i>Conjugate Hasofer-Lind-Rackwitz-Fiessler</i>
CS	<i>Conditional Sampling</i>
DGPR	<i>Dynamic Gaussian Process Regression</i>
FEM	Método dos Elementos Finitos
FORM	Método de Confiabilidade de Primeira Ordem
GSS	<i>Generalized Subset Simulation</i>
HL-RF	<i>Hasofer-Lind-Rackwitz-Fiessler</i>
iHL-RF	<i>Improved Hasofer-Lind-Rackwitz-Fiessler</i>
IHS	<i>Improved Latin Hypercube</i>
IHS-EIS	<i>Improved Latin Hypercube-Effective Importance Sampling</i>
i.i.d.	Independentes e identicamente distribuídos
IS	<i>Importance Sampling</i>
MC	Monte Carlo
MCMC	<i>Markov Chain Monte Carlo</i>
MCS	Monte Carlo
MFD	<i>Modified Feasible Direction</i>
MH	Metropolis-Hastings
M-IS	<i>Multisphere-based Importance Sampling</i>
MPP	Ponto mais provável de falha
nHL-RF	<i>New Hasofer-Lind-Rackwitz-Fiessler</i>
PDF	Função de densidade de probabilidades
RBDO	Otimização de Projeto Baseada em Confiabilidade
RBIS	<i>Radial-based Importance Sampling</i>
RSM	<i>Response Surface Method</i>

SepMC	<i>Separable Monte Carlo</i>
SMC	Monte Carlo Seletivo
SORM	Método de Confiabilidade de Segunda Ordem
SQP	<i>Sequential Quadratic Programming</i>
SuS	<i>Subset Simulation</i>
VA	Variável aleatória

LISTA DE SÍMBOLOS

P_f	Probabilidade de falha
NFE	Número de avaliações da função de falha
$g(\mathbf{X})$	Função de falha
$\mathbf{X}$	Conjunto de variáveis aleatórias
D_f	Domínio de falha
$f_{\mathbf{X}}(\mathbf{x})$	PDF conjunta de $\mathbf{X}$
$\mathbf{U}$	Conjunto de variáveis aleatórias no espaço normal padrão
β	Índice de Confiabilidade
Φ	CDF normal padrão unidimensional
β_c	Índice de Confiabilidade de Cornell
G	Função de falha na forma de variável aleatória
μ_G	Média de G
σ_G	Desvio-padrão de G
$f_{\mathbf{U}}(\mathbf{u})$	PDF conjunta de $\mathbf{U}$
$g(\mathbf{U})$	Função de falha no espaço normal padrão
$\mathbf{u}^*$	Ponto mais provável de falha no espaço normal padrão
$\tilde{g}(\mathbf{U})$	Aproximação linear da função de falha no espaço normal padrão
$\mathbf{u}_p$	Ponto sobre o qual a série da aproximação $\tilde{g}(\mathbf{u})$ é centrada
$\tilde{P}_f$	Estimativa da probabilidade de falha
J_{ux}	Matriz jacobiana de $\mathbf{u}$ em relação a $\mathbf{x}$
J_{xu}	Matriz jacobiana de $\mathbf{x}$ em relação a $\mathbf{u}$
$\boldsymbol{\mu}$	Conjunto das médias das variáveis aleatórias
ρ	Coeficiente de correlação
$\mathbf{Z}$	Conjunto de variáveis aleatórias correlacionadas no espaço normal padrão
$\rho_{X_{ij}}$	Coeficiente de correlação entre as variáveis X_i e X_j
F_{X_i}	CDF da variável X_i
$F_{X_i}^{neq}$	CDF da normal equivalente da variável X_i
f_{X_i}	PDF da variável X_i
$f_{X_i}^{neq}$	PDF da normal equivalente da variável X_i
φ	PDF normal padrão unidimensional

$\sigma_{X_i}^{neq}$	Desvio-padrão da normal equivalente da variável X_i
$\mu_{X_i}^{neq}$	Média da normal equivalente da variável X_i
J_{zx}	Matriz jacobiana de de $\mathbf{z}$ em relação a $\mathbf{x}$
J_{xz}	Matriz jacobiana de de $\mathbf{x}$ em relação a $\mathbf{z}$
J_{uz}	Matriz jacobiana de de $\mathbf{u}$ em relação a $\mathbf{z}$
J_{zu}	Matriz jacobiana de de $\mathbf{z}$ em relação a $\mathbf{u}$
N	Tamanho amostral ou número de pontos amostrais da variável mais custosa
I	Função indicadora
NFP	Número de pontos falhos da amostra
p	Probabilidade de uma variável de Bernoulli ser igual a 1
E	Esperança
Var	Variância
CoV	Coeficiente de Variação
β_{MC}	Índice de Confiabilidade do Monte Carlo
$h_{\mathbf{x}}(\mathbf{x})$	Função de amostragem do <i>Importance Sampling</i>
E_h	Esperança com relação a $h_{\mathbf{x}}(\mathbf{x})$
D	Domínio do conjunto de variáveis $\mathbf{X}$
w	Pesos do <i>Importance Sampling</i>
Var_h	Variância com relação a $h_{\mathbf{x}}(\mathbf{x})$
φ_n	PDF normal padrão n -dimensional
$h_{\mathbf{u}}(\mathbf{u})$	Função de amostragem do <i>Importance Sampling</i> no espaço isoprobabilístico
M	Número de pontos amostrais da variável menos custosa
C	Parte da função de falha das variáveis de resistência
R	Parte da função de falha das variáveis de solicitação
$\hat{F}_C$	CDF experimental de C
b	Nível limite de falha para o <i>Subset Simulation</i>
b_i	Nível limite i para o <i>Subset Simulation</i>
F_i	Evento de falha i para o <i>Subset Simulation</i>
p_0	Probabilidade de falha fixada para cada nível do <i>Subset Simulation</i>
$f_{\pi}(\mathbf{X})$	PDF alvo do <i>Markov Chain Monte Carlo</i>
$q(\mathbf{x})$	Função proporcional à PDF alvo do <i>Markov Chain Monte Carlo</i>
$p^*(\mathbf{X}; \mathbf{v})$	PDF proposta para o <i>Markov Chain Monte Carlo</i>

r	Taxa de variação ponderada da PDF alvo para o <i>Markov Chain Monte Carlo</i>
ν	Ponto no qual a PDF proposta do <i>Markov Chain Monte Carlo</i> foi centrada
$\mathcal{F}$	Fronteira de Pareto
$\mathcal{S}$	Amostra na forma de uma matriz $N \times n$
$\mathcal{S}_p$	Conjunto de pontos dominantes da amostra
$\mathcal{S}_{ps}$	Conjunto de pontos dominantes seguros da amostra
$\mathcal{S}_s$	Conjunto de pontos seguros da amostra
$p_{dominant}$	Lista de pontos dominantes da amostra
$p_{dominated}$	Lista de pontos dominados da amostra
q	Carga distribuída no problema 1
W	Carga concentrada no problema 1
M_r	Momento resistente no problema 1

SUMÁRIO

1	INTRODUÇÃO	18
1.1	CONSIDERAÇÕES INICIAIS	18
1.2	OBJETIVOS	21
1.3	REFERENCIAL TEÓRICO	21
1.4	ORGANIZAÇÃO DA DISSERTAÇÃO	24
2	MÉTODOS DE CONFIABILIDADE ESTRUTURAL	25
2.1	MÉTODO DE CONFIABILIDADE DE PRIMEIRA ORDEM (FORM)	25
2.1.1	Descrição do método	27
2.1.2	Transformação para o espaço isoprobabilístico	28
2.1.3	Busca pelo MPP	31
2.1.4	Algoritmo do FORM	32
2.2	MÉTODOS DE SIMULAÇÃO	32
2.2.1	Método Monte Carlo (MC)	33
2.2.1.1	Características do estimador da probabilidade de falha do MC	33
2.2.1.2	Índice de Confiabilidade do MC	34
2.2.1.3	Algoritmo do MC	35
2.2.2	Importance Sampling (IS)	35
2.2.2.1	Características do estimador da probabilidade de falha do IS	36
2.2.2.2	Função de amostragem usando o MPP	37
2.2.2.3	Algoritmo do IS	38
2.2.3	Separable Monte Carlo (SepMC)	38
2.2.3.1	Características do estimador da probabilidade de falha do SepMC	40
2.2.3.2	Algoritmo do SepMC	40
2.2.4	Subset Simulation (SuS)	41
2.2.4.1	Markov Chain Monte Carlo (MCMC)	42
2.2.4.2	Algoritmo do SuS	45
3	MÉTODO PROPOSTO: MONTE CARLO SELETIVO (SMC)	47
3.1	DEFINIÇÃO DO PROBLEMA	47
3.2	OTIMALIDADE DE PARETO	48
3.3	CONCEITO DO MÉTODO	49

3.3.1	Algoritmo do SMC	50
3.3.2	Exemplo do método	52
3.4	LIMITAÇÕES	54
3.5	VALIDAÇÃO E COMPARAÇÃO COM OUTROS MÉTODOS	55
4	APLICAÇÕES	57
4.1	ESTUDO PARAMÉTRICO DO NÚMERO DE VARIÁVEIS	57
4.2	PROBLEMAS DE REFERÊNCIA	59
4.3	PROBLEMAS DE ESTRUTURAS	64
4.3.1	Problema 1: Viga biapoiada	64
4.3.2	Problema 2: Pórtico biengastado	69
5	CONCLUSÕES E SUGESTÕES	72
5.1	CONCLUSÕES DOS RESULTADOS OBTIDOS	72
5.2	SUGESTÕES PARA ESTUDOS FUTUROS	73
	REFERÊNCIAS	75

1 INTRODUÇÃO

Este trabalho consiste na dissertação de mestrado do autor, apresentada ao Programa de Pós-Graduação em Engenharia Civil da Universidade Federal de Pernambuco. O objetivo da presente pesquisa é propor um novo método baseado em simulação para uma análise de confiabilidade estrutural eficiente, utilizando um processo de filtragem para reduzir o número de avaliações necessárias para classificar os pontos da amostra em seguros ou falhos. O método é, então, comparado com outros métodos de simulação para verificar sua eficiência em diferentes problemas.

1.1 CONSIDERAÇÕES INICIAIS

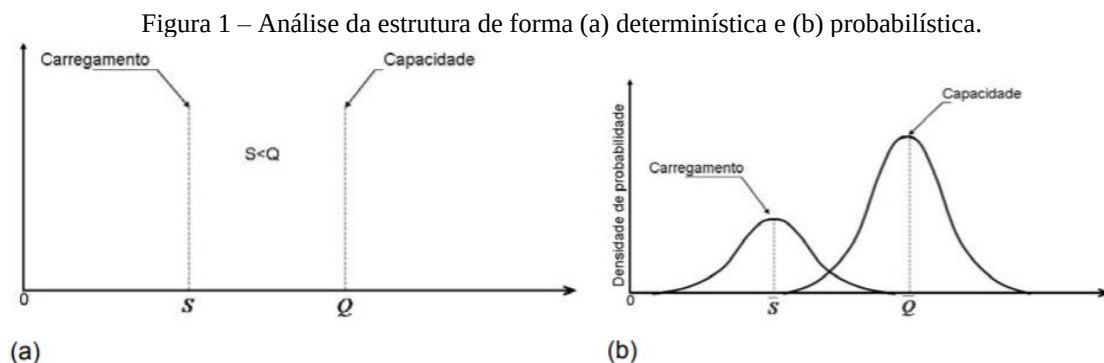
É fundamental para a engenharia civil buscar o equilíbrio entre segurança e economia, de modo que a infraestrutura proporcionada por ela atenda às exigências físicas e ambientais de cada construção e, ao mesmo tempo, seja economicamente eficiente. Nesse contexto, a engenharia estrutural deve considerar, no dimensionamento, o máximo de fatores possíveis que podem influenciar na correspondência dos cálculos com a realidade. Essa ideia guiou o desenvolvimento de metodologias de dimensionamento de estruturas mais refinadas, que consideram as incertezas presentes na teoria em relação ao que pode ocorrer na prática.

O interesse pelo comportamento de componentes estruturais está presente na humanidade há mais de dois milênios. No século III a.C., Arquimedes fez a primeira descrição matemática do fenômeno de equilíbrio, através do princípio de alavanca (DIJKSTERHUIS, 1987). A partir do século XVII d.C., a aplicação da matemática e física mecânica na indústria da construção se tornou mais direta, em especial após Galileo Galilei iniciar o estudo da resistência dos materiais em 1638 e Robert Hooke desenvolver a teoria da elasticidade em 1660. Apenas em 1773, entretanto, a análise estrutural se tornou um campo de estudo próprio, com a apresentação do *Essai sur une application des maximis règles et de minimis à quelques problèmes de statique relatifs à l'architecture* pelo físico francês Charles Augustin de Coulomb (KURRER, 2008).

No início do século XX, Mayer (1926) publicou pela primeira vez sobre a necessidade de analisar as incertezas associadas às características dos componentes estruturais (NOWAK e COLLINS, 2000). O modelo de cálculo para dimensionamento de estruturas atualmente adotado pelas normas brasileiras é o modelo semi-probabilístico dos estados-limites (ABNT,

2014; ABNT, 2008), no qual as incertezas acerca das variáveis usadas no dimensionamento são consideradas através de coeficientes de segurança parciais e de valores característicos, *i.e.*, valores cuja probabilidade de serem superados na prática é pré-determinada. A partir daí, o cálculo é feito de forma determinística, multiplicando cada variável de projeto por um coeficiente de segurança correspondente, de modo a aumentar as variáveis associadas às solicitações e reduzir as variáveis associadas às resistências (MADSEN *et al.*, 1986). Nesse cálculo as incertezas são consideradas de forma indireta, o que pode levar a resultados insatisfatórios do ponto de vista econômico e até mesmo de segurança, como demonstrado em estudos que comparam o dimensionamento baseado nas normas brasileiras com o dimensionamento considerando um modelo probabilístico (SANTOS *et al.*, 2014; SILVA, 2017; SOUZA JUNIOR, 2008). Algumas normas internacionais, apesar de usarem a mesma metodologia de dimensionamento que as normas brasileiras, já usam um modelo probabilístico para calibrar os coeficientes de segurança parciais, como é o caso dos Eurocodes (CEN, 2004; CEN, 2005).

Um modelo de dimensionamento probabilístico representa uma alternativa mais eficiente e permite uma maior noção das chances de a estrutura não atender aos requisitos de segurança. Ao considerar a variação nos valores das propriedades dos materiais e das ações envolvidas, podemos calcular a probabilidade de um determinado estado-limite ser atingido, seja ele de serviço ou último. Essa probabilidade, obtida através da análise de confiabilidade, é chamada de probabilidade de falha (P_f) e pode ser usada como parâmetro para um dimensionamento estrutural probabilístico (BECK, 2019). A Figura 1 ilustra a importância da consideração da aleatoriedade nas características da estrutura, sendo comparado com um exemplo determinístico que usa a média das variáveis de projeto. Nela, é possível ver que há uma região na qual o carregamento supera a capacidade da estrutura (Figura 1(b)).



Fonte: Rosowsky (1999).

A análise de confiabilidade estrutural é, portanto, uma forma de considerar essa aleatoriedade das propriedades das estruturas e seus carregamentos, permitindo dimensionamentos mais confiáveis do que o cálculo semi-probabilístico. As variações nas dimensões e resistências dos materiais, bem como nos valores das cargas permanentes e acidentais, são consideradas através de variáveis aleatórias (VAs). Com elas, uma função de falha, que depende do problema em questão, é construída para calcular a P_f da estrutura. Impondo um limite para essa probabilidade através da análise de risco, é possível determinar se essa estrutura é segura para os fins desejados ou não, onde o risco associado a um determinado problema considera a gravidade das consequências da falha e a frequência de ocorrência da mesma. De maneira geral, quanto maior a frequência de ocorrência e maior a gravidade, maior o risco e menor deve ser o limite imposto para a P_f (BECK, 2019). Dessa forma, a análise de confiabilidade estrutural tem sido objeto de pesquisa para aplicação em diversos problemas de engenharia, como, por exemplo, avaliação de corrosão (MUSTAFA, 2014; MOTTA *et al.*, 2021), análise de risco de barragens (PEREIRA, 2019) e otimização estrutural (MOTTA e AFONSO, 2016; SCHUELLER e JENSEN, 2008).

O método mais antigo para análise de confiabilidade é o método Monte Carlo (MC) (RUBINSTEIN, 1981), que consiste na realização de simulações dos possíveis valores das VAs do problema e contagem de quantas vezes a falha ocorreu. Anteriormente conhecido como “amostragem estatística”, o método renasceu a partir de 1945 com a invenção do primeiro computador eletrônico, inicialmente buscando soluções no estudo de reações nucleares em cadeia (METROPOLIS, 1987; WU, 2013). Apesar da simplicidade, o alto custo computacional exigido para problemas práticos da engenharia estrutural, tanto devido ao tamanho amostral quanto à complexidade das funções de falha, atraiu muitas pesquisas nas últimas décadas, levando ao desenvolvimento de diversas variantes para reduzir o custo computacional (CHEHADE e YOUNES, 2020).

Nesse trabalho, um novo método baseado no MC para realizar a análise de confiabilidade é proposto, visando os problemas cuja função de falha é localmente monotônica na região de interesse, uma vez que a maior parte dos problemas de engenharia estrutural tem essa característica. O método proposto, nomeado Monte Carlo Seletivo (SMC), visa reduzir o número de avaliações da função de falha (NFE) de uma amostra para, assim, reduzir o custo computacional exigido pelo MC tradicional sem perda de precisão.

1.2 OBJETIVOS

O objetivo geral dessa dissertação é propor e validar o método Monte Carlo Seletivo (SMC) como uma alternativa a outros métodos de análise de confiabilidade estrutural.

Os objetivos específicos são:

- Descrever o funcionamento do SMC e desenvolver um algoritmo para o mesmo;
- Analisar a influência do número de VAs da função de falha na redução do *NFE* promovido pelo SMC;
- Validar o SMC em problemas com funções de falha diversas, tanto lineares quanto não lineares;
- Aplicar o SMC em problemas de estruturas;
- Comparar o SMC com outros métodos baseados no MC tradicional, especificamente os métodos *Importance Sampling* (IS), *Subset Simulation* (SuS) e *Separable Monte Carlo* (SepMC);
- Analisar as vantagens e as limitações do SMC.

1.3 REFERENCIAL TEÓRICO

O cálculo da probabilidade de falha associada a um determinado estado-limite da estrutura dá origem ao chamado problema fundamental de confiabilidade (BECK, 2019). Esse cálculo é feito através de uma função de falha ou função de estado-limite $g(\mathbf{X})$, onde $\mathbf{X} = \{X_1, X_2, \dots, X_n\}$ são as variáveis aleatórias do problema. Uma vez definida essa função, o domínio de falha (D_f) do problema, *i.e.*, os valores de $\mathbf{X}$ para os quais a falha ocorre, é usualmente definido como a região em que $g(\mathbf{X}) \leq 0$. O cálculo da P_f é então feito através da função de densidade de probabilidades (PDF) conjunta de $\mathbf{X}$, $f_{\mathbf{X}}(\mathbf{x})$, conforme a equação

$$P_f = P(\mathbf{x} \in D_f) = \int_{D_f} f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (1.1)$$

Essa integral no espaço multidimensional pode ser de difícil solução, quando não impossível, e em muitos casos sequer é conhecida a PDF conjunta das VAs do problema, apenas as PDFs marginais. Assim, os métodos de confiabilidade buscam uma estimativa desse valor a partir das informações estatísticas e custo computacional disponíveis.

O método mais tradicional, MC, estima a probabilidade de falha através de uma amostra gerada por simulações, de modo que o tamanho amostral deve ser suficiente para que a precisão desejada seja atingida (DITLEVSEN e MADSEN, 1996). A ideia do MC é simples, se resumindo a realizar repetições de um determinado experimento, onde os valores das variáveis usadas em cada uma dessas repetições são gerados aleatoriamente através das informações estatísticas disponíveis. Apesar de robusto e simples, o MC pode se tornar inviável em muitos problemas por exigir um grande *NFE* para atingir uma precisão adequada, especialmente para problemas com uma pequena probabilidade de falha, levando a um alto custo computacional. Além disso, em problemas complicados de engenharia estrutural, a função de falha muitas vezes não tem uma forma fechada e métodos numéricos são aplicados para resolver o problema, como o Método dos Elementos Finitos (FEM), o que aumenta ainda mais esse custo (WU, 2013).

Métodos aproximados, como o Método de Confiabilidade de Primeira Ordem (FORM) (HASOFER e LIND, 1974; RACKWITZ e FIESSLER, 1978) e o Método de Confiabilidade de Segunda Ordem (SORM) (FIESSLER *et al.*, 1979) foram desenvolvidos para evitar o uso de simulações e contornar o problema do custo computacional do MC, mas acabaram apresentando dificuldades de convergência em determinados problemas, em especial para funções de falha altamente não lineares e/ou descontínuas (YANG *et al.*, 2006; YASEEN *et al.*, 2019). Além disso, quando o problema tem mais de um modo de falha, o FORM pode fornecer resultados enviesados e não conservadores, já que a busca pelo MPP desconsidera mais de um modo de falha e possivelmente a região mais relevante (DER KIUREGHIAN e DAKESSIAN, 1998; CHING *et al.*, 2009). Em casos em que o valor máximo das funções de falha é considerado, mesmo que elas sejam lineares, o operador máximo faz com que a função resultante seja altamente não linear e o FORM também apresenta dificuldade para convergir (TORII *et al.*, 2014). Esses não são problemas para o MC, que apesar de ser menos eficiente que o FORM e SORM em determinados problemas, tem as vantagens de ser mais preciso e estável (TORII *et al.*, 2019).

Posteriormente, diversos métodos baseados no MC foram desenvolvidos para reduzir a variância do mesmo sem perder totalmente suas vantagens, tais como o *Importance Sampling* (IS) (BOURGUND e BUCHER, 1986), *Subset Simulation* (SuS) (AU e BECK, 2001), *Separable Monte Carlo* (SepMC) (SMARSLOK *et al.*, 2008), *Orthogonal Plain Sampling* (HOHENBICHLER e RACKWITZ, 1988), *Line Sampling* (SCHUËLLER *et al.*, 2004) e métodos baseados em modelos substitutos (KAYMAZ, 2005; HURTADO e ALVAREZ, 2001; ZHAO *et al.*, 2015). Mais recentemente, variações desses métodos têm sido elaboradas com o

objetivo de eliminar seus pontos fracos para aplicações práticas de engenharia, em especial para aplicação em Otimização de Projeto Baseada em Confiabilidade (RBDO) (TRUONG e HA, 2020; CHAUDHURI *et al.*, 2020).

Nos últimos anos, diversos métodos foram propostos com o objetivo de reduzir o custo computacional do MC tradicional ou de métodos baseados no mesmo, principalmente para problemas com pequenas probabilidades de falha e funções de falha complexas. A Tabela 1 apresenta alguns desses métodos, onde a redução no *NFE* apresentada é em relação ao método de referência do respectivo estudo. Thedy e Liao (2021) apresentaram uma variação do *Radial-based Importance Sampling* (RBIS) que busca reduzir o *NFE* ao adicionar outra esfera para identificar os pontos seguros e desconsiderá-los nas avaliações da função de falha. Já o foco dos métodos de modelos substitutos foi melhorar a função de aprendizagem para reduzir o custo da construção desse modelo, como visto em Su *et al.* (2017), Meng *et al.* (2020) e Yun *et al.* (2018). É importante destacar que os modelos substitutos podem ser aplicados em qualquer um dos demais métodos, incluindo o proposto nesse trabalho. O melhor resultado na redução do *NFE* foi a proposta de Pan e Dias (2017), ficando entre 99,76% e 99,99% em relação ao MC tradicional. Entretanto, vale salientar que cada método foi testado em problemas diferentes e possuem limitações próprias que dependem dos conceitos nos quais foram baseados.

Tabela 1 – Métodos baseados em simulação MC encontrados na literatura recente.

Estudo	Método			Resultados	
	Nome	Sigla	Baseado em	Redução do <i>NFE</i>	Método de referência
Su <i>et al.</i> (2017)	<i>Dynamic Gaussian Process Regression surrogate model based on Monte Carlo</i>	DGPR-based MCS	Modelo substituto	0%-61,3%	<i>Response Surface Method</i> (RSM)
Meng <i>et al.</i> (2020)	<i>Active Weight Learning-Monte Carlo</i>	AWL-MCS	Modelo substituto	0%-87,74%	<i>Adaptive Kriging Monte Carlo</i> (AK-MCS)
Shayanfar <i>et al.</i> (2018)	<i>Cross-entropy based Adaptive Directional Importance Sampling</i>	ADIS	IS e simulação direcional	87,24%-99,99%	MC
Truong e Kim (2017)	<i>Improved Latin Hypercube-Effective Importance Sampling</i>	IHS-EIS	IS e <i>Improved Latin Hypercube</i> (IHS)	88%-99,85%	MC
Yun <i>et al.</i> (2018)	<i>Adaptive Kriging-Modified Importance Sampling</i>	AK-MIS	Modelo substituto e IS	30,63%-38,92%	<i>Adaptive Kriging-Importance Sampling</i> (AK-IS)
Pan e Dias (2017)	<i>Adaptive Support Vector Machine-Monte Carlo</i>	ASVM-MCS	Modelo substituto	99,76%-99,99%	MC

Cheng <i>et al.</i> (2022)	<i>Generalized Subset Simulation</i>	GSS	SuS	34,78%-45,45%	SuS
Thedy e Liao (2021)	<i>Multisphere-based Importance Sampling</i>	M-IS	<i>Radial-based Importance Sampling (RBIS)</i>	46,52%-99,91%	MC

Fonte: O autor (2022).

1.4 ORGANIZAÇÃO DA DISSERTAÇÃO

Essa dissertação é organizada em cinco capítulos, sendo eles resumidos nos próximos parágrafos.

No primeiro capítulo são feitas considerações iniciais importantes para o entendimento do trabalho, comentando de forma breve a história do dimensionamento estrutural e as vantagens do dimensionamento probabilístico em relação ao adotado pela ABNT atualmente. Também é nesse capítulo que os principais conceitos da análise de confiabilidade estrutural são apresentados, passando pelos trabalhos realizados nos últimos anos para resolver os problemas pendentes na aplicação prática do MC.

No Capítulo 2 são descritos os métodos MC, IS, SuS e SepMC, que são aplicados nesse trabalho para comparação com o SMC. Além deles, também são apresentados os dois métodos usados de forma indireta pelo IS e pelo SuS: o FORM e MCMC, respectivamente. O cálculo da probabilidade de falha por cada método é explicado destacando suas especificidades, finalizando com os respectivos algoritmos.

O Capítulo 3 inicia definindo o problema que o método proposto busca resolver e o conceito de otimalidade de Pareto, fundamental para o entendimento do método. Em seguida, o SMC é descrito e seu funcionamento é ilustrado com um exemplo simples. Nesse capítulo, o algoritmo do método é mostrado passo a passo e suas limitações são detalhadas.

No Capítulo 4, um estudo paramétrico é feito para avaliar o impacto da dimensão do problema na eficiência do SMC. Além disso, 13 exemplos de referência e dois problemas de engenharia estrutural são realizados e os resultados do *NFE* e da probabilidade de falha são comparados com os métodos IS, SuS, SepMC e MC tradicional. Os resultados são discutidos, evidenciando as vantagens e desvantagens do método proposto em relação aos demais.

Por fim, no Capítulo 5, as conclusões dessa dissertação são apresentadas, resumindo o desempenho do SMC de maneira geral. Também são deixadas sugestões para estudos futuros, sejam eles com foco no SMC ou mesmo em outros métodos baseados no MC tradicional.

2 MÉTODOS DE CONFIABILIDADE ESTRUTURAL

O objetivo da análise de confiabilidade na engenharia estrutural é calcular a probabilidade de falha, *i.e.*, a chance de uma estrutura atingir o estado-limite que se deseja observar, dependendo do problema (BECK, 2019). Para esse fim, as variáveis do problema, tais como resistências dos materiais e solicitações estruturais, são consideradas VAs. Além disso, é possível ainda avaliar essas incertezas e o estado-limite do problema como dependentes do tempo ou estáticos (ZHANG *et al.*, 2017).

Na análise de confiabilidade estrutural tradicional, dois grupos de métodos são comumente utilizados: os métodos de simulação, que são basicamente o MC e suas variações, e os métodos de transformação, que são o FOSM, FORM, SORM e suas variações (LOPEZ e BECK, 2012). Os métodos de simulação têm como vantagens a sua simplicidade de implementação, robustez e um nível de acurácia ajustável, mesmo para problemas não lineares, mas podem exigir um alto custo computacional. Os métodos de transformação, por sua vez, têm um custo computacional menor, mas também uma menor precisão e robustez, não podendo ser aplicados em determinados problemas (ROOS *et al.*, 2006).

2.1 MÉTODO DE CONFIABILIDADE DE PRIMEIRA ORDEM (FORM)

Os métodos de transformação são chamados assim por consistirem em transformar as VAs do problema, $\mathbf{X} = \{X_1, X_2, \dots, X_n\}$, em VAs normais padrão $\mathbf{U} = \{U_1, U_2, \dots, U_n\}$, para calcular a probabilidade de falha no chamado espaço isoprobabilístico (LOPEZ e BECK, 2012). Nesse espaço, o método busca o ponto mais provável de falha (MPP), que é o ponto pertencente ao domínio de falha que fica mais próximo à origem e, consequentemente, com maior probabilidade de ocorrência. A distância entre o MPP (a falha) e a origem, chamada de índice de confiabilidade (β), é um valor que indica o nível de segurança e é utilizado para estimar a probabilidade de falha no FORM (ROOS *et al.*, 2006). Através desse índice, a probabilidade de falha pode ser calculada pela equação

$$P_f = \Phi(-\beta) \quad (2.1)$$

sendo Φ a função de probabilidade cumulativa (CDF) normal padrão unidimensional.

É importante destacar que existe um caso especial desse índice, chamado de índice de confiabilidade de Cornell (β_C), que só é válido para o caso em que $G = g(\mathbf{X})$ tem distribuição

normal (SELMÍ *et al.*, 2010). O β_c usa apenas a média e o desvio-padrão para calcular a probabilidade de falha (CORNELL, 1971). A transformação de Hasofer e Lind (HASOFER e LIND, 1974) para levar G ao espaço normal padrão resulta em

$$U = \frac{G - \mu_G}{\sigma_G} \quad (2.2)$$

onde U é uma VA normal padrão e μ_G e σ_G são a média e desvio-padrão de G , respectivamente.

A probabilidade de falha se torna

$$P_f = P(G < 0) = P\left(U < -\frac{\mu_G}{\sigma_G}\right) = \Phi\left(-\frac{\mu_G}{\sigma_G}\right) \quad (2.3)$$

Assim, a menor distância entre a falha e a origem de U , caso G seja normal, é

$$\beta_c = \frac{\mu_G}{\sigma_G} \quad (2.4)$$

e a probabilidade de falha pode ser calculada por esse índice pela Equação 2.1.

De maneira geral, o β usa a PDF conjunta das variáveis $\mathbf{U}$, ou seja, de $\mathbf{X}$ no espaço isoprobabilístico (HASOFER e LIND, 1974; RACKWITZ e FIESSLER, 1978). Chamando essa PDF conjunta de $f_U(\mathbf{u})$, uma vez tendo a função de falha no espaço isoprobabilístico, $g(\mathbf{U})$, o valor de β é encontrado ao resolver o problema de otimização (SELMÍ *et al.*, 2010)

$$\text{Minimize } \|\mathbf{U}\| = \sqrt{\mathbf{U}^T \mathbf{U}} \quad (2.5)$$

$$\text{Sujeito a } g(\mathbf{U}) = 0$$

ou seja, minimizando a distância entre a origem e o domínio de falha no espaço normal padrão multivariado. Baseado nisso, β pode ser escrito como

$$\beta = \sqrt{\mathbf{u}^{*T} \mathbf{u}^*} \quad (2.6)$$

onde $\mathbf{u}^*$ representa o MPP. Obter o β dessa forma permite o cálculo da probabilidade de falha pela Equação 2.1 sem a necessidade de ter a PDF de G .

Entre os métodos de transformação, o FORM é uma escolha simples e eficiente que fornece um resultado exato da P_f para uma função de falha linear no espaço isoprobabilístico. Para uma função de falha não linear, o FORM pode dar um bom resultado aproximado para a P_f , podendo ainda ser melhorado com o uso do SORM (MELCHERS e BECK, 2017).

2.1.1 Descrição do método

O FORM é um método que se baseia na aproximação linear da função de falha obtida pelo primeiro termo da expansão em série de Taylor, mostrada na Equação 2.7, onde $\tilde{g}(\mathbf{u})$ é a aproximação de $g(\mathbf{u})$ e $\mathbf{u}_p$ é o ponto sobre o qual a série é centrada. Como todo método de transformação, ele precisa que as VAs $\mathbf{X}$, possivelmente correlacionadas, sejam transformadas em variáveis $\mathbf{U}$, com distribuições normais padrão e descorrelacionadas.

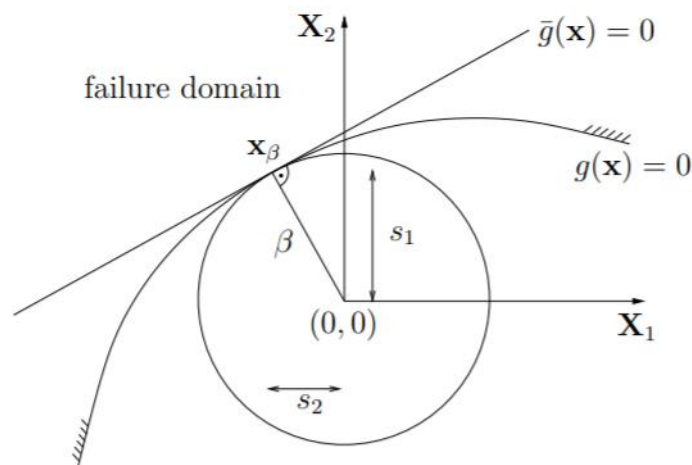
$$\tilde{g}(\mathbf{u}) = g(\mathbf{u}_p) + \nabla g(\mathbf{u}_p)(\mathbf{u} - \mathbf{u}_p) \quad (2.7)$$

O β dessa aproximação é igual ao de $g(\mathbf{u})$ quando ela é centrada no MPP (BECK, 2019). Isso nos permite estimar a P_f através da probabilidade de falha de $\tilde{g}(\mathbf{u})$, $\tilde{P}_f$, conforme a Equação 2.8.

$$P_f \approx \tilde{P}_f = \Phi(-\beta) \quad (2.8)$$

O erro associado à $\tilde{P}_f$ ocorre quando a função de falha é não linear, já que $\tilde{g}(\mathbf{u})$ tem uma fronteira de falha linear. Esse erro é exemplificado na Figura 2, onde o MPP foi escrito como $\mathbf{x}_\beta$ e a aproximação da função de falha como $\bar{g}(\mathbf{x})$.

Figura 2 – Aproximação linear da função de falha no espaço padrão reduzido.



Fonte: Unger e Roos (2004).

Assim, o FORM consiste em, após levar as VAs para o espaço isoprobabilístico, encontrar o MPP e β para calcular a $\tilde{P}_f$. O MPP é encontrado resolvendo um problema de otimização, já que se resume a um problema de minimização de $\|\mathbf{U}\|$, sendo o algoritmo mais tradicional usado para esse fim o Hasofer-Lind-Rackwitz-Fiessler (HASOFER e LIND, 1974; RACKWITZ e FIESSLER, 1978), conhecido como HL-RF.

O HL-RF é eficiente em diversos casos, mas pode apresentar problemas de convergência como oscilação periódica e soluções caóticas mesmo para funções de falha não lineares simples (YANG *et al.*, 2006). Outros algoritmos foram desenvolvidos buscando resolver esse problema, como os algoritmos *Improved HL-RF* (iHL-RF) (ZHANG e DER KIUREGHIAN, 1997), *Conjugate HL-RF* (CHL-RF) (KESHTEGAR e MIRI, 2014) e *New HL-RF* (nHL-RF) (SANTOS *et al.*, 2012).

2.1.2 Transformação para o espaço isoprobabilístico

Quando todas as variáveis possuem distribuição normal e são independentes entre si, a transformação de $\mathbf{X}$ em $\mathbf{U}$ pode ser feita pela transformação de Hasofer e Lind na sua forma matricial, resumida nas equações (BECK, 2019)

$$\mathbf{u} = J_{ux}(\mathbf{x} - \boldsymbol{\mu}) \quad (2.9)$$

$$\mathbf{x} = J_{xu}\mathbf{u} + \boldsymbol{\mu} \quad (2.10)$$

onde $\boldsymbol{\mu} = \{\mu_{x_1}, \mu_{x_2}, \dots, \mu_{x_n}\}$ é o vetor das médias das VAs no espaço original e J_{ux} e J_{xu} são as matrizes jacobianas de $\mathbf{u}$ em relação a $\mathbf{x}$ e de $\mathbf{x}$ em relação a $\mathbf{u}$, respectivamente (equações 2.11 e 2.12).

$$J_{ux} = \left[\frac{\partial u_i}{\partial x_j} \right]_{i=1,2,\dots,n; j=1,2,\dots,n} \quad (2.11)$$

$$J_{xu} = \left[\frac{\partial x_i}{\partial u_j} \right]_{i=1,2,\dots,n; j=1,2,\dots,n} \quad (2.12)$$

No caso mais geral, as VAs não necessariamente têm distribuição normal e podem ter coeficientes de correlação (ρ) não nulos. Para aproveitar os benefícios de simetria do espaço normal padrão sem correlação, é necessário usar a transformação de Nataf (NATAF, 1962) ou a transformação de Rosenblatt (ROSENBLATT, 1952).

A transformação de Nataf consiste em transformar as variáveis $\mathbf{X}$ nas variáveis com distribuição normal padrão possivelmente correlacionadas $\mathbf{Z}$ e, por fim, transformá-las nas variáveis descorrelacionadas no espaço normal padrão $\mathbf{U}$. As informações disponíveis para isso, em geral, são as PDFs marginais de $\mathbf{X}$ e os coeficientes de correlação entre pares de variáveis, $\rho_{X_{ij}}$, onde $i = 1, 2, \dots, n$ e $j = 1, 2, \dots, n$.

Antes da primeira parte da transformação de Nataf, caso as PDFs marginais não sejam normais, é necessário usar o Princípio da Aproximação Normal (DITLEVSEN, 1981). Esse princípio consiste em, para um dado valor x_i da variável X_i , encontrar uma PDF normal equivalente que tenha o mesmo conteúdo de probabilidade acumulado da PDF original nesse ponto. Como muitas PDFs podem atender a esse requisito, uma segunda condição é imposta: ter também o mesmo valor de PDF. Dessa forma, para obter a PDF normal equivalente, devem ser satisfeitas as equações

$$F_{X_i}(x_i) = F_{X_i}^{neq}(x_i) = \Phi(z_i) \quad (2.13)$$

$$f_{X_i}(x_i) = f_{X_i}^{neq}(x_i) = \frac{\varphi(z_i)}{\sigma_{X_i}^{neq}} \quad (2.14)$$

onde F_{X_i} é a CDF de X_i , f_{X_i} é a PDF de X_i , φ é a PDF normal padrão e o índice *neq* é usado para se referir à PDF normal equivalente.

Essas condições permitem a obtenção da média e desvio-padrão da PDF normal equivalente:

$$\mu_{X_i}^{neq} = x_i - z_i \sigma_{X_i}^{neq} \quad (2.15)$$

$$\sigma_{X_i}^{neq} = \frac{\varphi(z_i)}{f_{X_i}(x_i)} \quad (2.16)$$

onde $z_i = \Phi^{-1}(F_{X_i}(x_i))$, sendo Φ^{-1} a CDF normal padrão inversa.

Uma vez que as PDFs marginais foram convertidas para distribuições normais equivalentes, basta usar a transformação de Hasofer e Lind para concluir a primeira parte da transformação de Nataf. Nessa primeira parte acontece a transformação de $\mathbf{X}$ para $\mathbf{Z}$, onde os coeficientes de correlação devem ser também convertidos para o espaço normal padrão, o que pode ser realizado através de relações aproximadas entre seus valores para $\mathbf{X}$ e $\mathbf{Z}$ (DER KIUREGHIAN e LIU, 1986). Vale destacar que a PDF normal equivalente só é válida para o

ponto em questão, devendo ser recalculada cada vez que o ponto mudar. A transformação de Hasofer e Lind, nesse caso, se torna

$$\mathbf{z} = J_{zx}(\mathbf{x} - \boldsymbol{\mu}) \quad (2.17)$$

$$\mathbf{x} = J_{xz}\mathbf{z} + \boldsymbol{\mu} \quad (2.18)$$

A segunda parte acontece ao eliminar a correlação presente em $\mathbf{Z}$, levando a $\mathbf{U}$, através das equações

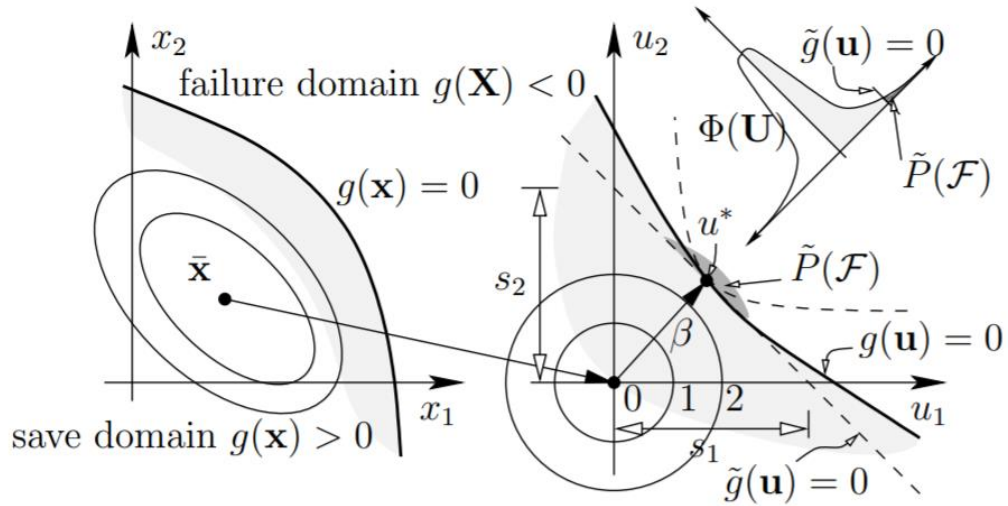
$$\mathbf{u} = J_{uz}\mathbf{z} \quad (2.19)$$

$$\mathbf{z} = J_{zu}\mathbf{u} \quad (2.20)$$

sendo J_{uz} e J_{zu} as matrizes jacobianas de $\mathbf{u}$ em relação a $\mathbf{z}$ e de $\mathbf{z}$ em relação a $\mathbf{u}$, respectivamente. Essas matrizes podem ser encontradas usando a decomposição ortogonal ou a decomposição de Cholesky (BECK, 2019).

Um exemplo da transformação de $\mathbf{X}$ para $\mathbf{U}$, considerando um espaço bidimensional, é mostrado na Figura 3.

Figura 3 – Transformação de $\mathbf{X}$ para $\mathbf{U}$, onde $\mathbf{X}$ são variáveis possivelmente correlacionadas no espaço original e $\mathbf{U}$ sem correlação no espaço isoprobabilístico.



Fonte: Roos *et al.* (2006).

2.1.3 Busca pelo MPP

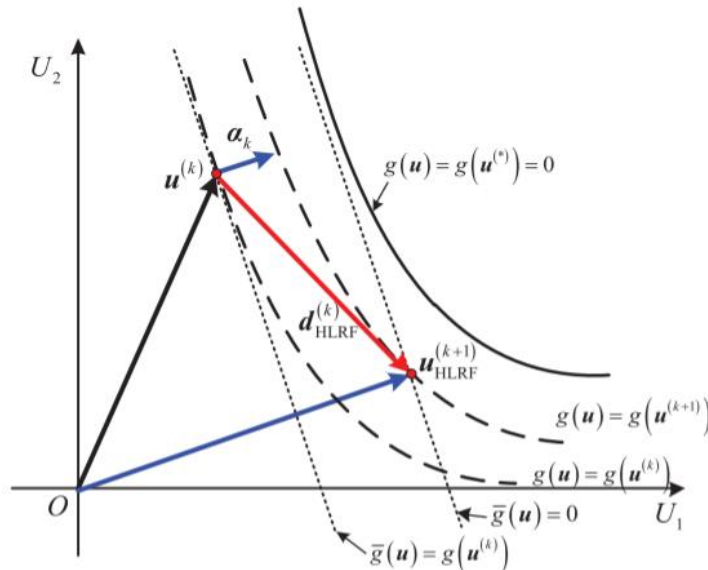
O algoritmo mais tradicional para buscar o MPP é o HL-RF (HASOFER e LIND, 1974; RACKWITZ e FIESSLER, 1978), que realiza essa tarefa no espaço isoprobabilístico. Ele usa duas informações fundamentais para guiar seu mecanismo: a primeira é o fato de que o MPP fica em $\tilde{g}(\mathbf{u}) = 0$ se essa aproximação linear estiver centrada nele e a segunda é que o gradiente de $g(\mathbf{u}^*)$ aponta na mesma direção de $\mathbf{u}^*$, onde $\mathbf{u}^*$ é o MPP. Com essas duas informações, a ideia do algoritmo é centrar a aproximação linear $\tilde{g}(\mathbf{u})$ em um ponto qualquer $\mathbf{u}_k$ e, a partir dele, procurar um novo ponto $\mathbf{u}_{k+1}$ que esteja em $\tilde{g}(\mathbf{u}_k) = 0$ para centrar uma nova aproximação, repetindo esse processo até que o ponto centrado esteja sobre $\tilde{g}(\mathbf{u}) = 0$. Para escolher entre os infinitos pontos que satisfazem $\tilde{g}(\mathbf{u}_k) = 0$, é imposta a condição de que o novo ponto deve ter a mesma direção de $\nabla g(\mathbf{u}_k)$.

Resumindo, o HL-RF é um processo iterativo, onde o novo ponto é buscado na fronteira de falha da aproximação usada para o ponto anterior e sempre na direção do gradiente da função de falha. A determinação de $\mathbf{u}_{k+1}$ é feita pela equação

$$\mathbf{u}_{k+1} = \frac{[\nabla g(\mathbf{u}_k)^T \mathbf{u}_k - g(\mathbf{u}_k)]}{\|\nabla g(\mathbf{u}_k)\|^2} \nabla g(\mathbf{u}_k) \quad (2.21)$$

A Figura 4 ilustra o funcionamento do HL-RF em um espaço bidimensional, onde α_k é um vetor unitário que indica o sentido de crescimento de $g(\mathbf{u})$ e a notação usada para a aproximação linear foi $\tilde{g}(\mathbf{u})$.

Figura 4 – Exemplo do funcionamento do HL-RF.



Fonte: Zhou et al. (2020).

Outros algoritmos de otimização podem ser usados no lugar do HL-RF, como, por exemplo, o *Sequential Quadratic Programming* (SQP) e o *Modified Feasible Direction* (MFD) (YOUN *et al.*, 2004). Além disso, também existem algoritmos construídos a partir de melhorias no HL-RF, como os presentes em Liu e Der Kiureghian (1991), Keshtegar e Miri (2014) e Santos *et al.* (2012).

2.1.4 Algoritmo do FORM

O algoritmo do FORM pode ser descrito nos seguintes passos, onde usualmente a média é escolhida como ponto inicial (BECK, 2019):

- 1 – Determine os coeficientes de correlação equivalentes de $\mathbf{Z}$ ($\rho_{z_{ij}}$, $i = 1, 2, \dots, n$ e $j = 1, 2, \dots, n$) e matrizes J_{uz} e J_{zu} ;
- 2 – Escolha o ponto inicial $\mathbf{x}_k$ para $k = 0$;
- 3 – Calcule ou atualize as matrizes J_{xz} e J_{zx} , devendo ser atualizadas a cada iteração;
- 4 – Faça a transformação de $\mathbf{x}$ para $\mathbf{u}$ no ponto $\mathbf{x}_k$;
- 5 – Avalie $g(\mathbf{x}_k)$;
- 6 – Calcule $\nabla g(\mathbf{x}_k)$ e transforme em $\nabla g(\mathbf{u}_k)$;
- 7 – Calcule $\mathbf{u}_{k+1}$ pelo HL-RF ou outro algoritmo de busca do MPP;
- 8 – Transforme $\mathbf{u}_{k+1}$ em $\mathbf{x}_{k+1}$;
- 9 – Verifique o critério de convergência adotado e, se descumprido, volte ao passo 3 com $k = k + 1$;
- 10 – Por fim, calcule o índice de confiabilidade $\beta = \|\mathbf{u}^*\|$ e a probabilidade de falha estimada pela Equação 2.8.

2.2 MÉTODOS DE SIMULAÇÃO

Os métodos de simulação consistem em realizar experimentos aleatórios em computador a partir das PDFs das VAs do problema em questão, de modo a gerar uma amostra sem precisar realizar os experimentos na prática. O primeiro método desse tipo foi o MC, a partir do qual surgiram todos os demais métodos de simulação (METROPOLIS, 1987; CHEHADE e YOUNES, 2020).

2.2.1 Método Monte Carlo (MC)

O método Monte Carlo (RUBINSTEIN, 1981) é o mais antigo dos métodos de simulação e foi um dos primeiros métodos a ser aplicado na confiabilidade estrutural (SHINOZUKA, 1983; AUGUSTI *et al.*, 1984). Ele também é conhecido como Monte Carlo Simples, Direto, Bruto ou Cru (BECK, 2019), para diferenciar dos métodos baseados no mesmo. O MC consiste em gerar uma amostra de tamanho N do vetor $\mathbf{X} = \{X_1, X_2, \dots, X_n\}$, de modo que os elementos da amostra estejam distribuídos de acordo com a PDF conjunta de $\mathbf{X}$, onde n é o número de VAs do problema. Com ela, a função de falha é avaliada N vezes e as informações estatísticas, como média e desvio-padrão, podem ser obtidas. Uma vez que a probabilidade de falha é dada por

$$P_f = \int_{D_f} f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (2.22)$$

é possível definir uma função indicadora $I(\mathbf{x})$ tal que

$$I(\mathbf{x}) = \begin{cases} 1, & \mathbf{x} \in D_f \\ 0, & \mathbf{x} \notin D_f \end{cases} \quad (2.23)$$

o que permite reescrever a P_f como uma integral sobre todo o domínio D :

$$P_f = \int_D I(\mathbf{x}) f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (2.24)$$

É possível notar que essa integral é igual à esperança da função $I(\mathbf{x})$, pela definição de esperança ($E(\cdot)$). Essa esperança pode ser estimada pela média de uma amostra finita com tamanho N , resultando em uma estimativa da probabilidade de falha dada pela equação

$$P_f = \int_D I(\mathbf{x}) f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} = E(I(\mathbf{x})) \approx \tilde{P}_f = \frac{1}{N} \sum_{k=1}^N I(\mathbf{x}_k) = \frac{NFP}{N} \quad (2.25)$$

onde $\mathbf{x}_k$ é o k -ésimo elemento da amostra e NFP é o número de pontos da amostra no domínio de falha, *i.e.*, onde $I(\mathbf{x}_k) = 1$.

2.2.1.1 Características do estimador da probabilidade de falha do MC

Enquanto a P_f é determinística para um dado problema, a $\tilde{P}_f$ é a média amostral da função $I(\mathbf{x})$, que tem distribuição de Bernoulli (BECK, 2019). Como a esperança de uma VA

com distribuição de Bernoulli é p , onde p é a probabilidade de ela ser igual a 1, então a esperança da $\tilde{P}_f$ é

$$E(\tilde{P}_f) = E\left(\frac{1}{N} \sum_{k=1}^N I(\mathbf{x}_k)\right) = \frac{N}{N} E(I(\mathbf{x})) = P_f \quad (2.26)$$

Isso significa que o estimador dado pelo MC para a P_f é considerado não tendencioso, pois quando $N \rightarrow \infty$, o estimador $\tilde{P}_f \rightarrow P_f$.

Já a variância de uma VA de Bernoulli é igual a $p(1-p)$, o que dá a variância do estimador $\tilde{P}_f$:

$$Var(\tilde{P}_f) = Var\left(\frac{1}{N} \sum_{k=1}^N I(\mathbf{x}_k)\right) = \frac{N}{N^2} Var(I(\mathbf{x})) = \frac{P_f(1-P_f)}{N} \quad (2.27)$$

Finalmente, o coeficiente de variação do MC pode ser calculado através das equações 2.26 e 2.27, conforme mostrado na Equação 2.28.

$$CoV(\tilde{P}_f) = \frac{\sqrt{Var(\tilde{P}_f)}}{E(\tilde{P}_f)} = \frac{\sqrt{P_f(1-P_f)}}{P_f \sqrt{N}} = \frac{\sqrt{P_f(1-P_f)}}{P_f \sqrt{N}} \frac{\sqrt{P_f}}{\sqrt{P_f}} = \frac{P_f \sqrt{(1-P_f)}}{P_f \sqrt{N} \sqrt{P_f}} = \sqrt{\frac{1-P_f}{NP_f}} \quad (2.28)$$

É possível notar que à medida que a $P_f \rightarrow 0$ para um valor constante do coeficiente de variação, N tende ao infinito. Assim, quando a probabilidade de falha é muito pequena, o MC exige uma amostra muito grande para aumentar sua precisão para o nível desejado. Isso em muitos casos pode ser proibitivo, sendo essa a principal motivação para o desenvolvimento de métodos alternativos a partir do MC.

2.2.1.2 Índice de Confiabilidade do MC

O MC calcula diretamente a $\tilde{P}_f$ sem passar pelo índice de confiabilidade β , diferentemente do FORM. Entretanto, pode-se definir um índice de confiabilidade equivalente para o MC (BECK, 2019):

$$\beta_{MC} = -\Phi^{-1}(\tilde{P}_f) \quad (2.29)$$

Uma vez que esse índice de confiabilidade é obtido de forma diferente daquele calculado pelo FORM, alguns autores escrevem o β obtido pelo FORM como β_{FORM} para diferenciar (ALI, 2012; PIRES *et al.*, 2019).

2.2.1.3 Algoritmo do MC

O algoritmo do MC tradicional pode ser resumido nos seguintes passos, admitindo que as VAs do problema são independentes (BECK, 2019):

- 1 – Gere N pontos amostrais $\mathbf{x}_k = \{x_{k,1}, x_{k,2}, \dots, x_{k,n}\}$ a partir das PDFs marginais de cada VA;
- 2 – Calcule a função de falha $g(\mathbf{x}_k)$ para cada ponto;
- 3 – Avalie a função indicadora $I(\mathbf{x}_k)$ para cada ponto, contando os pontos nos quais $g(\mathbf{x}_k) \leq 0$, ou seja, determinando o NFP ;
- 4 – Calcule a $\tilde{P}_f$ pela Equação 2.25;
- 5 – Calcule o coeficiente de variação da $\tilde{P}_f$ pela Equação 2.28.

2.2.2 Importance Sampling (IS)

Uma forma de reduzir a variância da $\tilde{P}_f$ e, conseqüentemente, o tamanho N da amostra necessário para uma determinada precisão no MC é através do IS. Esse método consiste em, ao invés de usar a PDF conjunta $f_{\mathbf{X}}(\mathbf{x})$ para gerar os elementos da amostra, usar uma função de amostragem $h_{\mathbf{X}}(\mathbf{x})$ que faça com que mais pontos estejam na região da falha (AU e WANG, 2014). Como a maior dificuldade em casos com probabilidade de falha pequena, comuns na confiabilidade estrutural, é que demandam uma amostra de tamanho grande para que alguns pontos caiam na região de falha, o objetivo do IS é fazer com que muitos pontos sejam gerados no domínio de falha mesmo para amostras menores.

Para tal feito, o integrando da equação da probabilidade de falha é multiplicado e dividido por $h_{\mathbf{X}}(\mathbf{x})$, como mostra a Equação 2.30.

$$P_f = \int_D I(\mathbf{x}) \frac{f_{\mathbf{X}}(\mathbf{x})}{h_{\mathbf{X}}(\mathbf{x})} h_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (2.30)$$

Com isso, a P_f fica igual à esperança da função $I(\mathbf{x}) \frac{f_{\mathbf{X}}(\mathbf{x})}{h_{\mathbf{X}}(\mathbf{x})}$ em relação a $h_{\mathbf{X}}(\mathbf{x})$, ou seja, admitindo que $\mathbf{X}$ é distribuído conforme $h_{\mathbf{X}}(\mathbf{x})$. De maneira similar ao MC, essa esperança pode ser estimada através de uma amostra finita de tamanho N , desde que tenha sido gerada por $h_{\mathbf{X}}(\mathbf{x})$ ao invés de $f_{\mathbf{X}}(\mathbf{x})$:

$$P_f = E_h \left(I(\mathbf{x}) \frac{f_X(\mathbf{x})}{h_X(\mathbf{x})} \right) \approx \tilde{P}_f = \frac{1}{N} \sum_{k=1}^N I(\mathbf{x}_k) \frac{f_X(\mathbf{x}_k)}{h_X(\mathbf{x}_k)}, \quad \mathbf{X} \sim h_X(\mathbf{x}) \quad (2.31)$$

Em suma, isso mostra que a $\tilde{P}_f$ pode ser obtida por uma função de amostragem no lugar da PDF original de $\mathbf{X}$, desde que os valores de $I(\mathbf{x})$ sejam multiplicados pelos pesos

$$w_k = \frac{f_X(\mathbf{x}_k)}{h_X(\mathbf{x}_k)}, \quad k = 1, 2, \dots, N \quad (2.32)$$

É importante salientar, entretanto, que nem toda função de amostragem leva a uma boa estimativa da probabilidade de falha, podendo até piorar a precisão da $\tilde{P}_f$ em comparação com o MC. É fundamental que $h_X(\mathbf{x})$ seja escolhido de tal forma que reduza ao máximo a variância da $\tilde{P}_f$.

2.2.2.1 Características do estimador da probabilidade de falha do IS

Da mesma forma que foi demonstrada para o MC, é fácil ver que o IS também leva a um estimador não enviesado da probabilidade de falha:

$$E_h(\tilde{P}_f) = P_f \quad (2.33)$$

A variância da $\tilde{P}_f$ estimada pelo IS se torna, portanto:

$$Var_h(\tilde{P}_f) = \frac{1}{N} \left(\int_D I(\mathbf{x})^2 \frac{f_X(\mathbf{x})^2}{h_X(\mathbf{x})^2} h_X(\mathbf{x}) d\mathbf{x} - P_f^2 \right) = \frac{1}{N} \int_D \frac{(I(\mathbf{x})f_X(\mathbf{x}) - P_f h_X(\mathbf{x}))^2}{h_X(\mathbf{x})^2} d\mathbf{x} \quad (2.34)$$

Isso significa que a variância dessa estimativa vai a zero quando

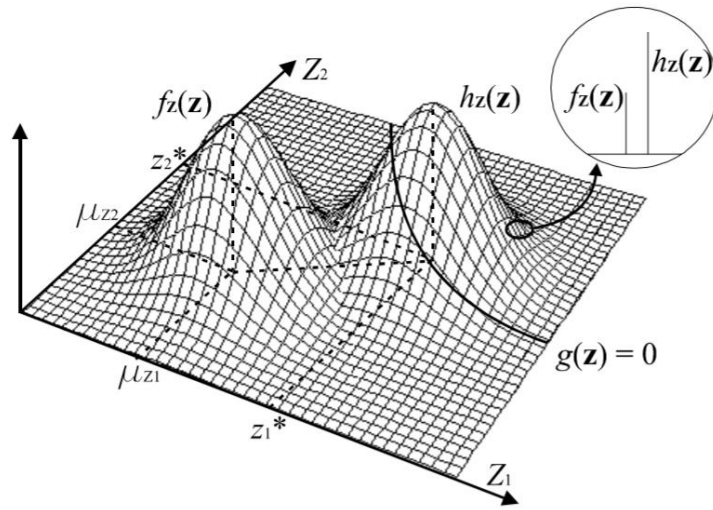
$$h_X(\mathbf{x}) = \frac{I(\mathbf{x})f_X(\mathbf{x})}{P_f} \quad (2.35)$$

Não é possível usar a Equação 2.35 na prática, já que o valor que se deseja conhecer está no denominador. Porém, é possível usar informações preliminares para fazer uma escolha adequada de $h_X(\mathbf{x})$, como saber a localização do MPP. Isso é explicado com mais detalhes na subseção 2.2.2.2.

2.2.2.2 Função de amostragem usando o MPP

Uma boa escolha para $h_X(\mathbf{x})$, de modo que mais pontos da amostra caiam na falha, é uma PDF centrada no MPP, que por sua vez pode ser obtido pelo HL-RF ou outro algoritmo de busca usado no FORM (BORGUND e BUCHER, 1986). A Figura 5 ilustra essa ideia em um espaço bidimensional, onde $f_Z(\mathbf{z})$ é a PDF original, $h_Z(\mathbf{z})$ é a função de amostragem, $g(\mathbf{z})$ é a função de falha e o ponto $\mathbf{z}^* = \{z_1^*, z_2^*\}$ é o MPP.

Figura 5 – IS com a função de amostragem centrada no MPP.



Fonte: Beck (2019).

Admitindo que $\mathbf{U}$ é o equivalente à variável $\mathbf{X}$ no espaço isoprobabilístico e que $\mathbf{u}^*$ é o MPP, a função de amostragem a ser usada pode ter a forma $h_U(\mathbf{u}) = \varphi_n(\mathbf{u} - \mathbf{u}^*)$, onde φ_n é a função de densidade de probabilidade conjunta normal padrão n -dimensional (BECK, 2019; ROOS *et al.*, 2006). A amostra é, então, gerada no espaço isoprobabilístico através de $h_U(\mathbf{u})$. Uma vez tendo a amostra e os pesos para cada ponto, $w(\mathbf{u}) = \frac{f_U(\mathbf{u})}{h_U(\mathbf{u})}$, a probabilidade de falha estimada pelo IS pode ser calculada diretamente no espaço isoprobabilístico pela equação

$$\tilde{P}_f = \sum_{i=1}^N I(\mathbf{u}_i) w(\mathbf{u}_i) \quad (2.36)$$

onde $\mathbf{u}_i$ é o equivalente ao ponto $\mathbf{x}_i$ no espaço isoprobabilístico e a função indicadora pode ser calculada no espaço original, pois $I(\mathbf{x}_i) = I(\mathbf{u}_i)$ para um ponto $\mathbf{x}_i$ qualquer.

Essa metodologia permite que aproximadamente metade dos pontos gerados caiam no domínio de falha quando a função de falha não é altamente não linear (BECK, 2019). Assim,

ocorre uma redução no coeficiente de variação da probabilidade falha ($CoV(\tilde{P}_f)$) em relação ao obtido pelo MC, se a função de amostragem for escolhida adequadamente.

2.2.2.3 Algoritmo do IS

O algoritmo do IS pode ser resumido nos seguintes passos, uma vez escolhendo a função de amostragem $h_X(\mathbf{x})$:

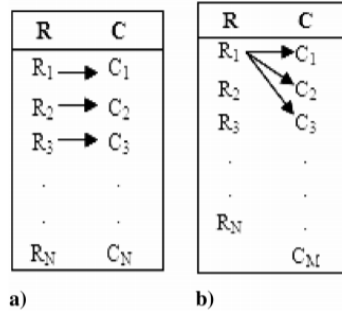
- 1 – Gere N pontos amostrais $\mathbf{x}_k = \{x_{k,1}, x_{k,2}, \dots, x_{k,n}\}$ a partir de $h_X(\mathbf{x})$;
- 2 – Calcule a função de falha $g(\mathbf{x}_k)$ para cada ponto;
- 3 – Avalie a função indicadora $I(\mathbf{x}_k)$ para cada ponto, contando os pontos nos quais $g(\mathbf{x}_k) \leq 0$, ou seja, determinando o NFP ;
- 4 – Calcule os pesos w_k usando a Equação 2.32;
- 5 – Calcule a $\tilde{P}_f$ pela Equação 2.31.

Nesse trabalho, foi acrescentado um critério de convergência no algoritmo para que ele repetisse o processo, aumentando N até que o $CoV(\tilde{P}_f)$ desejado fosse atingido. Nele, é feita uma verificação do $CoV(\tilde{P}_f)$ a cada grupo de pontos gerados (por exemplo, a cada mil pontos) e o algoritmo é interrompido quando o $CoV(\tilde{P}_f)$ ficar abaixo do nível fixado.

2.2.3 Separable Monte Carlo (SepMC)

Quando as VAs de resistência e as VAs de solicitação são independentes, elas podem ser amostradas separadamente. Ao invés de gerar N pontos de todas as VAs ao mesmo tempo, o SepMC (SMARSLOK *et al.*, 2008) consiste em gerar separadamente as amostras dessas variáveis, podendo essas amostras terem tamanhos distintos. Em outras palavras, sendo C a parte da função de falha com as VAs de resistência e R a parte da função de falha com as VAs de solicitação, geramos M elementos amostrais C_i ($i = 1, 2, \dots, M$) e N elementos amostrais R_j ($j = 1, 2, \dots, N$). A amostra total é dada pela combinação de cada elemento da amostra de C com cada elemento da amostra de R , resultando em $M \times N$ pontos amostrais (RAVISHANKAR *et al.*, 2010). A Figura 6 mostra a diferença entre o SepMC e o MC tradicional, onde nela R se refere a solicitação (*response*) e C se refere a resistência (*capacity*).

Figura 6 – Comparação da amostra gerada pelo (a) MC tradicional com a gerada pelo (b) SepMC.



Fonte: Ravishankar *et al.* (2010).

Usando a mesma ideia demonstrada para o MC, é fácil de ver que a probabilidade de falha do SepMC é dada por

$$\tilde{P}_f = \frac{1}{MN} \sum_{j=1}^N \sum_{i=1}^M I(g(C_i, R_j) < 0) \quad (2.37)$$

onde $g(C, R)$ é a função de falha e $I(C, R)$ é a função indicadora cujo valor é 1 quando ocorre a falha e 0 caso contrário.

O objetivo principal do SepMC é reduzir o tamanho amostral da parte mais custosa da função de falha (resistência ou solicitação) para reduzir o custo computacional envolvido nas avaliações da mesma. Para isso, ele aumenta o número de avaliações da função de falha em relação ao MC, obtendo uma precisão melhor que o MC com tamanho amostral N , mas uma precisão pior que o MC com tamanho amostral $M \times N$ (SMARSLOK *et al.*, 2008). Entretanto, o custo computacional do MC para uma amostra de tamanho $M \times N$ seria muito maior que o do SepMC, o que evidencia a vantagem do último nesse caso.

Em resumo, o SepMC é um método que simplifica a função de falha para uma forma $g(C, R) = C - R$ e calcula previamente sua parte mais complexa (geralmente a parte de solicitação R) uma quantidade menor de vezes do que o total de chamadas de $g(C, R)$. Com isso, há um benefício no esforço computacional dependendo do custo exigido pela parte mais complexa, dado que o número de vezes que ela é calculada é reduzido para N . Apesar do número total de avaliações da função de falha aumentar para $M \times N$ em comparação com o MC realizado com N simulações, as partes da função de falha envolvidas são calculadas previamente, de modo que a função de falha apenas chama valores já obtidos.

2.2.3.1 Características do estimador da probabilidade de falha do SepMC

É possível notar que, ao separar a função de falha em C e R , a probabilidade de falha se torna igual à esperança da CDF da parte de resistência em relação à de solicitação, ou seja, igual à probabilidade de a resistência ser menor que a solicitação. Em suma, isso pode ser escrito como $P_f = E_R(F_C(R))$. Se for construída uma CDF experimental $\hat{F}_C(c)$ de C a partir de uma amostra $C_i, i = 1, 2, \dots, M$, então a Equação 2.37 pode ser reescrita como

$$\tilde{P}_f = \frac{1}{N} \sum_{j=1}^N \hat{F}_C(R_j) \quad (2.38)$$

A esperança e a variância dessa estimativa são dadas por (SMARSLOK *et al.*, 2010)

$$E(\tilde{P}_f) = E(\hat{F}_C(R)) = P_f \quad (2.39)$$

$$Var(\tilde{P}_f) = \frac{1}{N} \left[\frac{1}{M} \xi_{R,R} + \theta \right] + \frac{(N-1)}{N} \left[\frac{1}{M} \xi_{R_1,R_2} \right] \quad (2.40)$$

onde

$$\xi_{R,R} = P_f - E(F_C(R)^2) \quad (2.41)$$

$$\theta = E(F_C(R)^2) - P_f^2 \quad (2.42)$$

$$\xi_{R_1,R_2} = E(F_C(\min(R_1, R_2))) - P_f^2 \quad (2.43)$$

Os termos $E(F_C(R)^2)$ e $E(F_C(\min(R_1, R_2)))$ são estimados pelas seguintes equações:

$$E(F_C(R)^2) = \frac{1}{N} \sum_{j=1}^N \left[\sum_{i=1}^M \frac{I(g(C_i, R_j) < 0)}{M} \right]^2 \quad (2.44)$$

$$E(F_C(\min(R_1, R_2))) = \frac{2}{N} \sum_{j=1}^{N/2} \sum_{i=1}^M \frac{I(g(C_i, \min(R_{2j+1}, R_{2j})) < 0)}{M} \quad (2.45)$$

2.2.3.2 Algoritmo do SepMC

O SepMC é muito similar ao MC tradicional, de modo que seu algoritmo é idêntico na maior parte das etapas. Uma vez que a função de falha tenha sido dividida em uma parte R e uma parte C , podemos resumir o passo a passo do método em:

- 1 – Gere N pontos amostrais r_k e M pontos amostrais c_l a partir das PDFs marginais de cada VA;
- 2 – Calcule a função de falha $g(r_k, c_l)$ para cada valor de $k = 1, 2, \dots, N$ e $l = 1, 2, \dots, M$;
- 3 – Avalie a função indicadora $I(r_k, c_l)$ para cada resultado de $g(r_k, c_l)$, contando os pontos nos quais $g(r_k, c_l) \leq 0$, ou seja, determinando o NFP ;
- 4 – Calcule a $\tilde{P}_f$ pela Equação 2.37.

2.2.4 Subset Simulation (SuS)

O SuS parte do princípio de que, se o domínio do problema for dividido em partes delimitadas por valores da função de falha $G = g(\mathbf{X})$, pode-se calcular a probabilidade de falha em cada parte para obter a probabilidade de falha total. Definindo a falha como $F = \{g(\mathbf{X}) < b\}$, são adotados níveis limites $b_1 > b_2 > b_3 > \dots > b_{m-1} > b_m = b$, de modo que o domínio fica dividido em partes cujos limites não se atravessam e a probabilidade de falha pode ser calculada como (AU e BECK, 2001)

$$P(G < b_m) = P((G < b_m) \cap (G < b_{m-1})) = P(G < b_m | G < b_{m-1})P(G < b_{m-1}) \quad (2.46)$$

Isso é possível pela característica que impomos aos níveis limites, de modo que a partição limitada por b_m está contida na partição limitada por b_{m-1} , levando a intersecção entre os eventos $G < b_m$ e $G < b_{m-1}$ a ser igual ao próprio evento $G < b_m$.

Se esse raciocínio for seguido e aplicado a $P(G < b_{m-1})$ na Equação 2.46, é obtida a equação

$$P(G < b) = P(G < b_m | G < b_{m-1})P(G < b_{m-1} | G < b_{m-2})P(G < b_{m-2}) \quad (2.47)$$

Fica claro, então, que essa ideia pode ser estendida a todos os níveis limites e é possível calcular a probabilidade de falha através de um produto de probabilidades condicionais:

$$P(G < b) = P(G < b_1) \prod_{i=2}^m P(G < b_i | G < b_{i-1}) \quad (2.48)$$

As probabilidades condicionais apresentadas podem ser encontradas através das suas PDFs condicionais, que são desconhecidas à priori. Se os eventos $G < b_i$ e $G < b_{i-1}$ forem nomeados como F_i e F_{i-1} , respectivamente, então

$$P(G(\mathbf{X}) < b_i | G(\mathbf{X}) < b_{i-1}) = P(F_i | F_{i-1}) = \int I(\mathbf{X} \in F_i) f(\mathbf{x} | F_{i-1}) d\mathbf{x} \quad (2.49)$$

onde $I(\mathbf{X} \in F_i)$ é a função indicadora cujo valor é 1 se $\mathbf{X} \in F_i$ e 0 no caso contrário e $f(\mathbf{x} | F_{i-1})$ é a PDF condicional de $\mathbf{X}$ dado que F_{i-1} aconteceu, dada por

$$f(\mathbf{x} | F_{i-1}) = \frac{I(\mathbf{X} \in F_{i-1}) f(\mathbf{x})}{P(F_{i-1})} \quad (2.50)$$

sendo $f(\mathbf{x})$ a PDF conjunta das variáveis $\mathbf{X}$.

A PDF condicional $f(\mathbf{x} | F_{i-1})$ representa a PDF $f(\mathbf{x})$ confinada no domínio de falha pela função indicadora e dividida pela área total do domínio de falha ($P(F_{i-1})$), por ser condicional a esse evento, de modo que sua integração ao longo do domínio de falha deve resultar em 1. É possível estimar o valor de $P(G(\mathbf{X}) < b_i | G(\mathbf{X}) < b_{i-1})$ através de uma amostra de tamanho N , da mesma forma que o MC tradicional:

$$P(F_i | F_{i-1}) \approx \tilde{P}(F_i | F_{i-1}) = \frac{1}{N} \sum_{j=1}^N I(\mathbf{X}_j \in F_i), \quad \mathbf{X} \sim f(\mathbf{x} | F_{i-1}) \quad (2.51)$$

Contanto que a amostra tenha sido obtida através de $f(\mathbf{x} | F_{i-1})$, a Equação 2.51 é válida (ZUEV, 2013). Essa PDF condicional é desconhecida, de modo que gerar uma amostra a partir dela diretamente é inviável, sendo necessária a utilização do *Markov Chain Monte Carlo* (MCMC). O MCMC não gera uma amostra com elementos independentes entre si, mas é capaz de dar uma boa estimativa para $P(F_i | F_{i-1})$ que tende a não ter viés à medida que o tamanho amostral aumenta (AU e WANG, 2014).

Nota-se que, para realizar essa metodologia, devemos fixar ou os níveis limites ou a probabilidade de falha desejada para cada nível. A segunda opção foi escolhida, sendo fixada a probabilidade $p_0 = 0,1$, ou seja, em cada nível exatamente 10% dos pontos amostrais estarão na falha. Isso é repetido até se chegar ao nível desejado, b_m .

2.2.4.1 Markov Chain Monte Carlo (MCMC)

O MCMC é um método iterativo que permite gerar amostras a partir de uma PDF desconhecida (AU e WANG, 2014). Para tal, é necessário ter uma função $q(\mathbf{x})$ proporcional a essa PDF alvo, que será chamada de $f_\pi(\mathbf{x})$. Essa condição é escrita como

$$f_{\pi}(\mathbf{x}) = \frac{q(\mathbf{x})}{\int_D q(\mathbf{x}) d\mathbf{x}} \quad (2.52)$$

onde D é o domínio de $\mathbf{X}$ e $\int_D q(\mathbf{x}) d\mathbf{x}$ é uma constante de proporcionalidade.

De maneira geral, a constante presente na Equação 2.52 é de difícil determinação, por ser uma integral multidimensional. O objetivo do MCMC é, portanto, gerar amostras de $f_{\pi}(\mathbf{x})$ a partir de $q(\mathbf{x})$, sem conhecer a constante de proporcionalidade entre essas funções. A amostra é construída partindo de um ponto inicial $\mathbf{X}_1$ que é usado para gerar o ponto seguinte ($\mathbf{X}_2$) e assim sucessivamente. No fim, uma amostra $\{\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_N\}$ é obtida, onde $\mathbf{X}_{k+1}$ é sempre gerado partindo de $\mathbf{X}_k$, $k = 1, 2, \dots, N$. Cada amostra construída dessa forma é chamada de *Markov chain*, sendo o ponto inicial chamado de semente.

No caso do SuS, a PDF desconhecida a partir da qual se quer gerar amostras é uma PDF condicional na forma presente na Equação 2.50, onde $i = 2, 3, \dots, m$. Fazendo a comparação com a Equação 2.52, é fácil ver que $f_{\pi}(\mathbf{x}) = f(\mathbf{x}|F_{i-1})$ e $q(\mathbf{x}) = I(\mathbf{X} \in F_{i-1})f(\mathbf{x})$ nessa situação.

Para que o MCMC funcione corretamente, a *Markov chain* deve satisfazer duas condições fundamentais: a equação conhecida como *detailed balance* e a teoria ergódica. Enquanto a primeira condição é sempre satisfeita pelo algoritmo do MCMC, a segunda depende de características do problema e de implementação do algoritmo, sendo necessário testar se está sendo satisfeita em cada caso.

Em suma, *detailed balance* (Equação 2.53) representa uma reversibilidade no processo, ou seja, a *Markov chain* pode ser gerada de trás para frente da mesma forma que de frente para trás. Ela é definida como

$$f_{X_{k+1}|X_k}(\mathbf{x}_{k+1}|\mathbf{x}_k)f_{\pi}(\mathbf{x}_k) = f_{X_{k+1}|X_k}(\mathbf{x}_k|\mathbf{x}_{k+1})f_{\pi}(\mathbf{x}_{k+1}) \quad (2.53)$$

onde $f_{X_{k+1}|X_k}(\mathbf{X}_{k+1}|\mathbf{X}_k)$ é chamada de PDF de transição e $\mathbf{x}_k$ é o k -ésimo elemento ($k = 1, 2, \dots, N$) da *Markov chain*.

Se a Equação 2.53 for satisfeita, a *Markov chain* terá como PDF estacionária a PDF alvo, da qual se deseja amostrar. Em outras palavras, dado que $\mathbf{X}_k$ é distribuído de acordo com $f_{\pi}(\mathbf{x})$, todos os elementos seguintes da amostra, de $k + 1$ até N , também serão distribuídos de acordo com $f_{\pi}(\mathbf{x})$. Essa propriedade é fundamental para que a amostra obedeça a PDF alvo mesmo ela sendo desconhecida.

Já a ergodicidade pode ser resumida como a capacidade da *Markov chain* explorar todo o domínio D sem ficar presa em nenhuma região específica. Pode ser demonstrado que, se a

teoria ergódica é satisfeita no problema em questão, a *Markov chain* vai tender a ser distribuída conforme $f_\pi(\mathbf{x})$ mesmo que $\mathbf{X}_1$ não o seja. Essa propriedade, descrita na Equação 2.54, é fundamental para que a amostra obedeça a PDF alvo independente do ponto inicial. Para mais detalhes desse conceito, ver Au e Wang (2014).

$$\lim_{k \rightarrow \infty} f_{\mathbf{X}_k}(\mathbf{x}) = f_\pi(\mathbf{x}) \quad (2.54)$$

Em adição a esse fato, a ergodicidade também é o que garante que a amostra gerada pelo MCMC vai tender a não ter viés mesmo não tendo elementos $\mathbf{X}_k$ independentes e identicamente distribuídos (i.i.d.). Essa propriedade é essencial para que a amostra possa ser usada para estimar uma probabilidade de falha, seja ela condicional ou não. A Equação 2.55 mostra essa propriedade para o caso em que $\mathbf{X}_1$ tem PDF igual a $f_\pi(\mathbf{x})$ e a Equação 2.56 mostra o caso em que $\mathbf{X}_1$ tem uma PDF qualquer.

$$E(\tilde{P}_f) = P_f \quad (2.55)$$

$$\lim_{N \rightarrow \infty} E(\tilde{P}_f) = P_f \quad (2.56)$$

Uma vez que essas propriedades são válidas, podemos aplicar o algoritmo do MCMC. Um algoritmo comum é o Metropolis-Hastings (MH) (HASTINGS, 1970), criado a partir do mais antigo algoritmo de MCMC, o algoritmo de Metropolis (METROPOLIS *et al.*, 1953). O MH é capaz de gerar um ponto $\mathbf{x}_{k+1}$ partindo de um ponto $\mathbf{x}_k$ usando os seguintes passos:

- 1 – Gere um valor u a partir da VA U distribuída uniformemente no intervalo $[0, 1]$ e gere um candidato $\mathbf{x}'$ a partir de uma PDF proposta $p^*(\mathbf{X}; \mathbf{x}_k)$;
- 2 – Calcule a taxa de crescimento ou decrescimento de $f_\pi(\mathbf{x})$, ponderada pela PDF proposta (Equação 2.57);

$$r = \frac{p^*(\mathbf{x}_k; \mathbf{x}')q(\mathbf{x}')}{p^*(\mathbf{x}'; \mathbf{x}_k)q(\mathbf{x}_k)} \quad (2.57)$$

- 3 – Defina o valor do próximo ponto ($\mathbf{x}_{k+1}$) seguindo a Equação 2.58.

$$\mathbf{x}_{k+1} = \begin{cases} \mathbf{x}', & u < r \\ \mathbf{x}_k, & u \geq r \end{cases} \quad (2.58)$$

onde a PDF proposta $p^*(\mathbf{X}; \mathbf{v})$ pode ser uma PDF qualquer desde que, para um dado $\mathbf{v}$, ela seja válida para qualquer valor de $\mathbf{X}$ e possa gerar uma amostra com elementos i.i.d. A variável $\mathbf{v}$, aqui, representa o ponto no qual $p^*(\mathbf{X}; \mathbf{v})$ foi centrada.

Esse algoritmo é repetido para cada valor de $k = 1, 2, \dots, N$ até formar a *Markov chain* $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$, que será uma amostra distribuída de acordo com a PDF alvo. Os primeiros elementos da amostra podem ser retirados, caso $\mathbf{X}_1$ não seja distribuído conforme $f_\pi(\mathbf{x})$, de modo a reduzir a distorção causada pelos elementos que não atingiram a PDF estacionária. Essa retirada é usualmente chamada de *burn-in* (AU e WANG, 2014).

É importante salientar que a qualidade da probabilidade de falha estimada pelo SuS depende das amostras geradas pelo MCMC, sendo necessário que o algoritmo funcione corretamente para o problema em questão. Além do Metropolis e do MH, outros algoritmos foram desenvolvidos à medida que a aplicação do MCMC aumentou e problemas mais complexos exigiram correções, como o *Conditional Sampling* (CS) e o *Adaptive Conditional Sampling* (ACS) (PAPAIOANNOU *et al.*, 2015).

2.2.4.2 Algoritmo do SuS

As informações de entrada para iniciar o SuS são a função de falha $g(\mathbf{X})$, o tamanho N_n da amostra de cada nível e a porcentagem de pontos falhos em cada nível p_0 . Sendo m o número de níveis, o algoritmo do SuS pode ser resumido nos seguintes passos (ZHANG e HESTHAVEN, 2020):

- 1 – Defina $f(\mathbf{x}|F_0) = f(\mathbf{x})$;
- 2 – Enquanto $g_i > 0$, onde i representa o número do nível:
 - 2.1 – Faça $i = i + 1$;
 - 2.2 – Gere N_n pontos $\mathbf{x}_k, k = 1, 2, \dots, N_n$, a partir de $f(\mathbf{x}|F_{i-1})$ usando o MCMC (na iteração $i = 1$ é usado o MC tradicional) com a semente de cada *Markov chain* sendo um dos pontos que falharam na iteração anterior;
 - 2.3 – Avalie $g(\mathbf{x}_k)$ para todos os pontos da amostra e coloque os pontos em uma ordem decrescente;
 - 2.4 – Defina g_i como o menor valor de $g(\mathbf{x})$ que ficou dentro dos p_0 mais próximos da falha (por exemplo, dos 10% mais próximos da falha);

2.5 – Defina $F_i = \{\mathbf{x} | g(\mathbf{x}) \leq g_i\}$;

3 – Contabilize o número de pontos que satisfazem $g(\mathbf{x}) < 0$, ou seja, calcule o NFP ;

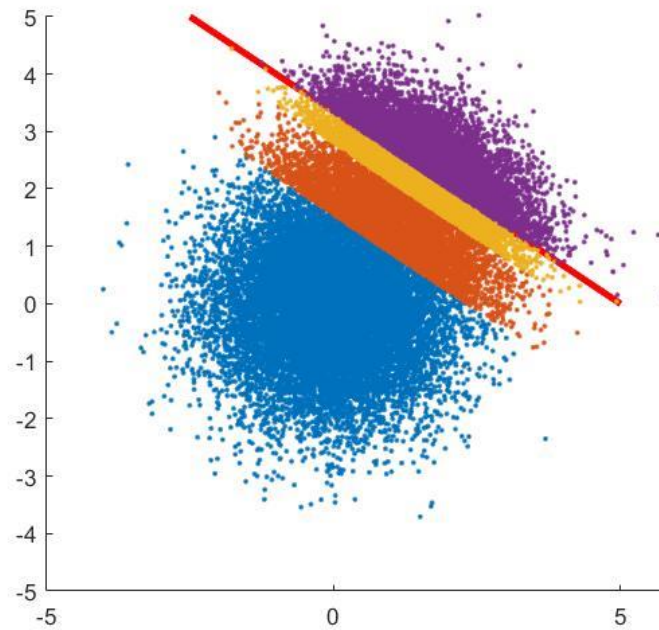
4 – Por fim, calcule $P_m = P(G < b_m | G < b_{m-1}) = NFP/N_n$ e obtenha a probabilidade de falha pelo produto de todas as probabilidades de falha condicionais (fixadas como p_0):

$$\tilde{P}_f = p_0^{m-1} P_m \quad (2.59)$$

Nesse caso, a facilidade de fixar $P(G < b_i | G < b_{i-1}) = p_0$, $i = 2, 3, \dots, m - 1$, é evidente na equação da probabilidade de falha estimada (Equação 2.59). É importante perceber também que são geradas $p_0 \times N_n$ *Markov chains* por amostra de tamanho N_n , sendo o tamanho de cada uma igual a p_0^{-1} (AU e WANG, 2014).

Para ilustrar, a Figura 7 mostra o resultado do SuS aplicado à função de falha $g(\mathbf{x}) = 10 - 2x_1 - 3x_2$, iniciando a amostra no nível zero (cor azul) até gerar pontos dentro do domínio de falha (cor roxa), sendo a linha vermelha o limite do domínio de falha. É possível ver claramente os pontos gerados em cada nível e como a amostra é levada até a região de falha.

Figura 7 – Exemplo de amostra gerada pelo SuS.



Fonte: O autor (2022).

3 MÉTODO PROPOSTO: MONTE CARLO SELETIVO (SMC)

O presente trabalho tem como objetivo propor um novo método de simulação para reduzir o número de avaliações da função de falha exigidos pelo MC, dado que o esforço computacional é o principal problema desse tipo de método. Considerando que em problemas complexos, como estruturas com muitos graus de liberdade calculadas por métodos numéricos, o número total de avaliações tem um custo que pode inviabilizar a resolução do problema, esse método traz uma solução alternativa aos métodos de redução de variância para essa questão.

O método proposto, chamado Monte Carlo Seletivo (SMC), tem como principal objetivo desconsiderar, no cálculo, os pontos que sabidamente não estão na falha. Em problemas em que a probabilidade de falha é pequena, como é comum na engenharia estrutural, uma parcela muito reduzida dos pontos cai no domínio de falha. Por conseguinte, quando o SMC consegue desconsiderar a maioria dos pontos que estão no domínio seguro, o *NFE* pode se aproximar do *NFP*, reduzindo o esforço computacional em relação ao MC tradicional.

3.1 DEFINIÇÃO DO PROBLEMA

Usualmente, nos problemas de engenharia estrutural, o sentido do eixo (no espaço n -dimensional) de cada variável X_i , $i = 1, 2, \dots, n$, no qual a falha ocorre é previamente conhecido. As VAs associadas às propriedades dos materiais, como tensão de escoamento, módulo de elasticidade, resistência à compressão do concreto e dimensões da seção, tem o domínio de falha no sentido negativo do seu eixo a partir da média, enquanto as VAs associadas às solicitações, como carregamentos verticais e horizontais, tem o domínio de falha no sentido positivo do seu eixo a partir da média. Quando esse sentido não é conhecido, ele geralmente pode ser identificado de maneira simples, através do cálculo do gradiente na média ou observando o comportamento da variável.

Com base nisso, podemos usar essa informação para buscar o domínio de falha e avaliar somente os pontos da amostra localizados nele, evitando chamadas desnecessárias de $g(\mathbf{X})$. Isso é possível quando a função de falha é monotônica, de modo que podemos saber se um ponto é o pior caso do outro apenas observando os valores de $\mathbf{X}$. Nessa situação, para encontrar o domínio de falha, basta buscar os piores casos de cada variável X_i , *i.e.*, os valores mais extremos no sentido de falha da respectiva VA. Isso se mantém verdade tanto para funções lineares quanto não lineares, desde que a monotonicidade da função seja satisfeita.

Para os fins desse trabalho, uma função multivariada $f(x_1, x_2, \dots, x_n)$ qualquer é considerada monotônica não-decrescente se, para quaisquer dois pontos $(x_1, x_2, \dots, x_n)$ e $(x_1^*, x_2^*, \dots, x_n^*)$ nos quais $x_i \leq x_i^*$, $i = 1, 2, \dots, n$, então $f(x_1, x_2, \dots, x_n) \leq f(x_1^*, x_2^*, \dots, x_n^*)$. Essa ideia é estendida para funções monotônicas não-crescentes, nas quais $f(x_1, x_2, \dots, x_n) \geq f(x_1^*, x_2^*, \dots, x_n^*)$ para quaisquer $x_i \leq x_i^*$, $i = 1, 2, \dots, n$. Fazer a verificação de monotonicidade em todo o domínio é complicada e, por isso, não foi feita nesse trabalho. Contudo, esse comportamento pode ser observado de forma aproximada através de uma amostra (por exemplo, com o uso do MC) em uma determinada região do domínio, o que foi adotado aqui para classificar as funções.

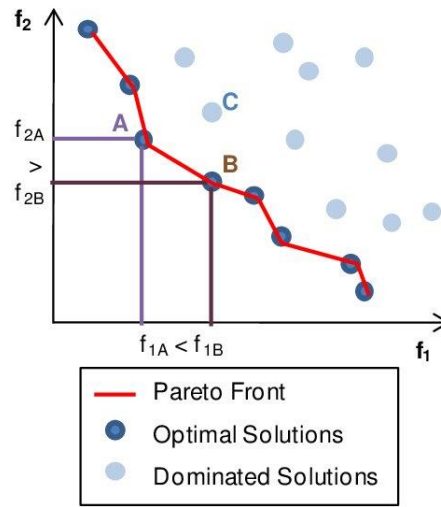
Os problemas alvos do método proposto, portanto, são aqueles com funções de falha monotônicas e baixa probabilidade de falha, que representam a maior parte dos problemas de engenharia estrutural. O benefício pode ser observado pelo número de avaliações de $g(X)$ exigido em comparação com o MC tradicional e outros métodos.

3.2 OTIMALIDADE DE PARETO

Para encontrar os pontos da amostra que representam os piores casos, o conceito de otimalidade de Pareto (MESSAC e MULLUR, 2007) pode ser usado. Esse conceito, que é aplicado em problemas de otimização multiobjetivo, serve para encontrar pontos cujo movimento no sentido de melhora de uma das funções gera necessariamente piora de outra função. A região que contém esses pontos é chamada de fronteira de Pareto, onde os pontos presentes nela são pontos ótimos em relação aos demais pontos do domínio. Um ponto que pertence a essa fronteira é dito não-dominado, enquanto os pontos que representam uma piora dele para todas as funções são ditos dominados pelo mesmo. Vale salientar que cada ponto dominante tem seus pontos dominados, que podem ou não coincidir com os de outro ponto dominante.

Um exemplo da fronteira de Pareto para otimização de duas funções objetivo é mostrado na Figura 8. Conforme mostrado pelos pontos A e B, considerando dois pontos na fronteira de Pareto, um ponto pode ser melhor que outro para uma das funções, mas não para as duas funções ao mesmo tempo. Já para os pontos dominados, como o ponto C, existe necessariamente outro ponto, na fronteira de Pareto ou não, que é melhor que ele para todas as funções.

Figura 8 – Exemplo de fronteira de Pareto.



Fonte: Villa e Labayrade (2011).

Considerando que temos uma amostra de tamanho N , a ideia de ótimo de Pareto é usada aqui para otimizar as diferentes VAs do problema ao mesmo tempo, onde otimizar, nesse contexto, significa buscar o pior caso. Em outras palavras, o objetivo ao buscar a fronteira de Pareto é encontrar os pontos mais próximos da falha do que todos os demais da amostra. Formalmente, podemos escrever que um determinado ponto $\mathbf{x}_k$ da amostra faz parte da fronteira de Pareto se não existe nenhum outro ponto $\mathbf{x}_j$, $j = 1, \dots, k-1, k+1, \dots, N$, tal que:

$$x_{j,i} \leq x_{k,i}, i = 1, 2, \dots, n \quad (3.1)$$

$$x_{j,i} < x_{k,i}, \text{ para ao menos um valor de } i \quad (3.2)$$

onde n é a quantidade de VAs do problema.

Um ponto $\mathbf{x}_j$ qualquer domina outro ponto $\mathbf{x}_k$ se as equações 3.1 e 3.2 forem satisfeitas, sendo a notação que representa essa relação entre os dois pontos escrita como $\mathbf{x}_j > \mathbf{x}_k$.

3.3 CONCEITO DO MÉTODO

Da mesma forma que o MC, o SMC gera uma amostra de tamanho N a partir das distribuições marginais das VAs do vetor $\mathbf{X}$. Porém, o SMC usa informações prévias sobre a localização dos pontos para chamar $g(\mathbf{X})$ um número de vezes próximo do NFP , que é bem menor do que N . Essa metodologia é possível quando $g(\mathbf{X})$ é monotônica, de modo que podemos identificar quais os pontos mais extremos no sentido da falha simplesmente observando os valores de cada VA dos pontos da amostra.

O conceito de otimalidade de Pareto é aplicado para determinar os pontos mais extremos da amostra, *i.e.*, a fronteira de Pareto, onde as VAs são tratadas como as funções objetivo a serem otimizadas. A fronteira de Pareto é formada pelos pontos que otimizam os valores de X_i , $i = 1, 2, \dots, n$, mas o sentido da falha pode ser de aumento para algumas variáveis e redução para outras, a partir da média. Para uniformizar o sentido do valor ótimo, os sinais daquelas cuja falha fica no sentido negativo do eixo são invertidos no algoritmo de filtragem da fronteira de Pareto. Como já mencionado, o sentido de falha de cada VA pode ser obtido por meio do gradiente na média ou observando o comportamento da mesma.

Uma vez encontrando a fronteira de Pareto, a função de falha é avaliada para os pontos contidos nela e a quantidade deles no domínio de falha é contabilizada. Levando em conta que os pontos dominados são casos melhores que seu ponto dominante para todas as VAs, é possível concluir que, se um ponto não-dominado é seguro, todos os pontos dominados por ele também são seguros. Essa é a ideia principal do SMC, que é usada para desconsiderar os pontos que são seguros. A partir disso, sobra uma amostra reduzida, eliminando da amostra original os pontos na falha que foram contados e os pontos que já se sabe que são seguros. O processo é então repetido para essa nova amostra, encontrando uma nova fronteira de Pareto, adicionando a quantidade de pontos falhos à encontrada anteriormente e eliminando da amostra os pontos seguros e os pontos falhos. Isso continua até que todos os pontos tenham sido eliminados da amostra e o *NFP* tenha sido totalmente contado.

Resumindo, o SMC é um processo iterativo que vai encontrando os pontos falhos progressivamente até ter o *NFP* da amostra original e, conseqüentemente, a probabilidade de falha do problema. Se a amostra for a mesma que foi usada no MC, a probabilidade de falha obtida via SMC é igual à obtida pelo MC, desde que as funções de falha sejam monotônicas na região de interesse. Vale salientar que o SMC também pode ser usado para problemas com função de falha não monotônica, mas o resultado se torna uma apenas uma aproximação do valor real do *NFP*, cuja qualidade depende do problema e da precisão desejada.

3.3.1 Algoritmo do SMC

Lembrando que n é o número de VAs do problema e chamando a fronteira de Pareto de $\mathcal{F}$, o algoritmo do SMC é descrito a seguir:

- 1 – Defina o sentido de falha, podendo ser obtido através do gradiente da função de falha (sentido oposto) ou de observação do comportamento das VAs;

2 – Gere uma amostra $\mathbf{S}$, que é uma matriz $N \times n$ onde as linhas são os pontos $\mathbf{x}_k$, $k = 1, 2, \dots, N$;

3 – Faça $NFP = 0$;

4 – Enquanto $\mathbf{S}$ não for uma matriz vazia:

4.1 – Determine os pontos dominantes da amostra:

$$\mathbf{S}_p = \{\mathbf{x}_k \in \mathbf{S} | \mathbf{x}_k \in \mathcal{F}\} \quad (3.3)$$

4.2 – Calcule a função de falha nos pontos dominantes para determinar quantos falharam:

$$NFP = NFP + \text{sum}(g(\mathbf{S}_p) \leq 0) \quad (3.4)$$

4.3 – Encontre os pontos seguros da fronteira de Pareto:

$$\mathbf{S}_{ps} = \{\mathbf{x}_k \in \mathbf{S}_p | g(\mathbf{x}_k) > 0\} \quad (3.5)$$

4.4 – Determine todos os pontos dominados por $\mathbf{S}_{ps}$:

$$\mathbf{S}_s = \{\mathbf{x}_k \in \mathbf{S} | \mathbf{S}_{ps} \succ \mathbf{x}_k\} \quad (3.6)$$

4.5 – Retire $\mathbf{S}_p$ e $\mathbf{S}_s$ de $\mathbf{S}$;

5 – Calcule a probabilidade de falha, dividindo os pontos falhos calculados pelo tamanho original da amostra:

$$\tilde{P}_f = NFP/N \quad (3.7)$$

A determinação dos pontos dominantes e dominados (passos 4.1 e 4.4) pode ser feita por qualquer algoritmo que encontre a fronteira de Pareto. Nesse trabalho foi usado o algoritmo apresentado a seguir, onde o sentido de falha de cada VA deve ser previamente conhecido. No fim do processo, são obtidas duas listas contendo os índices dos pontos dominantes ($p_{dominant}$) e dominados ($p_{dominated}$) dentro da amostra.

1 – Crie uma lista com os índices de cada ponto na matriz $\mathbf{S}$, chamada de $p_{dominant}$;

2 – Inverta, em $\mathbf{S}$, os sinais das VAs cuja falha é no sentido negativo do seu respectivo eixo;

3 – Enquanto houverem pontos dominados por pelo menos um outro ponto:

3.1 – Adote o último ponto da amostra como o ponto de referência;

3.2 – Identifique quais pontos são dominados pelo ponto de referência, ou seja, quais pontos tem valores menores ou iguais para todas as variáveis;

3.3 – Elimine os pontos dominados pelo ponto de referência em $\mathbf{S}$ e elimine os índices desses pontos em $p_{dominant}$;

3.4 – Adote o próximo ponto de referência como o ponto imediatamente acima do ponto de referência atual;

4 – Para cada ponto dominante, cujo índice dentro do $\mathbf{S}$ original está em $p_{dominant}$, crie um vetor com os índices no $\mathbf{S}$ original dos pontos dominados por ele, ou seja, com valores menores ou iguais para todas as variáveis;

5 – Coloque os vetores obtidos no passo 4 em células e reúna essas células em uma lista nomeada como $p_{dominated}$.

Fica claro que o sucesso do SMC depende do sentido de falha definido, tanto na busca da fronteira de Pareto quanto na redução do NFE . Assim, é fundamental que haja certeza nessa etapa do método, independente da forma escolhida para determinar esse sentido.

3.3.2 Exemplo do método

Para melhor visualização do método proposto, o algoritmo descrito será exemplificado usando uma função de falha simples com duas variáveis (Equação 3.8), ambas possuindo distribuição normal padrão. Definindo o tamanho da amostra como $N = 10^4$, a matriz $\mathbf{S}$ é gerada com dimensões $[10^4, 2]$. A partir dessas informações, podemos seguir com o algoritmo descrito na subseção 3.3.1 e buscar a fronteira de Pareto.

$$g(\mathbf{X}) = 10 - 2X_1 - 3X_2 \quad (3.8)$$

Uma vez que os pontos dominantes da amostra original são encontrados, a função de falha é avaliada para cada um e eles são separados entre pontos falhos e seguros. Esses pontos, para a primeira iteração, são mostrados na Figura 9, onde os que falharam foram representados por cruzeiros vermelhas e os seguros por círculos verdes. A amostra $\mathbf{S}$ é então reduzida, removendo os pontos dominantes e os pontos dominados pelos pontos dominantes seguros (simbolizados por “x”). O resultado é uma nova amostra $\mathbf{S}$ composta pelos pontos restantes, que estão em azul na Figura 9.

O processo é então repetido, encontrando uma nova fronteira de Pareto para a amostra reduzida e eliminando novamente os pontos dominantes e os pontos dominados por pontos seguros. Esse processo continua até que $\mathbf{S}$ seja uma matriz vazia, o que acontece na 4ª iteração. As iterações 2 a 4 são mostradas na Figura 10, onde é possível ver a fronteira de Pareto de cada

uma, os pontos dominados que foram retirados e o tamanho da amostra sendo reduzido. Ao final, todos os pontos foram avaliados ou eliminados.

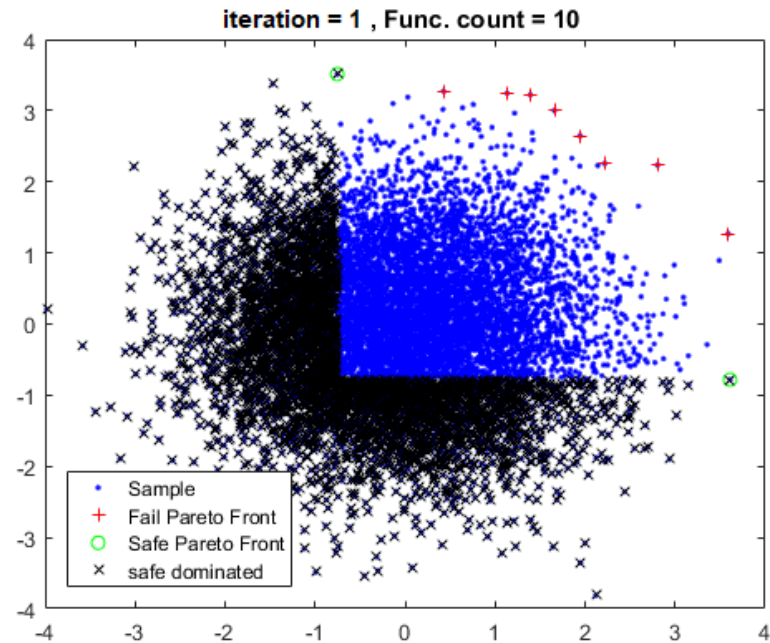
A cada iteração, o número de pontos dominantes falhos é contado e somado ao obtido na iteração anterior, de modo que o NFP foi obtido sem a necessidade de avaliar a amostra inteira. Nesse exemplo, o NFP foi 23 e, portanto, a $\tilde{P}_f$ é igual a 0,23%. Esse é exatamente o resultado obtido pelo MC tradicional, mas com apenas 44 avaliações de $g(\mathbf{X})$. O histórico das iterações desse exemplo pode ser visto na Tabela 2.

Tabela 2 – Resultados de cada iteração do SMC para $N = 10^4$.

Iteração	Pontos na fronteira de Pareto	Pontos de Pareto falhos	Pontos de Pareto seguros	Tamanho da amostra reduzida
1	10	8	2	6009
2	11	8	3	214
3	14	7	7	38
4	9	0	9	0
Total	44	23	21	-

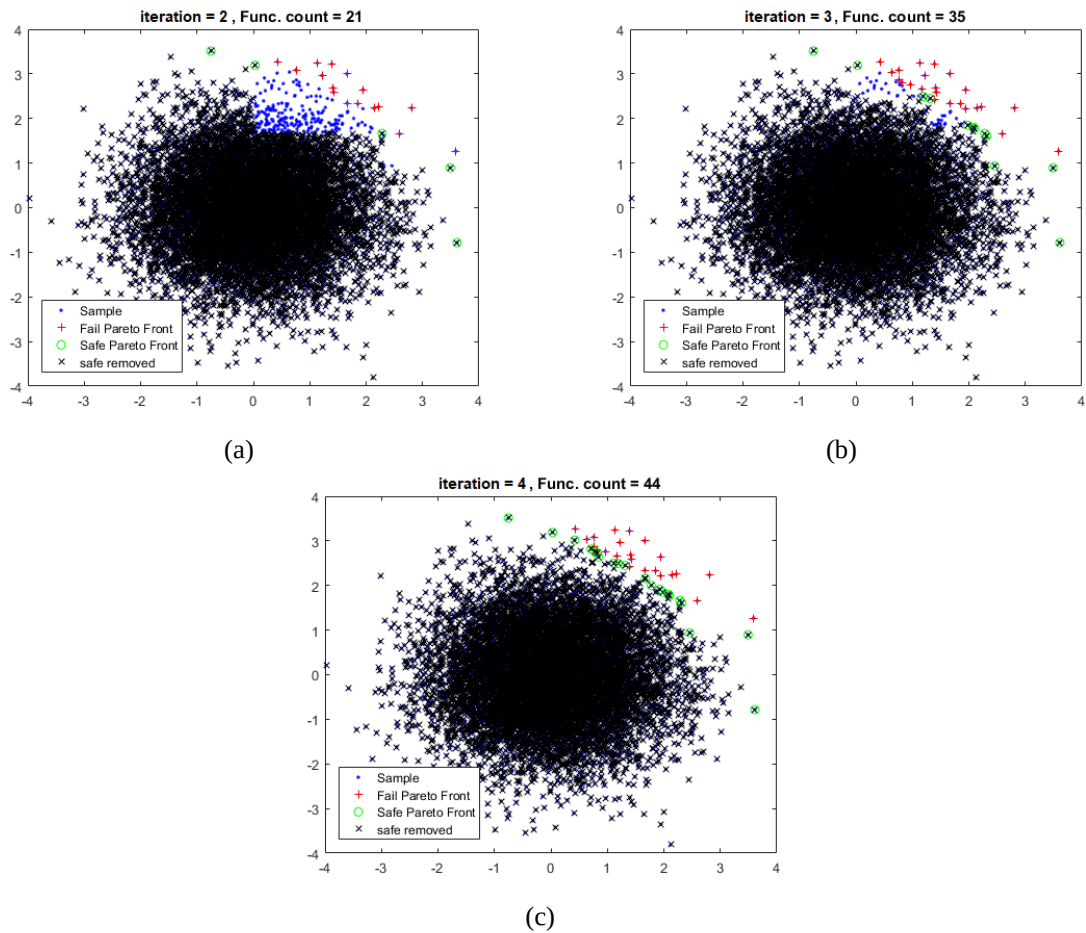
Fonte: O autor (2022).

Figura 9 – Primeira iteração do SMC.



Fonte: O autor (2022).

Figura 10 – Iterações do SMC: (a) iteração 2, (b) iteração 3 e (c) iteração 4.



Fonte: O autor (2022).

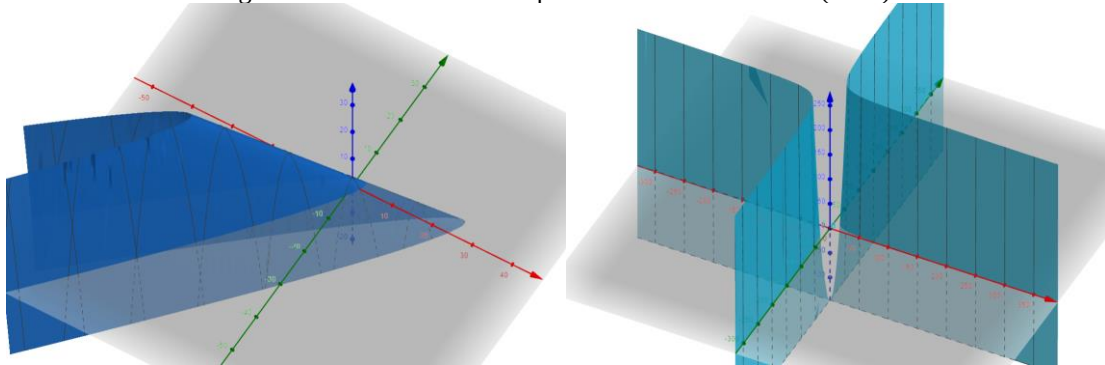
3.4 LIMITAÇÕES

Como já mencionado, se a função de falha for localmente monotônica na região de interesse, *i.e.*, na região com maior concentração de densidade de probabilidade no domínio de falha, o SMC encontra o mesmo valor de NFP que o MC tradicional para uma mesma amostra. Por consequência, os dois métodos obtêm a mesma probabilidade de falha nesse caso. O SMC funciona tanto em funções de falha lineares quanto não lineares, desde que a condição de monotonicidade seja satisfeita.

A principal limitação do SMC é, portanto, que o comportamento da função de falha deve ser tal que permita haver um caminho claro da média para o domínio de falha sem oscilações. Em especial, o SMC tem dificuldade com problemas cuja região de falha é estreita ou se apresenta em mais de um sentido, casos que são ilustrados pelos exemplos 2 e 11 presentes em Santos *et al.* (2012), respectivamente (Figura 11). Na primeira situação existem pontos seguros mais distantes no sentido de falha que muitos pontos falhos, ou seja, existem pontos

falhos dominados por pontos seguros, o que faz com que muitos pontos falhos sejam eliminados da amostra sem ser contabilizados. Já na segunda situação, o SMC acaba desconsiderando as partes do domínio de falha que estão em sentidos diferentes do sentido de falha estabelecido, o que subestima o *NFP* do problema.

Figura 11 – Problemas 2 e 11 presentes em Santos *et al.* (2012).



Fonte: O autor (2022).

Apesar disso, a maioria dos problemas de engenharia estrutural cumpre a condição de monotonicidade, de modo que o SMC possui vasta utilidade prática. Além disso, como o método trabalha diretamente com a amostra usada no MC, ele não apresenta os problemas de convergência do FORM que podem prejudicar o uso do IS centrado no MPP, os problemas para separar as VAs da função de falha no SepMC e os eventuais problemas do SuS no uso do MCMC quando a ergodicidade não pode ser verificada. Por conseguinte, ele não tem as mesmas limitações de outros métodos de simulação, uma vantagem que pode ser determinante dependendo do problema analisado.

3.5 VALIDAÇÃO E COMPARAÇÃO COM OUTROS MÉTODOS

A validação do SMC foi feita comparando seu resultado com o MC tradicional para o exemplo da subseção 3.3.2 (Equação 3.8), fixando $CoV(\tilde{P}_f) = 5\%$ para que o *NFE* necessário para essa medida de precisão seja mensurado. Como a amostra usada é a mesma para ambos os métodos e a função de falha é monotônica, o *NFP* resultante também deve ser o mesmo. Além dessa validação, também foi feita a comparação com os métodos IS, SuS e SepMC, também fixando $CoV(\tilde{P}_f) = 5\%$ para avaliar a eficiência do SMC na redução do *NFE* em comparação com eles. Para o SepMC, a comparação foi feita para a parte da função de falha com tamanho amostral N , considerando que é a parte mais custosa, enquanto o M foi fixado como 10^6 . Os resultados obtidos para todos os métodos estão descritos na Tabela 3.

Tabela 3 – Resultados da amostra do exemplo.

Método	Amostra	NFE^*	NFE/N	$\tilde{P}_f$	$CoV(\tilde{P}_f)$
MC	$N = 142000$	142000	100%	$2,81 \times 10^{-3}$	5%
IS	$N = 1310$	1322	100%	$2,77 \times 10^{-3}$	5%
SepMC	$N = 3500, M = 10^6$	3500	100%	$2,8 \times 10^{-3}$	5%
SuS	$N = 45780$	45780	100%	$2,72 \times 10^{-3}$	5%
SMC	$N = 142000$	491	0,35%	$2,81 \times 10^{-3}$	5%

*Para o SepMC, o NFE se refere a N .

Fonte: O autor (2022).

Como esperado, a $\tilde{P}_f$ calculada pelo SMC é igual à calculada pelo MC para esse exemplo. Do total de pontos amostrais gerados para o MC ter a precisão desejada, $N = 142000$, o SMC só precisou calcular 491 pontos, o que representa 0,35% da amostra. Isso evidencia a vantagem do método proposto, capaz de fazer o NFE ter um valor próximo do NFP , que nesse caso foi igual a 399. O benefício dessa redução do NFE fica mais evidente ao comparar com os demais métodos.

Enquanto o MC precisou realizar 142000 chamadas da função de falha para atingir $CoV(\tilde{P}_f) = 5\%$, o SuS precisou de 45780, onde houveram três níveis com 15260 pontos em cada um ($b_1 = 5,34$, $b_2 = 1,61$ e $b_3 = 0$). Já o IS precisou de 1322 chamadas, onde esse total inclui as exigidas pelo FORM para encontrar o MPP. O SepMC, para o mesmo $CoV(\tilde{P}_f)$, necessitou de 3500 avaliações da parte mais custosa, que podem ser tratadas como o NFE ao considerar que as M avaliações da parte menos custosa não influenciam de forma significativa no esforço computacional. Como pode ser observado, o SMC obteve o menor NFE dentre todos os métodos considerados, demonstrando sua relevância.

4 APLICAÇÕES

Com o objetivo de avaliar melhor a eficiência e robustez do método proposto, dois problemas de engenharia estrutural e 13 problemas de referência de confiabilidade com funções de falha diversas foram analisados. A meta do SMC é reduzir o número de avaliações necessário em problemas de confiabilidade, então esse foi o principal parâmetro de comparação com outros métodos de simulação. Assim, para todos os métodos, foi calculado o NFE exigido para uma mesma medida de precisão da $\tilde{P}_f$, fixando $CoV(\tilde{P}_f) = 5\%$. Além disso, um estudo paramétrico do número de VAs do problema foi realizado para analisar a influência da dimensão de $\mathbf{X}$ na eficiência do SMC.

Os problemas de engenharia estrutural escolhidos foram uma viga biapoiada presente em Ghali *et al.* (2014) e um pórtico simples presente em Schuëller *et al.* (1989) e Au e Wang (2014). Dentre os demais problemas, 12 foram escolhidos a partir da lista de problemas de confiabilidade reunida por Santos *et al.* (2012) e um foi retirado de Motta *et al.* (2021). Em todos, o erro relativo da $\tilde{P}_f$ em relação à probabilidade de falha de referência foi considerado, para medir a diferença entre a $\tilde{P}_f$ obtida pelo SMC e a obtida pelos demais métodos para os problemas analisados. No total, o SMC foi aplicado em 16 problemas de confiabilidade, visando observar o comportamento do método em situações diversas.

Nesse trabalho, além do algoritmo do SMC desenvolvido, foram usados os algoritmos dos métodos MC e IS presentes no módulo PyRe do Python (HACKL, 2013; BOURINET, 2010; DER KIUREGHIAN *et al.*, 2006) e o algoritmo do SuS desenvolvido pelo *Engineering Risk Analysis Group* da *Technische Universität München* (ERA GROUP, 2018). O algoritmo do SepMC foi implementado no Python pelo autor, com base na descrição do método presente em Smarslok *et al.* (2010).

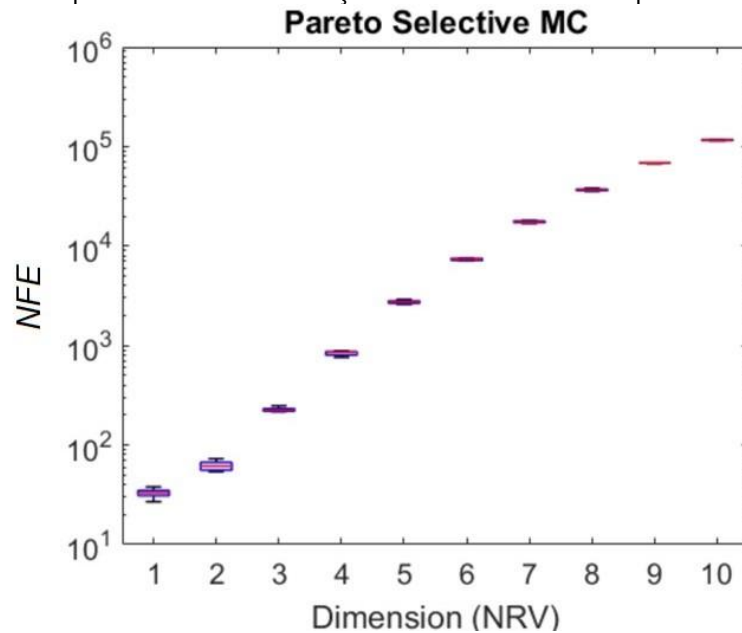
4.1 ESTUDO PARAMÉTRICO DO NÚMERO DE VARIÁVEIS

É importante estudar o comportamento do SMC para problemas com diferentes dimensões, pois o aumento do número de VAs provoca um aumento na proporção de pontos de Pareto em relação ao tamanho total da amostra. Para avaliar a influência do número de VAs na eficiência do SMC, um exemplo cuja função de falha é dada pela Equação 4.1 foi considerado, onde as variáveis são X_i , $i = 1, 2, \dots, n$, todas com distribuição normal padrão. Variando n de 1 a 10 e adotando o tamanho total da amostra como $N = 10^6$, foram realizadas 100 análises de

confiabilidade para cada valor de n . O NFE foi medido para cada uma dessas análises e os resultados são mostrados na Figura 12, onde n foi escrito como NRV . Independentemente do valor de n , $g(\mathbf{X})$ será uma variável normal com média igual a 4 e desvio padrão igual a 1. Assim, a probabilidade de falha exata é $P_f = 3 \times 10^{-5}$ e o número esperado de falhas é 30 ($N = 10^6$), que é o limite esperado de pontos calculados pelo método, alcançado para $n = 1$.

$$g(\mathbf{X}) = 4 - \frac{\sum_{i=1}^n X_i}{\sqrt{n}} \quad (4.1)$$

Figura 12 – Número de pontos avaliados em relação ao número de variáveis para uma amostra com $N = 10^6$.



Fonte: O autor (2022).

Como pode ser observado, para uma maior quantidade de VAs, a taxa de pontos da amostra que são avaliados pela função de falha aumenta. No caso onde a função de falha tem apenas duas variáveis, a quantidade de pontos avaliados ficou abaixo de 0,01% do total de pontos da amostra, enquanto no caso mais extremo, com 10 variáveis, esse valor subiu para pouco mais de 10%. A redução ainda é significativa, já que, mesmo para 10 dimensões, houve uma redução próxima de 90% do NFE em relação ao MC tradicional. Entretanto, é importante considerar essa característica do SMC em todas as aplicações, pois esse nem sempre será o caso e o número de dimensões pode prejudicar a eficiência do método. Em suma, notamos que o SMC é mais eficiente para problemas com menos variáveis envolvidas, mas ainda pode ser vantajoso para problemas com muitas variáveis, sendo importante verificar seu benefício nessas circunstâncias.

4.2 PROBLEMAS DE REFERÊNCIA

Para mensurar melhor a vantagem do método proposto, 13 exemplos apresentados na literatura foram realizados com funções de falha diversas, tanto lineares quanto não lineares. Desses, 12 foram selecionados a partir dos problemas listados em Santos *et al.* (2012) e o último foi escolhido de Motta *et al.* (2021). As funções de falha e as VAs de cada exemplo são descritas na Tabela 4. É importante destacar que as funções de falha dos exemplos II, III, X, XI e XIII não são localmente monotônicas na região de interesse, de modo que o método proposto não necessariamente encontra o *NFP* obtido pelo MC neles. Além do SMC, cada problema foi resolvido também usando os métodos IS, SuS, SepMC e MC tradicional.

Tabela 4 – Problemas de referência.

Problema	$g(\mathbf{X})$	Variáveis aleatórias	Referência
I	$0,1(X_1 - X_2)^2 - \frac{(X_1 + X_2)}{\sqrt{2}} + 2,5$	$X_1, X_2 \sim N(0, 1)$	Borri e Speranzini (1997)
II	$2 - X_2 - 0,1X_1^2 + 0,06X_1^3$	$X_1, X_2 \sim N(0, 1)$	Grooteman (2008)
III	$2 + 0,015 \sum_{i=1}^9 X_i^2 - X_{10}$	$X_1, X_2, \dots, X_{10} \sim N(0, 1)$	Grooteman (2008)
IV	$X_1^3 + X_2^3 - 18$	$X_1, X_2 \sim N(10, 5)$	Santosh <i>et al.</i> (2006)
V	$X_1^3 + X_2^3 - 18$	$X_1 \sim N(10, 5)$ $X_2 \sim N(9, 9, 5)$	Santosh <i>et al.</i> (2006)
VI	$X_1^3 + X_2^3 - 67,5$	$X_1 \sim N(10, 5)$ $X_2 \sim N(9, 9, 5)$	Santosh <i>et al.</i> (2006)
VII	$2,2257 - \frac{0,025\sqrt{2}}{27}(X_1 + X_2)^3 + 0,2357(X_1 + X_2)$	$X_1, X_2 \sim N(10, 3)$	Wang e Grandhi (1996)
VIII	$X_1X_2 - 2000X_3$	$X_1 \sim N(0,32, 0,032)$ $X_2 \sim N(1400000, 70000)$ $X_3 \sim \text{Lognormal}(100, 40)$	Santosh <i>et al.</i> (2006)
IX	$X_1 + 2X_2 + 3X_3 + X_4 - 5X_5 - 5X_6$	$X_1, X_2, X_3, X_4 \sim \text{Lognormal}(120, 12)$ $X_5 \sim \text{Lognormal}(50, 15)$ $X_6 \sim \text{Lognormal}(40, 12)$	Mahadevan e Pan (2001)
X	$X_1 + 2X_2 + 2X_3 + X_4 - 5X_5 - 5X_6$ $+ 0,001 \sum_{i=1}^6 \text{sen}(100X_i)$	$X_1, X_2, X_3, X_4 \sim \text{Lognormal}(120, 12)$ $X_5 \sim \text{Lognormal}(50, 15)$ $X_6 \sim \text{Lognormal}(40, 12)$	Liu e Der Kiureghian (1991)

XI	$-240758,1777 + 10467,364X_1 + 11410,63X_2 + 3505,3015X_3 - 246,81X_1^2 - 285,3275X_2^2 - 195,46X_3^2$	$X_1 \sim \text{Lognormal}(21,2, 0,1)$ $X_2 \sim \text{Lognormal}(20, 0,2)$ $X_3 \sim \text{Lognormal}(9,2, 0,1)$	Mahadevan e Pan (2001)
XII	$X_1X_2 - 78,12X_3$	$X_1 \sim N(2 \times 10^7, 0,5 \times 10^7)$ $X_2 \sim N(10^{-4}, 0,2 \times 10^{-4})$ $X_3 \sim \text{Gumbel}(4, 1)$	Santosh <i>et al.</i> (2006)
XIII	$12 - X_1 - X_2 - \sqrt{ X_1 - 5 }$	$X_1, X_2 \sim N(4, 1)$	Motta <i>et al.</i> (2021)

Fonte: O autor (2022).

Foi fixado $M = 10^6$ para o SepMC, com exceção dos problemas cuja separação das variáveis não foi possível, e o valor de N foi variado até alcançar a precisão desejada, sempre considerando que R é a parte mais custosa de $g(\mathbf{X}) = g(\mathbf{C}, R)$. Apesar do tamanho total da amostra gerada para o SepMC ser $M \times N$ e, conseqüentemente, o NFE total ser maior que o do MC tradicional, o número de avaliações de R é menor. Quanto mais custosa for R , maior é o benefício do método, que é medido aqui pelo N exigido para um determinado $CoV(\tilde{P}_f)$ em comparação com outros métodos.

Os resultados são apresentados na Tabela 5, onde foi fixado $CoV(\tilde{P}_f) = 5\%$ para todos os problemas, de modo a medir o NFE necessário para cada método. Para medir o erro relativo, foi considerada como probabilidade de falha exata aquela obtida pelo MC tradicional com $N = 10^7$. Uma vez que o menor número possível que o NFE do SMC pode assumir é o NFP do MC, essa informação é mostrada para cada problema, além do número de dimensões e a probabilidade de falha exata.

Tabela 5 – Resultados dos problemas de referência para $CoV(\tilde{P}_f) = 5\%$.

Problema	Dados*	Método	Amostra	NFE^{**}	$\tilde{P}_f$	Erro
I	$P_f = 4,200 \times 10^{-3}$, $n = 2$, $NFP = 350$	MC	$N = 85330$	85330	$4,102 \times 10^{-3}$	2,33%
		IS	$N = 1530$	1542	$4,152 \times 10^{-3}$	1,14%
		SepMC	$N = 80000$, $M = 80000$	80000	$4,332 \times 10^{-3}$	3,15%
		SuS	$N = 42210$	42210	$4,379 \times 10^{-3}$	4,27%
		SMC	$N = 85330$	407	$4,102 \times 10^{-3}$	2,33%
II	$P_f = 3,448 \times 10^{-2}$, $n = 2$, $NFP = 405$	MC	$N = 11550$	11550	$3,506 \times 10^{-2}$	1,67%
		IS	$N = 2250$	2262	$3,493 \times 10^{-2}$	1,29%
		SepMC	$N = 7000$, $M = 10^6$	7000	$3,692 \times 10^{-2}$	7,06%
		SuS	$N = 12080$	12080	$3,566 \times 10^{-2}$	3,41%
		SMC	$N = 11550$	444	$3,481 \times 10^{-2}$	0,94%

III	$P_f = 1,652 \times 10^{-2}$, $n = 10$, $NFP = 394$	MC	$N = 24560$	24560	$1,604 \times 10^{-2}$	2,92%
		IS	$N = 1070$	1098	$1,608 \times 10^{-2}$	2,68%
		SepMC	$N = 18000$, $M = 10^6$	18000	$1,610 \times 10^{-2}$	2,55%
		SuS	$N = 17240$	17240	$1,668 \times 10^{-2}$	0,96%
		SMC	$N = 24560$	10343	$1,600 \times 10^{-2}$	3,16%
IV	$P_f = 5,492 \times 10^{-3}$, $n = 2$, $NFP = 389$	MC	$N = 71090$	71090	$5,472 \times 10^{-3}$	0,36%
		IS	$N = 1770$	1818	$5,504 \times 10^{-3}$	0,23%
		SepMC	$N = 5000$, $M = 10^6$	5000	$5,851 \times 10^{-3}$	6,54%
		SuS	$N = 40860$	40860	$5,420 \times 10^{-3}$	1,30%
		SMC	$N = 71090$	447	$5,472 \times 10^{-3}$	0,36%
V	$P_f = 5,659 \times 10^{-3}$, $n = 2$, $NFP = 404$	MC	$N = 71790$	71790	$5,628 \times 10^{-3}$	0,54%
		IS	$N = 1860$	2160	$5,618 \times 10^{-3}$	0,72%
		SepMC	$N = 5000$, $M = 10^6$	5000	$5,624 \times 10^{-3}$	0,61%
		SuS	$N = 40020$	40020	$5,588 \times 10^{-3}$	1,25%
		SMC	$N = 71790$	465	$5,628 \times 10^{-3}$	0,54%
VI	$P_f = 1,306 \times 10^{-2}$, $n = 2$, $NFP = 342$	MC	$N = 27600$	27600	$1,239 \times 10^{-2}$	5,15%
		IS	$N = 1870$	2764	$1,273 \times 10^{-2}$	2,55%
		SepMC	$N = 3000$, $M = 10^6$	3000	$1,278 \times 10^{-2}$	2,16%
		SuS	$N = 20060$	20060	$1,387 \times 10^{-2}$	6,18%
		SMC	$N = 27600$	395	$1,239 \times 10^{-2}$	5,15%
VII	$P_f = 0,965$, $n = 2$, $NFP = 96$	MC	$N = 100$	100	0,960	0,50%
		IS	$N = 4610$	4649	0,962	0,29%
		SepMC	$N = 20$, $M = 20$	20	0,940	2,57%
		SuS	$N = 3600$	3600	0,966	0,13%
		SMC	$N = 100$	98	0,960	0,50%
VIII	$P_f = 1,482 \times 10^{-2}$, $n = 3$, $NFP = 357$	MC	$N = 24250$	24250	$1,472 \times 10^{-2}$	0,65%
		IS	$N = 1010$	1054	$1,522 \times 10^{-2}$	2,72%
		SepMC	$N = 15000$, $M = 10^6$	15000	$1,485 \times 10^{-2}$	0,23%
		SuS	$N = 19080$	19080	$1,548 \times 10^{-2}$	4,48%
		SMC	$N = 24250$	555	$1,472 \times 10^{-2}$	0,65%
IX	$P_f = 1,646 \times 10^{-3}$, $n = 6$, $NFP = 408$	MC	$N = 244220$	244220	$1,671 \times 10^{-3}$	1,50%
		IS	$N = 2680$	2804	$1,643 \times 10^{-3}$	0,21%
		SepMC	$N = 100000$, $M = 100000$	100000	$1,646 \times 10^{-3}$	0,02%
		SuS	$N = 52500$	52500	$1,646 \times 10^{-3}$	0,02%
		SMC	$N = 244220$	5704	$1,671 \times 10^{-3}$	1,50%
X	$P_f = 1,223 \times 10^{-3}$, $n = 6$,	MC	$N = 31870$	31870	$1,230 \times 10^{-2}$	0,60%
		IS	$N = 1230$	1783	$1,237 \times 10^{-2}$	1,17%
		SepMC	$N = 20000$, $M = 20000$	20000	$1,227 \times 10^{-2}$	0,35%

XI	$P_f = 0,529,$ $n = 3,$ $NFP = 219$	$NFP = 392$	SuS	$N = 20940$	20940	$1,297 \times 10^{-2}$	6,08%
			SMC	$N = 31870$	2709	$1,236 \times 10^{-2}$	1,09%
			MC	$N = 400$	400	0,548	3,52%
			IS	$N = 1100$	1294	0,530	0,12%
			SepMC	$N = 110,$ $M = 10^6$	110	0,525	0,83%
			SuS	$N = 3600$	3600	0,525	0,83%
XII	$P_f = 6,644 \times 10^{-4},$ $n = 3,$ $NFP = 393$		SMC	$N = 600$	267	0,401	24,25%
			MC	$N = 582520$	582520	$6,747 \times 10^{-4}$	1,55%
			IS	$N = 2500$	2634	$6,322 \times 10^{-4}$	4,85%
			SepMC	$N = 1000,$ $M = 10^6$	1000	$7,290 \times 10^{-4}$	9,72%
			SuS	$N = 90920$	90920	$6,873 \times 10^{-4}$	3,45%
			SMC	$N = 582520$	752	$6,747 \times 10^{-4}$	1,55%
XIII	$P_f = 1,322 \times 10^{-2},$ $n = 2,$ $NFP = 421$		MC	$N = 30670$	30670	$1,373 \times 10^{-2}$	3,87%
			IS	$N = 1160$	1208	$1,387 \times 10^{-2}$	4,93%
			SepMC	$N = 10000,$ $M = 10000$	10000	$1,370 \times 10^{-2}$	3,65%
			SuS	$N = 20100$	20100	$1,336 \times 10^{-2}$	1,08%
			SMC	$N = 30670$	473	$1,366 \times 10^{-2}$	3,34%

*O NFP se refere ao MC.

**Para do SepMC, o NFE se refere a N .

Fonte: O autor (2022).

Como é possível observar pelo erro relativo, o SMC conseguiu calcular a probabilidade de falha com uma precisão semelhante aos demais métodos. A exceção a isso é o problema XI, cujo erro foi muito maior devido à forma da função de falha e sua não monotonicidade. Nos problemas com função de falha localmente monotônica (I, IV, V, VI, VII, VIII, IX e XII), nota-se que a $\tilde{P}_f$ foi igual à obtida pelo MC, o que mostra que nessas condições todos os pontos falhos são avaliados corretamente. Além disso, mesmo nos problemas II, III, X e XIII, que não satisfazem a monotonicidade, o SMC resultou em uma boa aproximação.

O NFE do SMC foi o menor dentre todos os métodos em 8 de 13 problemas, perdendo apenas para o SepMC (considerando $NFE = N$) nos problemas VII e XI e para o IS nos problemas III, IX e X. Vale salientar também que, no problema VII, a redução promovida pelo SMC foi pequena pelo fato da maior parte da amostra estar dentro do domínio de falha, de modo que poucos pontos seguros puderam ser eliminados da amostra. Desconsiderando esse caso especial, a redução do NFE ficou entre 55,5% e 99,87%, onde 10 problemas tiveram redução acima de 90%, o que demonstra a eficiência do método.

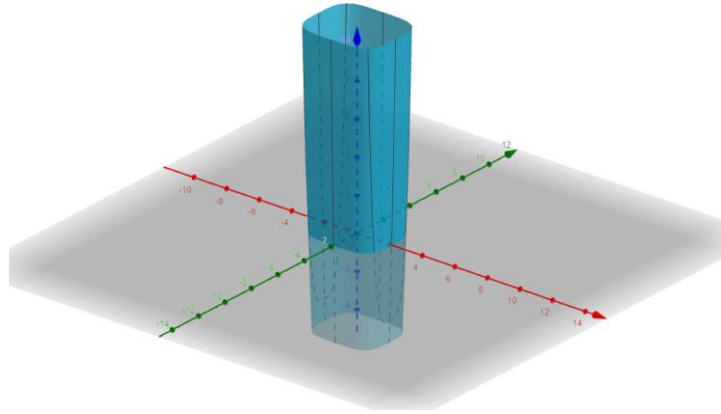
Nos problemas mostrados, com exceção de dois, o sentido do eixo no qual cada VA falha foi obtido com facilidade pelo cálculo do gradiente na média. Nos problemas II e XIII,

entretanto, foi necessária uma análise gráfica das funções de falha, dado que elas tinham um pequeno crescimento antes de decrescer em direção ao domínio de falha. Para problemas com mais de duas variáveis cuja determinação do sentido de falha pelo gradiente falhou, observar o gráfico de $g(\mathbf{X})$ não é possível, mas podem ser realizados testes variando a VA desejada e mantendo as demais com valor constante. De maneira geral, apesar do cálculo do gradiente funcionar para a maioria dos problemas alvos do SMC, a determinação do sentido de falha deve ser ponderada caso a caso, podendo ser necessário testar o comportamento de $g(\mathbf{X})$ para cada VA.

Vale salientar que, para os problemas V, VI e X, o FORM não conseguiu convergir em um limite de 100 iterações usando o algoritmo de busca HL-RF e, especialmente no problema X, nem mesmo usando o iHL-RF (SANTOS *et al.*, 2012). No problema XIII, cujo MPP tem coordenadas $\mathbf{x}^* = (5, 7)$, uma vez que ocorre uma singularidade no gradiente da função de falha quando $X_1 = 5$, o FORM não conseguiu convergir em um limite de 1000 iterações (MOTTA *et al.*, 2021). Como esperado, o SMC não teve dificuldades nesses exemplos, o que indica que ele pode ser também uma boa alternativa aos métodos de transformação para determinados problemas.

Por fim, foi observado que o SMC não obteve sucesso em alguns problemas da lista apresentada por Santos *et al.* (2012), o que era esperado, devido à forma dessas funções de falha. Além de não monotônicas, elas têm regiões de falha em mais de uma direção ou mudanças radicais no sentido do gradiente, de tal forma que ou o algoritmo não considerou todo o domínio de falha ou eliminou pontos no domínio de falha como se fossem seguros. No problema 21 (Figura 13), por exemplo, o domínio de falha se limita a uma pequena região em torno da origem, o que leva os pontos extremos em qualquer direção a serem seguros. É importante entender, no entanto, que esses problemas são exemplos matemáticos não associados a problemas práticos de engenharia estrutural e, em grande parte, os problemas práticos de confiabilidade estrutural satisfazem o requisito exigido pelo SMC na região de interesse.

Figura 13 – Problema 21 presente em Santos *et al.* (2012).



Fonte: O autor (2022).

4.3 PROBLEMAS DE ESTRUTURAS

Para examinar a metodologia em funções de falha associadas à engenharia civil, o SMC foi aplicado a dois problemas de estruturas: uma viga biapoiada e um pórtico simples composto de dois pilares e uma viga. Para ambos os problemas, foi analisada a falha pelo momento fletor solicitante superar o resistente, tendo a viga um único modo de falha e o pórtico três. Os resultados foram comparados com os métodos MC, IS, SuS e SepMC, todos para um $CoV(\tilde{P}_f) = 5\%$.

4.3.1 Problema 1: Viga biapoiada

O primeiro problema estrutural, retirado de Ghali *et al.* (2014), é uma viga simplesmente apoiada com vão $L = 6$ m, uma carga distribuída q e uma carga concentrada W atuando no meio do vão (Figura 14). Considerando que o momento fletor resistente é M_r , a função de falha é dada por

$$G(q, W, M_r) = M_r - \left(\frac{qL^2}{8} + \frac{WL}{4} \right) \quad (4.2)$$

onde as variáveis q , W e M_r tem distribuição normal e não possuem correlação. Seus parâmetros estatísticos são mostrados na Tabela 6.

Figura 14 – Viga biapoiada do problema 1.

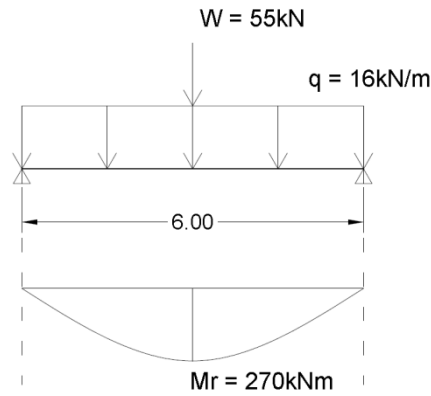
Fonte: Ghali *et al.* (2014).

Tabela 6 – Parâmetros estatísticos das variáveis do problema 1.

	q	W	M_r
Média	16 kN/m	55 kN	270 kNm
CoV	10%	15%	12%

Fonte: O autor (2022).

Assim como nos problemas resolvidos na subseção 4.2, o SMC foi comparado com os métodos MC, IS, SuS e SepMC, fixando o $CoV(\tilde{P}_f)$ em 5% para calcular o NFE exigido por cada método para uma mesma medida de precisão.

Para o SepMC, a função de falha foi reescrita como

$$g(C, R) = C - R \quad (4.3)$$

onde C e R tem tamanhos amostrais iguais a M e N , respectivamente, e são dados por

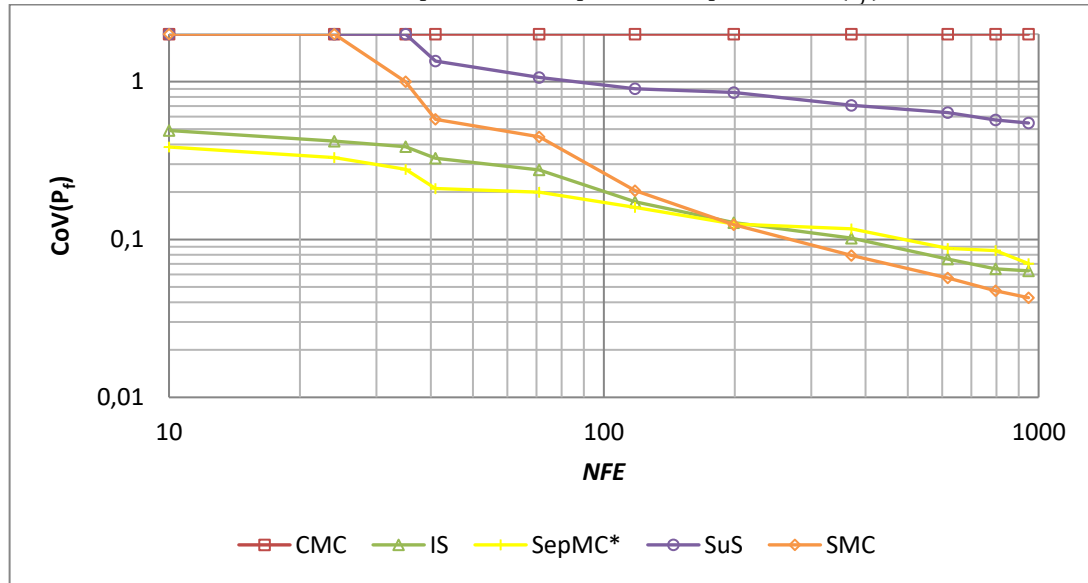
$$C(M_r) = M_r \quad (4.4)$$

$$R(q, W) = \left(\frac{ql^2}{8} + \frac{Wl}{4} \right) \quad (4.5)$$

Comparando a redução no $CoV(\tilde{P}_f)$ para todos os métodos à medida que o NFE aumenta, foi obtido o gráfico da Figura 15. No caso do SepMC, foi medido o número de avaliações de R fixando $M = 10^6$. Os pontos que obtiveram $\tilde{P}_f = 0$ e, com isso, $CoV(\tilde{P}_f) = \infty$ foram posicionados no topo do gráfico. É possível ver que o $CoV(\tilde{P}_f)$ do SMC decresce mais rápido que os demais métodos, sendo o menor de todos para valores de NFE acima de 199. Além disso, o IS e o SepMC tiveram valores próximos do SMC para mais de 199 avaliações,

tendo os três métodos uma diferença menor que 0,06 entre si no $CoV(\tilde{P}_f)$ ao ser aproximado de 1000 avaliações.

Figura 15 – $CoV(\tilde{P}_f)$ versus o número de avaliações da função de falha calculado pelo SMC em comparação com o MC, IS, SuS e SepMC, onde os pontos no topo têm $CoV(\tilde{P}_f) = \infty$.



*Para o SepMC, $M = 10^6$ e $NFE = N$.

Fonte: O autor (2022).

Os resultados desse problema, adotando $CoV(\tilde{P}_f) = 5\%$ e $M = 10^6$ para o SepMC, são mostrados na Tabela 7. O valor exato da probabilidade de falha, dado por Ghali *et al.* (2014), é $P_f = 0,557 \times 10^{-3}$.

Tabela 7 – Análise dos resultados do problema da viga biapoada.					
Referência: Ghali <i>et al.</i> (2014)			$P_f = 0,557 \times 10^{-3}$		
Método	Amostra	NFE^*	$\tilde{P}_f$	$CoV(\tilde{P}_f)$	Erro
CMC	$N = 698000$	698000	$0,573 \times 10^{-3}$	5%	2,87%
IS	$N = 1530$	1544	$0,538 \times 10^{-3}$	5%	3,41%
SepMC	$N = 2200, M = 10^6$	2200	$0,538 \times 10^{-3}$	5%	3,41%
SuS	$N = 88600$	88600	$0,530 \times 10^{-3}$	5%	5,03%
SMC	$N = 698000$	747	$0,573 \times 10^{-3}$	5%	2,87%

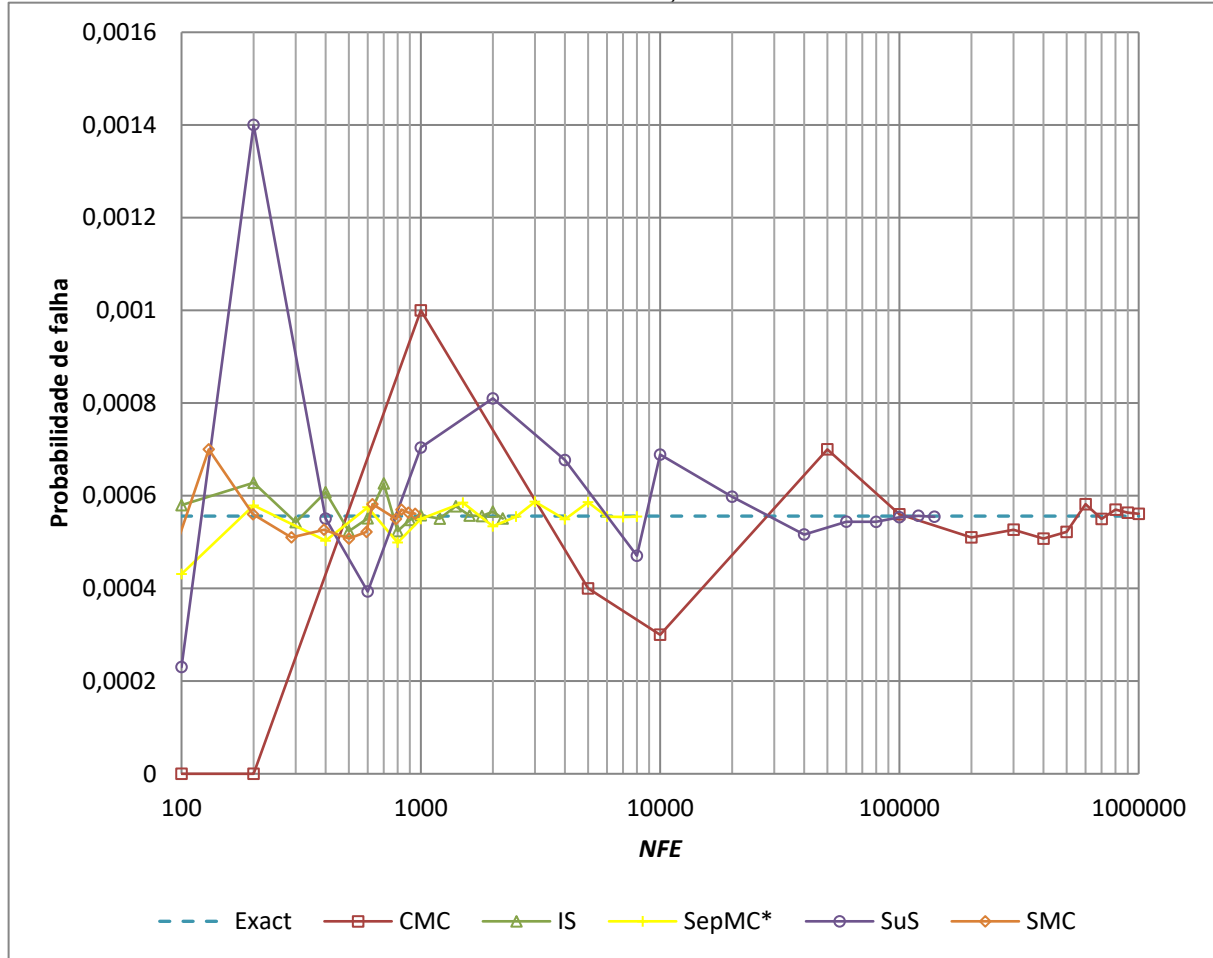
*Para do SepMC, o NFE se refere a N .

Fonte: O autor (2022).

O SMC obteve o mesmo resultado do MC tradicional, considerando a mesma amostra de tamanho N , mas obteve um NFE em torno de 0,1% do tamanho total da amostra. Enquanto o SMC precisou de 747 avaliações da função de falha, o IS precisou de 1544, o SepMC de 2200, o SuS de 88600 e o MC de 698000. A Figura 16 mostra uma comparação da convergência

da $\tilde{P}_f$ de cada método em relação ao NFE , onde o SMC foi o primeiro método a convergir para o valor da P_f .

Figura 16 – Convergência da $\tilde{P}_f$ para cada método.

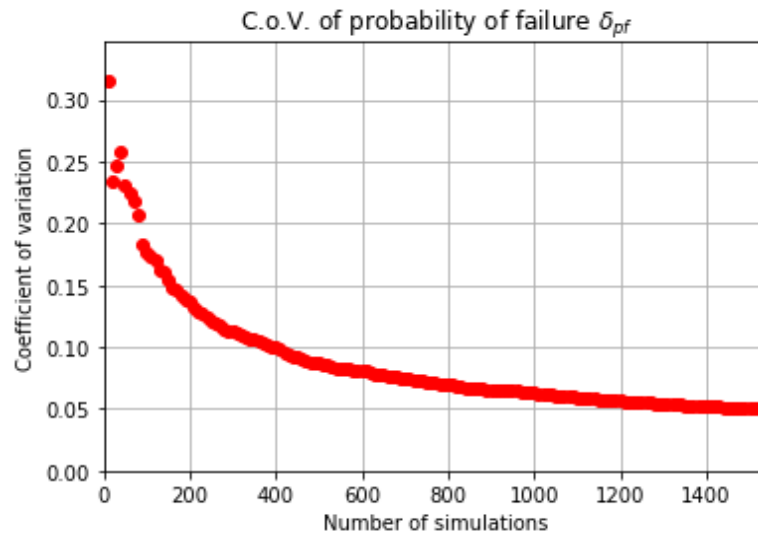


*Para o SepMC, $M = 10^6$ e $NFE = N$.

Fonte: O autor (2022).

Como a metodologia escolhida para o IS foi centrar a função de amostragem no MPP, o NFE do IS é a soma das N avaliações de $g(\mathbf{X})$ da amostra com as avaliações realizadas pelo FORM. Assim, além das avaliações para a amostra de tamanho $N = 1530$, o FORM precisou de 14 avaliações para encontrar o MPP, totalizando $NFE = 1544$. Dessa forma, apesar da amostra do IS ser menor que a do SMC, o último requer menos chamadas de $g(\mathbf{X})$. Uma visão mais detalhada dos valores do $CoV(\tilde{P}_f)$ do IS, desde a primeira simulação até atingir 5%, é mostrada na Figura 17.

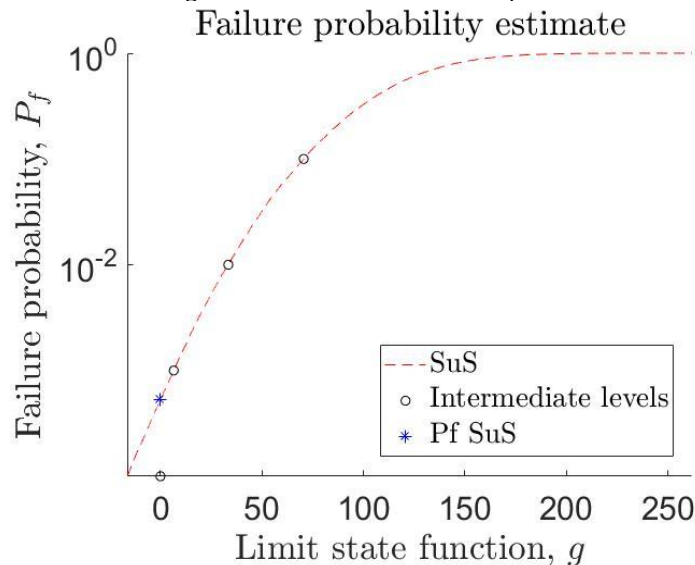
Figura 17 – $CoV(\tilde{P}_f)$ do IS à medida que a simulação avança.



Fonte: O autor (2022).

A despeito do SuS ter exigido mais avaliações de $g(\mathbf{X})$ que os métodos IS, SepMC e SMC, ele gerou uma redução de 87% no tamanho da amostra em relação ao MC tradicional. Quatro níveis foram necessários para a amostra gerada pelo SuS chegar no domínio de falha, sendo eles limitados por $b_1 = 70,66$, $b_2 = 34,16$, $b_3 = 6,21$ e $b_4 = 0$. Esses limites foram calculados de modo a ter uma probabilidade de falha condicional por nível igual a $p_0 = 10\%$, com exceção do último nível, no qual foi imposto o limite $b_4 = 0$. Assim, através do MCMC, foram gerados 22150 pontos amostrais por nível, obtendo $NFE = 88600$. A Figura 18 mostra a convergência da probabilidade de falha pelo SuS à medida que os níveis se aproximam do domínio de falha, sendo essa curva equivalente à CDF da função de falha estimada pelo SuS.

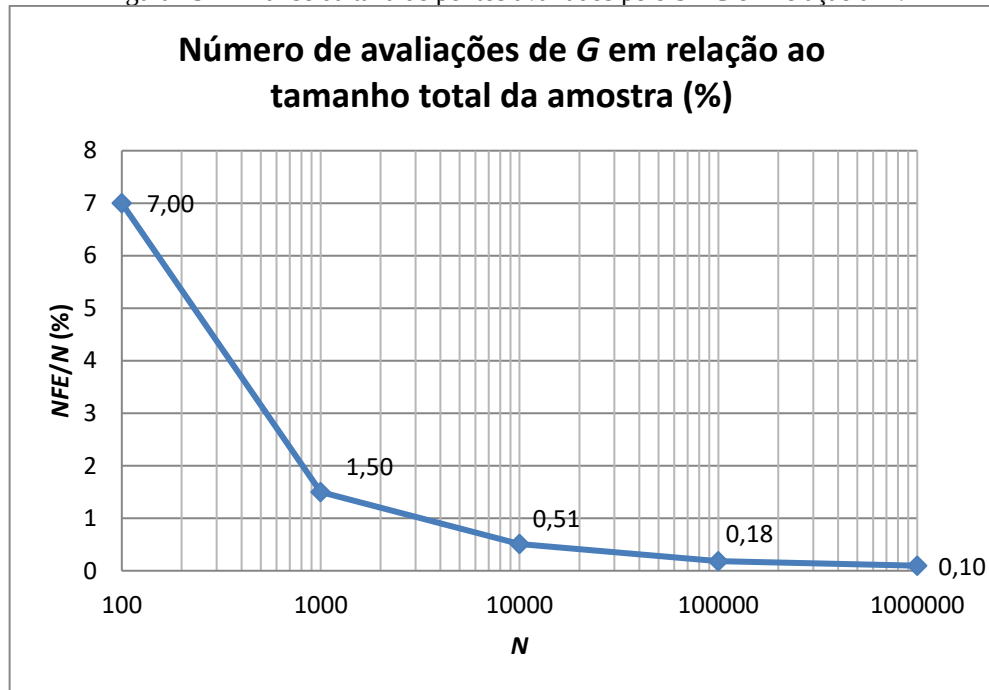
Figura 18 – CDF de G estimada pelo SuS.



Fonte: O autor (2022).

Se analisarmos o SMC para $N = 10^2, 10^3, 10^4, 10^5, 10^6$, podemos notar que, além da solução desse problema continuar igual ao MC, a proporção de avaliações em relação ao tamanho total da amostra permanece abaixo de 10% e diminui à medida que N aumenta. Assim, mesmo para um $CoV(\tilde{P}_f)$ menor que 5% ($N = 10^6$), a quantidade de avaliações de $g(\mathbf{X})$ e, por conseguinte, da variável R , ainda fica menor que todos os demais métodos. A Figura 19 mostra a análise dessa proporção para o problema apresentado.

Figura 19 – Análise da taxa de pontos avaliados pelo SMC em relação a N .

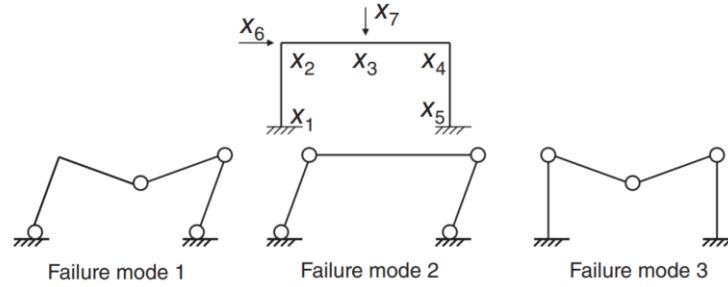


Fonte: O autor (2022).

4.3.2 Problema 2: Pórtico biengastado

O último problema estrutural é um pórtico simples biengastado com vão $L = 10$ m e altura dos pilares igual a 5 m, retirado de Au e Wang (2014) e Schuëller *et al.* (1989). Duas cargas concentradas são aplicadas na estrutura: X_6 , que é aplicado no topo do pilar da esquerda, e X_7 , que é aplicado no centro do vão. Os momentos fletores máximos que os nós do pórtico podem suportar são dados por X_1, X_2, X_3, X_4 e X_5 , conforme Figura 20. Os parâmetros estatísticos das variáveis do problema estão presentes na Tabela 8, onde as cargas concentradas tem distribuição máxima de Gumbel e os momentos fletores tem distribuição lognormal. Como resultado de referência para esse problema foi considerado o obtido por Schuëller *et al.* (1989) por integração numérica adaptativa com $N = 10^6$, $P_f = 2 \times 10^{-2}$.

Figura 20 – Modos de falha do pórtico simples do problema 2.



Fonte: Au e Wang (2014).

Tabela 8 – Parâmetros das distribuições das variáveis do problema 2.

	$X_1, \dots, X_5$	X_6	X_7
Distribuição	Lognormal	Gumbel	Gumbel
Média	60 kNm	20 kN	25 kN
CoV	10%	30%	30%

Fonte: O autor (2022).

Nesse problema, como pode ser visto na Figura 20, três modos de falha são possíveis dependendo da ordem em que as rótulas plásticas são formadas, ou seja, da ordem em que os nós atingem o momento máximo resistente. Uma função de falha é definida para cada um desses modos, representando a taxa entre a parte de solicitação (referente às cargas aplicadas) e a parte de resistência (referente aos momentos fletores resistentes nos nós) da equação de equilíbrio de momentos. Essas funções são dadas por

$$g_1(\mathbf{X}) = \frac{5X_6 + 5X_7}{X_1 + 2X_3 + 2X_4 + X_5} \quad (4.6)$$

$$g_2(\mathbf{X}) = \frac{5X_6}{X_1 + 2X_2 + X_4 + X_5} \quad (4.7)$$

$$g_3(\mathbf{X}) = \frac{5X_7}{X_2 + X_3 + X_4} \quad (4.8)$$

onde um valor de g_i maior que 1 significa que a solicitação supera a resistência, resultando em falha. Como qualquer um dos modos de falha causa o colapso da estrutura, só o maior valor das três funções de falha precisa ser avaliado, mesmo que mais de uma falha ocorra ao mesmo tempo. Assim, a falha é definida como $F = \{g > 1\}$, onde g é o maior valor entre g_1 , g_2 , e g_3 (equações 4.6 a 4.8). O problema pode, então, ser resumido na função

$$g = \max(g_1, g_2, g_3) \quad (4.9)$$

Adotando $CoV(\tilde{P}_f) = 5\%$ para todos os métodos, obtemos o número de avaliações da função de falha e da parte mais custosa exigidas por cada método para uma mesma medida de precisão. Os resultados são apresentados na Tabela 9. No caso do FORM realizado no IS, a busca utilizando a Equação 4.9 levou ao MPP de g_1 (Equação 4.6), de modo que a amostra foi centrada nesse ponto.

Tabela 9 – Resultados de probabilidade de falha e avaliações de função do problema 2.					
Referência: Schuëller <i>et al.</i> (1989)			$P_f = 2 \times 10^{-2}$		
Método	Amostra	NFE*	$\tilde{P}_f$	$CoV(\tilde{P}_f)$	Erro
MC	$N = 20010$	20010	$2,009 \times 10^{-2}$	5%	0,45%
IS	$N = 3150$	3284	$2,047 \times 10^{-2}$	5%	2,35%
SepMC	$N = 15000,$ $M = 15000$	15000	$1,923 \times 10^{-2}$	5%	3,85%
SuS	$N = 16100$	16100	$2,046 \times 10^{-2}$	5%	2,30%
SMC	$N = 20010$	2052	$2,009 \times 10^{-2}$	5%	0,45%

*Para do SepMC, o NFE se refere a N .

Fonte: O autor (2022).

Enquanto o MC tradicional precisou de 20010 avaliações de $g(\mathbf{X})$ e os outros métodos entre 3284 e 16100, o SMC precisou apenas de 2052, correspondendo a 10,26% do tamanho total da amostra (N). A solução foi semelhante à do problema 1, dado que as três funções de falha desse problema têm o mesmo sentido de falha. A única diferença é que os pontos dominantes tiveram que ser calculados três vezes cada um, para encontrar o valor máximo g .

5 CONCLUSÕES E SUGESTÕES

O método proposto, SMC, funciona buscando os pontos mais extremos no sentido de falha e delimitando o domínio de falha a partir deles, de modo a não precisar avaliar a amostra inteira para calcular a probabilidade de falha. Assim, seu objetivo é reduzir o número de avaliações da função de falha exigido pelo MC tradicional e, por conseguinte, reduzir o custo computacional envolvido nos métodos de simulação.

5.1 CONCLUSÕES DOS RESULTADOS OBTIDOS

A partir das aplicações e análise do método proposto, foi possível notar que ele apresentou uma redução no *NFE*, em relação ao MC, maior que o SuS em todos os problemas observados, maior que o IS em 12 e que o SepMC em 13 de 15 problemas, incluindo dois problemas de engenharia estrutural e problemas não monotônicos. Em problemas localmente monotônicos na região de interesse, o SMC obteve uma probabilidade de falha igual à obtida pelo MC tradicional para uma mesma amostra, o que está de acordo com o esperado. Mesmo em problemas não monotônicos, a $\tilde{P}_f$ calculada teve um valor próximo da calculada pelo MC tradicional e da P_f de referência, com exceção do problema XI, cuja forma da função de falha é mais irregular que as demais. No geral, desconsiderando o problema VII, cuja probabilidade de falha é muito próxima de 1 (situação que não é o foco do método), o SMC obteve uma redução no *NFE* entre 55,5% e 99,9%.

Porém, como esperado, o SMC teve dificuldades em problemas cuja forma da função de falha, além de não monotônica, é muito irregular. Em especial, problemas cuja falha está em mais de uma direção ou em uma região muito limitada, como os problemas 2, 11 e 21 presentes em Santos *et al.* (2012) (figuras 11 e 13). Esse comportamento é esperado, dado que o método foi desenvolvido para problemas com um sentido de falha claro para cada VA. Esse comportamento deve ser levado em conta antes de aplicar o SMC.

É importante destacar que o número de VAs do problema mostrou um impacto considerável na eficiência do SMC no problema estudado, mas o método se manteve vantajoso para todos os valores de n . Para duas VAs, o número de pontos da amostra avaliados em relação ao total ficou abaixo de 0,01%, enquanto para 10 variáveis essa taxa aumentou para aproximadamente 10%. Dessa forma, quanto menor for o número de variáveis do problema, mais eficiente o SMC se torna. É importante ressaltar que isso não significa necessariamente

que o método perde utilidade em problemas com muitas variáveis, dado que em todos os casos a redução se manteve acima de 90% no problema investigado.

A comparação com os outros métodos baseados no MC levou às seguintes conclusões:

- De maneira geral, o SMC tem uma maior redução no *NFE* que o SuS e o IS, apesar de precisar de uma amostra maior. Em casos nos quais o tamanho amostral é um problema, a combinação do SMC com algum método de redução de variância, como os próprios SuS e IS, pode ser a solução.
- O SMC tem a vantagem, em relação ao IS com função de amostragem centrada no MPP, de não depender da convergência do FORM para seu funcionamento.
- O SMC se mostrou melhor na redução do número de avaliações de R do que o SepMC na maioria dos problemas observados, para um M fixado em 10^6 . Entretanto, o SepMC não tem as mesmas limitações que o SMC, como observado no problema XI.
- A principal diferença do SepMC e SMC é que o primeiro calcula a parte mais custosa previamente para reduzir o número de vezes em que ela é avaliada, enquanto o segundo realiza o cálculo dessa parte diretamente na função de falha, pois $g(\mathbf{X})$ é chamada um número muito menor de vezes. É possível, ainda, combinar os dois métodos, já que o preço do SepMC para reduzir o custo computacional é um número maior de avaliações de g ($M \times N$ vezes), que pode ser reduzido pelo SMC.

5.2 SUGESTÕES PARA ESTUDOS FUTUROS

Considerando as limitações do SMC, estudos futuros podem focar em contorná-las para aumentar a robustez do método. O SMC teve dificuldade em problemas nos quais a região de falha não está em um sentido claro ou em um único sentido, o que está de acordo com o requisito de monotonicidade do método. Assim, as estratégias podem focar em considerar a mudança no sentido do gradiente da função de falha entre a média e o domínio de falha. É recomendado um estudo mais aprofundado de outros métodos de simulação, buscando técnicas que ajudem nessa questão.

Além disso, é importante observar o comportamento do método em contextos diferentes dos apresentados nesse trabalho, como funções com mais modos de falha e combinações do método proposto com outros métodos baseados no MC. Problemas mais complicados também

precisam ser avaliados, tanto em número de variáveis quanto no uso de métodos numéricos para a obtenção da função de falha, como o FEM. Nessas situações, os métodos com os quais o SMC foi comparado também podem aumentar seus benefícios, o que deve ser levado em conta. Além disso, é importante investigar e utilizar métodos eficientes para a obtenção dos pontos de Pareto (filtragem), visto que para problemas de grande dimensão e tamanho amostral esse processo pode ser um tanto custoso.

Por fim, é válido considerar o uso do SMC em problemas de otimização baseada em confiabilidade (RBDO), de modo a ter evidências da qualidade do método proposto em problemas desse campo. Da mesma forma que foi feito no presente trabalho, o erro relativo e o *NFE* podem ser analisados, sendo ainda interessante medir o tempo computacional. É interessante que exemplos focados em engenharia estrutural sejam escolhidos, buscando uma visão mais prática da aplicação do SMC na RBDO.

REFERÊNCIAS

- ALI, O. **Time-dependent reliability of FRP strengthened reinforced concrete beams under coupled corrosion and changing loading effects**. 2012. 325 p. Tese (Doutorado em Ciências de Engenharia) – Université d’Angers, Angers, França, 2012.
- ASSOCIAÇÃO BRASILEIRA DE NORMAS TÉCNICAS. **NBR 8800**: Projeto de estruturas de aço e de estruturas mistas de aço e concreto de edifícios. Rio de Janeiro: ABNT, 2008.
- ASSOCIAÇÃO BRASILEIRA DE NORMAS TÉCNICAS. **NBR 6118**: Projeto de estruturas de concreto — Procedimento. Rio de Janeiro: ABNT, 2014.
- AU, S. K.; BECK, J. L. Estimation of small failure probabilities in high dimensions by subset simulation. **Probabilistic Engineering Mechanics**, v. 16, n. 4, p. 263–277, 2001.
- AU, S. K.; WANG, Y. **Engineering Risk Assessment with Subset Simulation**. Hoboken, EUA: John Wiley & Sons, 2014.
- AUGUSTI, G.; BARATTA, A.; CASCIATI, F. **Probabilistic Methods in Structural Engineering**. Londres, Reino Unido: CRC Press, 1984.
- BECK, A. T. **Confiabilidade e Segurança das Estruturas**. 1. ed. São Paulo: GEN LTC, 2019.
- BORRI, A.; SPERANZINI, E. Structural reliability analysis using a standard deterministic finite element code. **Structural Safety**, v. 19, p. 361–382, 1997.
- BOURINET, J. M. **FERUM 4.1 User’s Guide**. Clermont-Ferrand, França: Institute Français de Mécanique Avancée (IFMA), 2010.
- CHAUDHURI, A.; KRAMER, B.; WILLCOX, K. E. Information Reuse for Importance Sampling in Reliability-Based Design Optimization. **Reliability Engineering and System Safety**, v. 201, 2020.
- BOURGUND, U.; BUCHER, C. G. **Importance sampling procedure using design points (ISPUD)**: A user's manual, Report No. 8-86, Innsbruck, Áustria: Institute of Engineering Mechanics, University of Innsbruck, 1986.
- CEN. **Eurocode 2**: Design of concrete structures, Part 1-1: General rules and rules for buildings, EN 1992-1-1:2004. CEN, Brussels, 2004.
- CEN. **Eurocode 3**: Design of steel structures — Part 1-1: General rules and rules for buildings, EN 1993-1-1, CEN, Brussels, 2005.
- CHEHADE, F. E. H.; YOUNES, R. Structural reliability software and calculation tools: a review. **Innovative Infrastructure Solutions**, v. 5, n. 1, 2020.
- CHENG, K. *et al.* Estimation of small failure probability using generalized subset simulation. **Mechanical Systems and Signal Processing**, v. 163, 2022.

CHING, J.; PHOON, K. K.; HU, Y. G. Efficient Evaluation of Reliability for Slopes with Circular Slip Surfaces Using Importance Sampling. **Journal of Geotechnical and Geoenvironmental Engineering**, v. 135, n. 6, p. 768–777, 2009.

CORNELL, A.C. First Order uncertainty analysis of soils deformation and stability. In: CONFERENCE ON APPLICATIONS OF STATISTICS AND PROBABILITY TO SOIL AND STRUCTURAL ENGINEERING, 1., 1971, Hong Kong, China. **Proceedings...** Hong Kong: Hong Kong University Press, 1971.

DER KIUREGHIAN, A.; DAKESSIAN, T. Multiple design points in first and second order reliability. **Structural Safety**, v. 20, n. 1, p. 37–49, 1998.

DER KIUREGHIAN, A.; HAUKAAS, T.; FUJIMURA, K. Structural reliability software at the University of California, Berkeley. **Structural Safety**, v. 28, n. 1–2, p. 44–67, 2006.

DER KIUREGHIAN, A.; LIU, P. Structural reliability under incomplete probability information. **Journal of Engineering Mechanics**, v. 112, n. 1, p. 85–104, 1986.

DIJKSTERHUIS, E. J. **Archimedes**. Princeton, EUA: Princeton University Press, 1987.

DITLEVSEN, O. Principle of normal tail approximation. **Journal of the Engineering Mechanics Division**, v. 107, n. 6, p. 1191–1209, 1981.

DITLEVSEN, O.; MADSEN, H. O. **Structural Reliability Methods**. Chichester, Reino Unido: John Wiley & Sons, 1996.

ERA GROUP. Subset Simulation for reliability analysis. Technische Universität München, 2021. Disponível em: <https://www.cee.ed.tum.de/era/software/reliability/subset-simulation/>. Acesso em: 14 fev. de 2022.

FIESSLER, B.; NEUMANN, H. J.; RACKWITZ, R. Quadratic limit states in structural reliability. **Journal of the Engineering Mechanics Division**, v. 105, n. 4, p. 661–676, 1979.

GHALI, A.; NEVILLE, A.; BROWN, T. G. **Structural analysis: a unified classical and matrix approach**. 5. ed. Boca Raton, EUA: CRC Press, 2014.

GROOTEMAN, F. Adaptive radial-based importance sampling method for structural reliability. **Structural Safety**, v. 30, p. 533–542, 2008.

HACKL, J. **Generic Framework for Stochastic Modeling of Reinforced Concrete Deterioration Caused by Corrosion**. 2013. 142 f. Dissertação (Mestrado em Engenharia Civil e Ambiental) – Norwegian University of Science and Technology, Trondheim, Noruega, 2013.

HASOFER, A. M.; LIND, N. C. Exact and invariant second moment code format. **Journal of the Engineering Mechanics Division**, v. 100, n. 1, p. 111–121, 1974.

HASTINGS, W. K. Monte carlo sampling methods using markov chains and their applications. **Biometrika**, v. 57, p. 97–109, 1970.

HOHENBICHLER, R.; RACKWITZ, R. Improvement of second-order reliability estimates by importance sampling. **Journal of Engineering Mechanics**, v. 114, p. 2195–2199, 1988.

HURTADO, J. E.; ALVAREZ, D. A. Neural-network-based reliability analysis: a comparative study. **Computer Methods in Applied Mechanics and Engineering**, v. 191, n. 1, p. 113–132, 2001.

KAYMAZ, I. Application of kriging method to structural reliability problems. **Structural Safety**, v. 27, n. 2, p. 133–151, 2005.

KESHTEGAR, B.; MIRI, M. Reliability analysis of corroded pipes using conjugate HL–RF algorithm based on average shear stress yield criterion. **Engineering Failure Analysis**, v. 46, p. 104–117, 2014.

KURRER, K. **The History of the Theory of Structures: From Arch Analysis to Computational Mechanics**. 1. ed. Berlim, Alemanha: Ernst Sohn, 2008.

LAPLACE, P. S. **Théorie analytique de probabilités**. Paris, França: Courcier, 1812.

LIU, P. L.; DER KIUREGHIAN, A. Optimization algorithms for structural reliability. **Structural Safety**, v. 9, n. 3, p. 161–177, 1991.

LOPEZ, R. H.; BECK, A. T. Reliability-based design optimization strategies based on form: a review. **Journal of the Brazilian Society of Mechanical Sciences and Engineering**, v. 34, p. 506–514, 2012.

MADSEN, H. O.; KRENK, S.; LIND, N. C. **Methods of Structural Safety**. Englewood Cliffs, EUA: Prentice Hall, 1986.

MAHADEVAN, S.; PAN, S. Multiple linearization methods for nonlinear reliability analysis. **Journal of Engineering Mechanics**, v. 127, n. 11, p. 1165–1172, 2001.

MAYER, M. **Die Sicherheit der Bauwerke**. Berlim, Alemanha: Springer, 1926.

MELCHERS, R. E. Importance sampling in structural systems. **Structural Safety**, v. 6, n. 1, p. 3–10, 1989.

MELCHERS, R. E.; BECK, A. T. **Structural reliability analysis and prediction**. 2. ed. Chichester, Reino Unido: John Wiley & Sons, 2017.

MENG, Z. *et al.* An active weight learning method for efficient reliability assessment with small failure probability. **Structural and Multidisciplinary Optimization**, v. 61, n. 3, p. 1157–1170, 2020.

MESSAC, A.; MULLUR, A. A. Multiobjective Optimization: Concepts and Methods. In: ARORA, J. S. **Optimization Of Structural And Mechanical Systems**. Iowa City, EUA: World Scientific Publishing Company. 2007. p. 121–147.

METROPOLIS, N. The beginning of the Monte Carlo method. **Los Alamos Science Special Issue**, v. 15, p. 125–130, 1987.

- METROPOLIS, N.; ROSENBLUTH, A. W.; ROSENBLUTH, M. N. Equations of state calculations by fast computing machines. **Journal of Chemical Physics**, v. 21, p. 1087–1091, 1953.
- MOTTA, R. S.; AFONSO, S. M. B. An efficient procedure for structural reliability-based robust design optimization. **Structural and Multidisciplinary Optimization**, v. 54, n. 3, p. 511–530, 2016.
- MOTTA, R. S. *et al.* Reliability analysis of ovalized deep-water pipelines with corrosion defects. **Marine Structures**, v. 77, 2021.
- MUSTAFFA, Z. Developments in Reliability-Based Assessment of Corrosion. *In*: ALIOFKHAZRAEI, M. **Developments in Corrosion Protection**. Londres, Reino Unido: InTech. 2014. p. 681-698.
- NATAF, A. Détermination des distributions de probabilités dont les marges sont données. **Comptes Rendus de l'Académie des Sciences**, v. 225, p. 42–43, 1962.
- NOWAK, A. S.; COLLINS, K. R. **Reliability of Structures**. Michigan, EUA: MacGraw Hill, 2000.
- PAN, Q.; DIAS, D. An efficient reliability method combining adaptive support vector machine and Monte Carlo simulation. **Structural Safety**, v. 67, p. 85–95, 2017.
- PAPAIIOANNOU, I. *et al.* MCMC algorithms for Subset Simulation. **Probabilistic Engineering Mechanics**, v. 41, p. 89–103, 2015.
- PEREIRA, R. M. R. **Probabilistic-based structural safety analysis of concrete gravity dams**. 2019. 321 f. Tese (Doutorado em Engenharia Civil) – Universidade Nova de Lisboa, Lisboa, Portugal, 2019.
- PIRES, K. O. *et al.* Reliability analysis of built concrete dam. **Revista IBRACON de Estruturas e Materiais**, v. 12, n. 3, p. 551–579, 2019.
- RACKWITZ, R.; FIESSLER, B. Structural reliability under combined random load sequence. **Computers and Structures**, v. 9, n. 5, p. 489–494, 1978.
- RAVISHANKAR, B. *et al.* Error estimation and error reduction in Separable Monte-Carlo method. **AIAA Journal**, v. 48, n. 11, p. 2624–2630, 2010.
- ROOS, D.; ADAM, U.; BAYER, V. Design reliability analysis. *In*: CAD-FEM USERS' MEETING 2006 – INTERNATIONAL CONGRESS ON FEM TECHNOLOGY WITH 2006 GERMAN ANSYS CONFERENCE, 24., 2006, Stuttgart, Alemanha. **Proceedings...** Stuttgart: CADFEM, 2006.
- ROSENBLATT, M. Remarks on a multivariate transformation. **Annals of Mathematical Statistics**, v. 23, n. 3, p. 470–472, 1952.
- ROSOWSKY, D. V. Structural Reliability. *In*: CHEN, W. F. **Structural Engineering Handbook**. Boca Raton, EUA: CRC Press LLC, 1999. p. 1473-1512.

RUBINSTEIN, R. Y. **Simulation and the Monte-Carlo Method**. Nova York, EUA: John Wiley & Sons, 1981.

SANTOS, D. M.; STUCCHI, F. R.; BECK, A. T. Reliability of beams designed in accordance with brazilian codes. **Revista IBRACON de Estruturas e Materiais**, v. 7, n. 5, p. 723–746, 2014.

SANTOS, S. R.; MATIOLI, L. C.; BECK, A. T. New optimization algorithms for structural reliability. **Computer Modeling in Engineering Sciences**, v. 83, p. 23–55, 2012.

SANTOSH, T. V. *et al.* Optimum step length selection rule in modified HLRF method for structural reliability. **International Journal of Pressure Vessels and Piping**, v. 83, p. 742–748, 2006.

SCHUËLLER, G. I. *et al.* On efficient computational schemes to calculate structural failure probabilities. **Probabilistic Engineering Mechanics**, v. 4, n. 1, p. 10-18, 1989.

SCHUËLLER, G. I.; JENSEN, H. A. Computational methods in optimization considering uncertainties—an overview. **Computer Methods in Applied Mechanics and Engineering**, v. 198, n. 1, p. 2–13, 2008.

SCHUËLLER, G. I.; PRADLWARTER, H. J.; KOUTSOURELAKIS, P. S. A critical appraisal of reliability estimation procedures for high dimensions. **Probabilistic Engineering Mechanics**, v. 19, n. 4, p. 463–474, 2004.

SELMİ, M. *et al.* Reliability analyses of slope stability. **European Journal of Environmental and Civil Engineering**, v. 14, p. 1227–1257, 2010.

SHAYANFAR, M. A. *et al.* An adaptive directional importance sampling method for structural reliability analysis. **Structural Safety**, v. 70, p. 14–20, 2018.

SHINOZUKA, M. Basic analysis of structural safety. **Journal of Structural Engineering**, v. 109, n. 3, p. 721–740, 1983.

SILVA, G. R. **Análise da confiabilidade da ligação laje-pilar interno sob punção de acordo com a NBR 6118:2014**. 2017. 174 p. Dissertação (Mestrado em Engenharia) – Universidade Federal do Rio Grande do Sul, Porto Alegre, 2017.

SMARSLOK, B. P.; ALEXANDER, D.; HAFTKA, R.; CARRARO, L.; GINSBOURGER, D. Separable Monte Carlo Simulation Applied to Laminated Composite Plates Reliability. *In*: AIAA/ASME/ASCE/AHS/ASC STRUCTURES, STRUCTURAL DYNAMICS, AND MATERIALS CONFERENCE, 49., AIAA/ASME/AHS ADAPTIVE STRUCTURES CONFERENCE, 16., 2008, Schaumburg, EUA. **Proceedings...** Schaumburg: American Institute of Aeronautics and Astronautics, 2008.

SMARSLOK, B. P. *et al.* Improving accuracy of failure probability estimates with Separable Monte Carlo. **International Journal of Reliability and Safety**, v. 4, n. 4, p. 393–414, 2010.

SOUZA JUNIOR, A. C. **Aplicação de confiabilidade na calibração dos coeficientes parciais de segurança de normas brasileiras de projeto estrutural**. 2008. 149 f.

Dissertação (Mestrado) – Escola de Engenharia de São Carlos, Universidade de São Paulo, São Carlos, 2008.

SU, G.; PENG, L.; HU, L. A gaussian process-based dynamic surrogate model for complex engineering structural reliability analysis. **Structural Safety**, v. 68, p. 97–109, 2017.

THEDY, J.; LIAO, K. W. Multisphere-based importance sampling for structural reliability. **Structural Safety**, v. 91, 2021.

TORII, A. J. *et al.* A performance measure approach for risk optimization. **Structural and Multidisciplinary Optimization**, v. 60, p. 927–947, 2019.

TORII, A. J.; LOPEZ, R. H.; MIGUEL, L. F. F. A Generalization of the sequential optimization and reliability assessment method for RBDO problems. *In: INTERNATIONAL SYMPOSIUM ON UNCERTAINTY QUANTIFICATION AND STOCHASTIC MODELLING*, 2., 2014, Rouen, França. **Proceedings...** Rouen: Springer, 2014.

TRUONG, V. H.; HA, M. H. Reliability-based design optimization of steel frames using direct design. *In: IOP CONFERENCE SERIES: MATERIALS SCIENCE AND ENGINEERING*, 2020, Hanoi, Vietnã. **Proceedings...** Hanoi: IOP Science, 2020. v. 869.

TRUONG, V. H.; KIM, S. E. An efficient method of system reliability analysis of steel cable-stayed bridges. **Advances in Engineering Software**, v. 114, p. 295–311, 2017.

UNGER, J.; ROOS, D. Investigation and benchmark of algorithms for reliability analysis. *In: WEIMARER OPTIMIERUNGS-UND STOCHASTIKTAGE*, 1., 2004, Weimar, Alemanha. **Proceedings...** Weimar: DYNARDO, 2004.

VILLA, C.; LABAYRADE, R. Energy efficiency vs subjective comfort: a multiobjective optimisation method under uncertainty. *In: CONFERENCE OF INTERNATIONAL BUILDING PERFORMANCE SIMULATION ASSOCIATION*, 12., 2011, Sydney, Austrália. **Proceedings...** Sydney, Austrália: International Building Performance Simulation Association, 2011, p. 1905-1912.

WANG, L.; GRANDHI, R. V. Safety index calculations using intervening variables for structural reliability. **Computers and Structures**, v. 59, p. 1139–1148, 1996.

WU, B. **Reliability Analysis of Dynamic Systems: Efficient Probabilistic Methods and Aerospace Applications**. 1. ed. Waltham, EUA: Academic Press, 2013.

YANG, D. X.; LI, G.; CHENG, G. D. Convergence analysis of First Order Reliability Method using chaos theory. **Computers & Structures**, v. 84, n. 8-9, p. 563–571, 2006.

YASEEN, Z. M.; ALDLEMY, M. S.; SADEGH, M. O. Non-gradient probabilistic gaussian global-best harmony search optimization for First-Order Reliability Method. **Engineering with Computers**, v. 36, n. 4, p. 1–12, 2019.

YOUN, B. D.; CHOI, K. K.; DU, L. Adaptive probability analysis using an enhanced hybrid mean value method. **Structural and Multidisciplinary Optimization**, v. 29, n. 2, p. 134–148, 2004.

YUN, W.; LU, Z.; JIANG, X. An efficient reliability analysis method combining adaptive kriging and modified importance sampling for small failure probability. **Structural and Multidisciplinary Optimization**, v. 58, n. 4, p. 1383–1393, 2018.

ZHANG, Y.; DER KIUREGHIAN, A. **Finite element reliability methods for inelastic structures**. Berkeley, EUA: Department of civil and environmental engineering, University of California, 1997. (Report UCB/SEMM - 97/05).

ZHANG, D. *et al.* Time-dependent reliability analysis through response surface method. **Journal of Mechanical Design**, v. 139, n. 4, 2017.

ZHANG, Z.; HESTHAVEN, J. S. Rare event simulation for large-scale structures with local nonlinearities. **Computer Methods in Applied Mechanics and Engineering**, v. 366, 2020.

ZHAO, H. *et al.* An efficient reliability method combining adaptive importance sampling and kriging metamodel. **Applied Mathematical Modelling**, v. 39, n. 7, p. 1853–1866, 2015.

ZHOU, S. T. *et al.* An improved first order reliability method based on modified armijo rule and interpolation-based backtracking scheme. **Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability**, v. 235, n. 2, p. 209–229, 2020.

ZUEV, K. M. Subset Simulation Method for Rare Event Estimation: An Introduction. *In*: BEER, M.; KOUGIOUMTZOGLOU, I.; PATELLI, E.; AU, I. K. **Encyclopedia of Earthquake Engineering**. Berlim, Alemanha: Springer, 2013.