
Algoritmo Evolutivo Multiobjetivo para o Escalonamento de Produção com Restrição na Indústria Farmacêutica

Debora Toshie Kohara



UNIVERSIDADE FEDERAL DE UBERLÂNDIA
FACULDADE DE COMPUTAÇÃO
PROGRAMA DE PÓS-GRADUAÇÃO EM CIÊNCIA DA COMPUTAÇÃO

Uberlândia
2024

Debora Toshie Kohara

**Algoritmo Evolutivo Multiobjetivo para o
Escalonamento de Produção com Restrição na
Indústria Farmacêutica**

Dissertação de mestrado apresentada ao Programa de Pós-graduação da Faculdade de Computação da Universidade Federal de Uberlândia como parte dos requisitos para a obtenção do título de Mestre em Ciência da Computação.

Área de concentração: Ciência da Computação

Orientador: Luiz Gustavo Almeida Martins

Coorientador: Gina Maira Barbosa de Oliveira

Uberlândia

2024

Ficha Catalográfica Online do Sistema de Bibliotecas da UFU
com dados informados pelo(a) próprio(a) autor(a).

K79 2024	Kohara, Débora Toshie, 1994- Algoritmo Evolutivo Multiobjetivo para o Escalonamento de Produção com Restrição na Indústria Farmacêutica [recurso eletrônico] : - / Débora Toshie Kohara. - 2024. Orientadora: Luiz Gustavo Almeida Martins. Coorientadora: Gina Maira Barbosa de Oliveira. Dissertação (Mestrado) - Universidade Federal de Uberlândia, Pós-graduação em Ciência da Computação. Modo de acesso: Internet. Disponível em: http://doi.org/10.14393/ufu.di.2024.637 Inclui bibliografia. Inclui ilustrações. 1. Computação. I. Martins, Luiz Gustavo Almeida, 1974-, (Orient.). II. Oliveira, Gina Maira Barbosa de, 1967-, (Coorient.). III. Universidade Federal de Uberlândia. Pós-graduação em Ciência da Computação. IV. Título. CDU: 681.3
-------------	---

Bibliotecários responsáveis pela estrutura de acordo com o AACR2:

Gizele Cristine Nunes do Couto - CRB6/2091
Nelson Marcos Ferreira - CRB6/3074



UNIVERSIDADE FEDERAL DE UBERLÂNDIA
Coordenação do Programa de Pós-Graduação em Ciência da
Computação

Av. João Naves de Ávila, 2121, Bloco 1A, Sala 243 - Bairro Santa Mônica, Uberlândia-MG,
CEP 38400-902

Telefone: (34) 3239-4470 - www.ppgco.facom.ufu.br - cpgfacom@ufu.br



ATA DE DEFESA - PÓS-GRADUAÇÃO

Programa de Pós-Graduação em:	Ciência da Computação				
Defesa de:	Dissertação de Mestrado, 35/2024, PPGCO				
Data:	30 de agosto de 2024	Hora de início:	14:00	Hora de encerramento:	15:55
Matrícula do Discente:	12112CCP007				
Nome do Discente:	Débora Toshie Kohara				
Título do Trabalho:	Algoritmo Evolutivo Multiobjetivo para o Escalonamento de Produção com Restrição na Indústria Farmacêutica				
Área de concentração:	Ciência da Computação				
Linha de pesquisa:	Inteligência Artificial				
Projeto de Pesquisa de vinculação:	Computação Bio-inspirada e Aprendizagem de Máquina na Resolução de Problemas.				

Reuniu-se por videoconferência, a Banca Examinadora, designada pelo Colegiado do Programa de Pós-graduação em Ciência da Computação, assim composta: Professores Doutores: Paulo Henrique Ribeiro Gabriel - FACOM/UFU, Alexandre Cláudio Botazzo Delbem - ICMC-USP e Luiz Gustavo Almeida Martins - FACOM/UFU, orientador da candidata.

Os examinadores participaram desde as seguintes localidades: Alexandre Cláudio Botazzo Delbem - São Carlos/SP. Os outros membros da banca e a aluna participaram da cidade de Uberlândia.

Iniciando os trabalhos o presidente da mesa, Prof. Dr. Luiz Gustavo Almeida Martins, apresentou a Comissão Examinadora e o candidato, agradeceu a presença do público, e concedeu ao Discente a palavra para a exposição do seu trabalho. A duração da apresentação do Discente e o tempo de arguição e resposta foram conforme as normas do Programa.

A seguir o senhor presidente concedeu a palavra, pela ordem sucessivamente, aos examinadores, que passaram a arguir á candidata. Ultimada a arguição, que se desenvolveu dentro dos termos regimentais, a Banca, em sessão secreta, atribuiu o resultado final, considerando o candidato:

Aprovado

Esta defesa faz parte dos requisitos necessários à obtenção do título de Mestre.

O competente diploma será expedido após cumprimento dos demais requisitos,

conforme as normas do Programa, a legislação pertinente e a regulamentação interna da UFU.

Nada mais havendo a tratar foram encerrados os trabalhos. Foi lavrada a presente ata que após lida e achada conforme foi assinada pela Banca Examinadora.



Documento assinado eletronicamente por **ALEXANDRE CLAUDIO BOTAZZO DELBEM, Usuário Externo**, em 02/09/2024, às 11:04, conforme horário oficial de Brasília, com fundamento no art. 6º, § 1º, do [Decreto nº 8.539, de 8 de outubro de 2015](#).



Documento assinado eletronicamente por **Luiz Gustavo Almeida Martins, Professor(a) do Magistério Superior**, em 02/09/2024, às 12:29, conforme horário oficial de Brasília, com fundamento no art. 6º, § 1º, do [Decreto nº 8.539, de 8 de outubro de 2015](#).



Documento assinado eletronicamente por **Paulo Henrique Ribeiro Gabriel, Professor(a) do Magistério Superior**, em 02/09/2024, às 13:32, conforme horário oficial de Brasília, com fundamento no art. 6º, § 1º, do [Decreto nº 8.539, de 8 de outubro de 2015](#).



A autenticidade deste documento pode ser conferida no site https://www.sei.ufu.br/sei/controlador_externo.php?acao=documento_conferir&id_orgao_acesso_externo=0, informando o código verificador **5635938** e o código CRC **6795BC7B**.

Gostaria de dedicar esta tese à todos que me auxiliaram e acreditaram em mim.

Agradecimentos

Gostaria de expressar minha profunda gratidão a todos que contribuíram para o sucesso desta dissertação de mestrado. Sem o apoio e a colaboração de cada um de vocês, este trabalho não teria sido possível. À minha família, que sempre esteve ao meu lado, oferecendo amor, incentivo e compreensão. Seus sacrifícios e encorajamento foram fundamentais para minha jornada acadêmica. Ao meu orientador, Prof. Luiz, cuja orientação sábia e paciência foram essenciais. Suas sugestões e direcionamento foram essenciais e ampliaram a minha visão. À Gina, minha coorientadora, pelo apoio e pelas discussões construtivas. Seu domínio enriqueceu minha pesquisa. Aos órgãos de fomento, CAPES e FAPEMIG, por acreditarem em meu potencial e pelo suporte financeiro concedido. Suas bolsas e recursos foram vitais para a realização deste estudo. Aos amigos e colegas, pelas discussões, trocas de ideias e momentos compartilhados. Suas contribuições foram inestimáveis. Por fim, agradeço a todos que, direta ou indiretamente, contribuíram para este trabalho. Que esta dissertação possa contribuir para o avanço do conhecimento em nossa área.

“Se uma máquina é esperada para ser infalível, ela não pode também ser inteligente.”

Alan Turing

Resumo

O escalonamento de bateladas em processos químicos em ambientes reais, por exemplo com variabilidade de demanda e restrições de vendas perdidas, tem sido investigados. Sendo tipicamente não lineares com funções de avaliação com alto custo computacional, dificultando a convergência de soluções. São caracterizados por múltiplos objetivos e restrições conflitantes entre si, aumentando a complexidade da exploração de soluções viáveis. Isto dificulta a adoção de certas técnicas, porém os Algoritmos Evolutivos Multiobjetivo (MOEA) têm sido explorados devido aos seus operadores flexíveis e eficientes, aplicáveis a diversos problemas não lineares multiobjetivo. Neste trabalho, novas estratégias de inicialização da população, seleção de indivíduos e mutação baseada em busca local são incorporadas a um MOEA, baseado no NSGA-II, aplicadas ao escalonamento de bateladas em uma linha de produção da indústria farmacêutica como um problema multiobjetivo com restrições. A qualidade das soluções não-dominadas foi avaliada, considerando o número de soluções válidas (NS) e métricas multiobjetivo, tais como: hipervolume (hv), taxa de erro (E) e distância geracional invertida (IGD+). Resultados experimentais mostraram que a abordagem evolutiva melhorou significativamente o modelo de referência, tanto em ambientes estáticos (onde um único conjunto de demandas é usado na avaliação), quanto dinâmicos (onde as demandas mudam ao longo das gerações do MOEA). Nos ambientes estáticos, nossa abordagem, em comparação à referência, melhorou 23,3% na média do E, 82,5% no IGD+, 0,3% no hv e teve uma piora de 30,2% no NS, enquanto que nos ambientes dinâmicos, melhorou em 38,7% no E, 83,7% no IGD+, 86,6% no NS e 0,8% no hv. Observamos que a integração de uma busca local ao operador de mutação reduziu em 10,5% no E, 0,1% no IGD+, 0,1% no hv e 10,7% no NS, comparado com o algoritmo que adota apenas a inicialização e o método de seleção propostos. A análise de dominância indicou que a busca local no número de bateladas é eficaz em cenários dinâmicos.

Palavras-chave: Otimização Multiobjetivo com restrição. Algoritmos Genéticos Multiobjetivos. Escalonamento de Bateladas. Manufatura Farmacêutica.

Abstract

Batch scheduling in chemical processes under real-world conditions, such as demand variability and lost sales constraints, has been investigated. These problems are typically non-linear with computationally expensive evaluation functions, making solution convergence challenging. They are characterized by multiple conflicting objectives and constraints, increasing the complexity of exploring viable solutions. This complicates the adoption of certain techniques, but Multiobjective Evolutionary Algorithms (MOEAs) have been explored due to their flexible and efficient operators, applicable to various non-linear multiobjective problems. In this work, new strategies for population initialization, individual selection, and local search-based mutation are incorporated into an NSGA-II-based MOEA, applied to batch scheduling in a pharmaceutical production line as a multiobjective constrained problem. The quality of non-dominated solutions was evaluated considering the number of valid solutions (NS) and multiobjective metrics such as hypervolume (hv), error rate (E), and inverted generational distance (IGD+). Experimental results showed that the evolutionary approach significantly improved the reference model in both static environments (where a single demand set is used for evaluation) and dynamic environments (where demands change over MOEA generations). In static environments, our approach, compared to the reference, improved by 23.3% the average of E, 82.5% in IGD+, 0.3% in hv, and had a 30.2% deterioration in NS. In dynamic environments, it improved by 38.7% in E, 83.7% in IGD+, 86.6% in NS, and 0.8% in hv. We observed that integrating local search into the mutation operator reduced E by 10.5%, IGD+ by 0.1%, hv by 0.1%, and NS by 10.7%, compared to the algorithm using only the proposed initialization and selection methods. Dominance analysis indicated that local search in the number of batches is effective in dynamic scenarios.

Keywords: Constrained Multiobjective Optimization. Multiobjective Genetic Algorithms. Batch Scheduling. Pharmaceutical Manufacturing.

Lista de ilustrações

Figura 1 – Fluxo de execução de um Algoritmo Genético.	39
Figura 2 – Representações de indivíduos	41
Figura 3 – População de indivíduos	42
Figura 4 – Conjunto de soluções e fronteira de não dominadas.	46
Figura 5 – Classificação de ordenação de não dominância.	49
Figura 6 – Cálculo da <i>Crowding Distance</i>	49
Figura 7 – Exemplo de planejamento de produção.	53
Figura 8 – Evolução das métricas de estoque	55
Figura 9 – Representação de um indivíduo.	62
Figura 10 – Fluxo de execução do modelo de referência	62
Figura 11 – Cruzamento Uniforme de referência.	64
Figura 12 – Operador de mutação de referência.	64
Figura 13 – Operadores de cruzamento baseados em pontos de corte avaliados. . . .	69
Figura 14 – Operador de Mutação Otimizado.	70
Figura 15 – Operações de mutação baseadas em heurística.	73
Figura 16 – Operador de mutação com busca local de bateladas.	74
Figura 17 – Distribuição de métricas hv e E por geração	94
Figura 18 – Distribuição de métricas IGD+ e NS por geração	94
Figura 19 – Distribuição de métricas hv e E por geração	102
Figura 20 – Distribuição de métricas IGD+ e NS por geração	102
Figura 21 – Melhores fronteiras de não dominados em relação ao hv.	103
Figura 22 – Melhores fronteiras de não dominados em relação ao NS.	104
Figura 23 – Melhores fronteiras de não dominados em relação ao IGD+.	104
Figura 24 – Melhores fronteiras de não dominados em relação ao E.	105
Figura 25 – Métricas hv e E	106
Figura 26 – Métricas IGD+ e NS	106
Figura 27 – Métrica Tempo médio (min)	106
Figura 28 – Evolução das métricas hv, E, IGD+, NS e tempo ao longo das gerações.	106

Lista de tabelas

Tabela 1 – Dados de processo	53
Tabela 2 – Tempo de troca (em dias) para a troca entre produtos (<i>setup</i>).	53
Tabela 3 – Distribuição Triangular para a demanda estocástica e Meta de Estoque (kg).	56
Tabela 4 – Métricas avaliação para os testes de inicialização aleatória.	82
Tabela 5 – Métrica CS das estratégias de inicialização aleatória.	82
Tabela 6 – Teste de hipóteses comparando p-ng5 em relação às estratégias de inicialização aleatória.	83
Tabela 7 – Métricas avaliação para os testes de sensibilidade de pMutG.	84
Tabela 8 – Métrica CS das estratégias de sensibilidade de pMutG.	84
Tabela 9 – Teste de hipóteses para as métricas comparando p-mMG7 em relação à outros	85
Tabela 10 – Métricas avaliação para os testes de sensibilidade de pAddGen.	85
Tabela 11 – Métrica CS das estratégias de sensibilidade de pAddGen.	86
Tabela 12 – Teste de hipóteses para as métricas comparando p-mMG7AG100 em relação à outros	86
Tabela 13 – Métricas avaliação de sensibilidade de pMutBat.	86
Tabela 14 – Métrica CS das estratégias de sensibilidade de pMutBat.	87
Tabela 15 – Teste de hipóteses para as métricas comparando p-mMB25 em relação à outros	87
Tabela 16 – Métricas avaliação de sensibilidade de pAddBat	88
Tabela 17 – Teste de hipóteses para as métricas comparando p-MAddB25 em relação à outros	88
Tabela 18 – Métricas avaliação para a sensibilidade de probabilidade de cruzamento.	90
Tabela 19 – Métrica CS das estratégias ref p-cC50 p-cC70	90
Tabela 20 – Teste de hipóteses para as métricas comparando p-cC70 em relação à outros	90
Tabela 21 – Métricas avaliação para os testes	91

Tabela 22 – Métrica CS das estratégias p-MAddB25 p-C1PV p-C1PF p-C2PF p-C2PFV p-C2PV	92
Tabela 23 – Teste de hipóteses para as métricas comparando p-C1PF em relação à outros	92
Tabela 24 – Métricas avaliação para os testes ref-s1 ini-rand	93
Tabela 25 – Desempenho do ini-rand em relação ao modelo de referência.	93
Tabela 26 – Análise comparativa entre o ini-rand e o modelo de referência.	93
Tabela 27 – Métricas avaliação para os testes	95
Tabela 28 – Métricas avaliação para os testes de inicialização de população heurísticas	96
Tabela 29 – Métrica CS das estratégias de inicialização de população.	96
Tabela 30 – Teste de hipóteses para as métricas comparando ini-h-p em relação às outras estratégias de inicialização.	97
Tabela 31 – Desempenho do algoritmo de acordo com a estratégia de mutação utilizada.	98
Tabela 32 – Métrica CS das estratégias ini-rand mut-h-pH25 mut-h-pH50 mut-h-pH75	98
Tabela 33 – Teste de hipóteses comparando mut-h-pH25 em relação às demais estratégias de mutação investigadas	99
Tabela 34 – Métricas avaliação para os testes de inicialização de população heurísticas	99
Tabela 35 – Métrica CS das estratégias de inicialização.	100
Tabela 36 – Teste de hipóteses para as métricas comparando ini-heu em relação à outras estratégias de inicialização	101
Tabela 37 – Métricas avaliação para os testes	103
Tabela 38 – Métricas avaliação para os testes	107
Tabela 39 – Métrica CS das estratégias	109
Tabela 40 – Teste de hipóteses para as métricas comparando ini-heu-g1000 em relação à pC90 em diferentes gerações	109
Tabela 41 – Teste de hipóteses para as métricas comparando ini-heu-g1000 em relação à pC80 para diferentes gerações	110
Tabela 42 – Teste de hipóteses para as métricas comparando ini-heu-g1000 em relação ao pC60 em diferentes gerações	110
Tabela 43 – Métricas avaliação para os testes	111
Tabela 44 – Métrica CS das estratégias	112
Tabela 45 – Teste de hipóteses para as métricas comparando pC90 em relação à outros	112
Tabela 46 – Desempenho das abordagens que empregam reinserção particionada. .	113
Tabela 47 – Métrica CS das estratégias	115
Tabela 48 – Teste de hipóteses para as métricas comparando pC90-g600 em relação à outros	115

Tabela 49 – Métricas avaliação para os testes	117
Tabela 50 – Análise de dominância entre as abordagens, considerando o percentual de uso da busca local nas operações de mutação.	118
Tabela 51 – Teste de hipóteses comparando d-pMG40 em relação às abordagens que utilizam busca local de bateladas nas operações de mutação. . . .	119
Tabela 52 – Teste de hipóteses para as métricas comparando d-BLB75 em relação com outros percentuais da média do número de bateladas.	119
Tabela 54 – Análise de dominância entre as abordagens evolutivas propostas e o modelo de referência.	120
Tabela 53 – Comparação entre as principais abordagens investigadas e o modelo de referência.	121
Tabela 55 – Teste de hipóteses comparando o algoritmo d-pMG40 com as melhores abordagens investigadas e o modelo de referência.	122
Tabela 56 – Desempenho dos modelos de amostragem em relação à abordagem evolutiva.	124
Tabela 57 – Desempenho das abordagens evolutivas considerando diferentes simulações de demanda.	125
Tabela 58 – Análise de dominância entre as abordagens evolutivas em ambientes dinâmicos.	126
Tabela 59 – Teste de hipóteses comparando o din-pMG40 em relação às demais abordagens em um ambiente com variação na demanda.	126
Tabela 60 – Desempenho do algoritmo em função da probabilidade de mutação de gene.	146
Tabela 61 – Métrica CS considerando as probabilidades de mutação de gene. . . .	147
Tabela 62 – Teste de hipóteses comparando d-pRe60 com as demais probabilidades de mutação de gene.	147
Tabela 63 – Métricas avaliação para os testes	149
Tabela 64 – Métrica CS considerando as probabilidades de adição de gene.	150
Tabela 65 – Teste de hipóteses comparando d-pMG40 com as demais probabilidades de adição de gene.	150
Tabela 66 – Desempenho do algoritmo em função da probabilidade de adição de gene adotada.	151
Tabela 67 – Métrica CS considerando as probabilidades de adição de gene.	152
Tabela 68 – Teste de hipóteses considerando as probabilidades de adição de gene. .	153

Lista de siglas

AG Algoritmos Genéticos

AGMO Algoritmos Genéticos Multiobjetivo

CS Cobertura entre dois Conjuntos (*Coverage of two Sets*)

E Taxa de erro

hv Hipervolume

IGD+ Distância Geracional Invertida + (*Inverted Generational Distance Plus*)

NS Número de Soluções não dominadas válidas

PCDA Perfis de Conjuntos de Demanda Aleatórios

PCDE Perfis de Conjuntos de Demanda Estáticos

PCDGA Perfis de Conjuntos de Demanda Geracionais Aleatórios

PIM Programação Inteira Mista

POM Problema de Otimização Multiobjetivo

POMR Problema de Otimização Multiobjetivo com Restrições

SMC Simulações de Monte Carlo

Sumário

1	INTRODUÇÃO	27
1.1	Motivação	31
1.2	Objetivos e Desafios da Pesquisa	32
1.3	Hipóteses e Questões da Pesquisa	33
1.4	Contribuições	34
1.5	Organização da Dissertação	34
2	FUNDAMENTAÇÃO TEÓRICA	37
2.1	Computação Evolutiva	37
2.1.1	Representação do Indivíduo	40
2.1.2	População	41
2.1.3	Função de Aptidão	42
2.1.4	Seleção dos Pais para Cruzamento	42
2.1.5	Operadores Genéticos	43
2.1.6	Reinserção da População	44
2.1.7	Critério de Parada	45
2.2	Problemas Multiobjetivos	46
2.2.1	Definições	47
2.2.2	Algoritmos Evolutivos Multiobjetivos	48
3	ESCALONAMENTO DA PRODUÇÃO FARMACÊUTICA .	51
3.1	Sequenciamento de Bateladas	51
3.2	Formulação do Problema Multiobjetivo	54
3.3	Trabalhos Relacionados	57
3.3.1	AGMO Aplicados em Problemas Multiobjetivo de Escalonamento com Restrições	57
3.3.2	Uso de Heurísticas no Aprimoramento das Otimizações	60
3.4	Modelo de Referência	61

4	DESENVOLVIMENTO	67
4.1	Algoritmo Baseado em Métodos Clássicos	67
4.2	Inicialização Baseada em Heurísticas	69
4.3	Cruzamento Produto-Batelada	71
4.4	Operadores de Mutação	72
4.4.1	Mutação Heurística	72
4.4.2	Mutação com Busca Local	74
4.5	Seleção Particionada	75
4.5.1	Reinserção	76
4.5.2	Seleção de Pais	76
5	EXPERIMENTOS E RESULTADOS	77
5.1	Ambiente de Experimentos	77
5.1.1	Metodologia para Cálculo das Métricas	78
5.1.2	Ambientes de Teste	79
5.1.3	Teste de Hipótese	80
5.2	Modelo com Inicialização Aleatória (ini-rand)	81
5.2.1	Inicialização de Indivíduos Aleatória	81
5.2.2	Mutação	83
5.2.3	Cruzamento	89
5.2.4	Melhor Modelo Randômico (ini-rand)	92
5.3	Modelo Heurístico (ini-heu)	95
5.3.1	Estratégias de Inicialização da População	95
5.3.2	Estratégias de Mutação	97
5.3.3	Integração da Heurística de Inicialização da População e Mutação	99
5.3.4	Análise de Gerações	101
5.3.5	Análise Gráfica ini-rand e ini-heu	103
5.4	Modelo baseado em Diversidade	105
5.4.1	Sensibilidade das Taxas de Cruzamento ao longo da Evolução	105
5.4.2	Seleção Particionada Baseada em Restrições	110
5.4.3	Sensibilidade em Relação à Mutação de Gene	115
5.4.4	Operador de Cruzamento Produto-Batelada	116
5.4.5	Operador de Mutação com Busca Local de Batelada	116
5.4.6	Análise Comparativa entre os Melhores Modelos	120
5.5	Avaliação da Eficiência de Exploração do AGMO	123
5.6	Sensibilidade das Abordagens a Variações de Demandas	124
6	CONCLUSÃO	127
6.1	Trabalhos Futuros	129
6.2	Contribuições Em Produção Bibliográfica	130

REFERÊNCIAS	133
--------------------	------------

APÊNDICES 143

APÊNDICE A	–	INVESTIGAÇÕES ADICIONAIS	145
A.1		Sensibilidade em Relação à Mutação de Gene	145
A.1.1		Análise da Probabilidade de Mutação de Gene	145
A.1.2		Análise da Probabilidade de Adição de Gene	148
A.2		Cruzamento Produto-Batelada	148

INTRODUÇÃO

As cadeias de suprimento podem ser classificadas entre três partes principais: produção, demanda e gestão de cadeia produtiva. Na produção, os sistemas de manufatura usualmente possuem capacidade produtiva restrita, diferentemente das demandas de produtos, são inerentemente voláteis e exposta a eventos estocásticos (JANKAUSKAS; FARID, 2019). Portanto projeções de mercado são utilizadas para realizar o planejamento da capacidade produtiva de modo a gerar estoques capazes de atender as demandas dentro do prazo esperado. Entretanto, a manutenção desses estoques possuem custos relacionados aos processos de compra, custos de ordem de compra e manutenção de estoque, bem como pode ocasionar perdas em decorrência da falta do estoque. Portanto, é necessário definir uma política de gestão e controle da cadeia produtiva de modo a otimizar o custo total de estoque apesar das incertezas na demanda (SAMAK-KULKARNI; RAJHANS, 2013). Uma das tarefas dessa gestão é a otimização do planejamento da capacidade de produção, ou seja, a definição da sequência de manufatura de produtos distintos e das quantidades a serem produzidas para cada linha de produção. Entretanto, pela natureza do problema percebe-se a presença de objetivos muitas vezes conflitantes como, por exemplo, o custo de manutenção do estoque e a capacidade de atendimento da demanda. Para aumentar a probabilidade de atendimento à demanda incerta dos clientes é necessário produzir grandes estoques dos produtos que, por sua vez, possuem impacto no custo total do estoque. Contudo, trabalhar com estoques mínimos para a redução de custos pode comprometer as vendas por um período, ocasionando atrasos na entrega ou perda de pedidos. Além disto, na modelagem de casos reais, é necessário considerar restrições do problema (ex: atender métricas de política de estoque) e eventos estocásticos que podem afetar significativamente o resultado da otimização, como é o caso de alterações sazonais de demanda. No caso da indústria farmacêutica, existem riscos adicionais durante o planejamento da capacidade produtiva de um novo produto e a sua sazonalidade. Em especial devido recursos financeiros e o tempo do processo de Pesquisa e Desenvolvimento (P&D) restritivos. O custo de P&D de um novo fármaco pode variar de 1,0 a 2,6 bilhões de dólares (NORMAN, 2016; HARRER et al., 2019; BATTELLE, 2015), e uma típica aprovação regulatória pode levar

cerca de 7 a 15 anos (HARRER et al., 2019; NORMAN, 2016; MAJOZI; SEID; LEE, 2015). Logo o planejamento de capacidade, mesmo em fases iniciais do desenvolvimento, é importante.

O escalonamento de produção em lotes de produtos (ou bateladas de produtos) em processos químicos tem sido objeto de pesquisa acadêmica e industrial desde o início da década de 1990, com foco em lidar com incertezas do mundo real, como demanda variável e capacidade de produção restrita (COTT; MACCHIETTO, 1989; KANAKA-MEDALA; REKLAITIS; VENKATASUBRAMANIAN, 1994; SHAH, 1998; HONKOMP; MOCKUS; REKLAITIS, 1997; HONKOMP; MOCKUS; REKLAITIS, 1999; HONKOMP et al., 2000). Nesse contexto, destaca-se a oportunidade de utilizar técnicas de escalonamento de tarefas para a otimização de múltiplos objetivos durante o sequenciamento de produção de distintos fármacos. O objetivo torna-se, portanto, encontrar um conjunto de soluções viáveis otimizadas em relação a múltiplos objetivos simultaneamente, proporcionando ao tomador de decisão flexibilidade na escolha de um plano de manufatura (sequência e quantidades dos produtos a serem produzidos) mais resiliente a determinados cenários de distribuições de demanda.

Diferentemente de problemas com apenas um objetivo (mono-objetivo), que possuem um espaço de busca com exploração mais simples, a meta da otimização multiobjetivo é encontrar um conjunto de soluções que superem as demais em pelo menos um dos objetivos considerados (soluções não dominadas). O conjunto ótimo de soluções não dominadas é denominado fronteira de Pareto. Dentre as características importantes para um algoritmo de otimização multiobjetivo, destacam-se a convergência, ou seja, sua capacidade de explorar o espaço de busca de modo a encontrar soluções ótimas para o problema de forma estável; a diversidade entre os indivíduos da população, resultando em um conjunto de soluções suficientemente distintas entre si; e o desempenho em relação ao tempo de execução adequado à aplicação do problema (LAUMANN et al., 2002). Frequentemente, para modelar problemas reais, é necessário considerar a presença de dois aspectos que aumentam a complexidade de problemas multiobjetivos, que serão investigados neste trabalho: problemas com restrições e avaliações de objetivos baseados em simulações com elevado custo computacional. A presença de restrições pode aumentar a complexidade do problema e dificultar a exploração eficiente do espaço de busca, uma vez que reduzem as regiões válidas desse espaço, podendo resultar em uma baixa diversidade entre as soluções encontradas e, conseqüentemente, uma baixa convergência do algoritmo (ZHANG; CHIONG, 2016; GAO et al., 2019; TIAN et al., 2021). Nesse contexto, estratégias que permitam uma exploração eficaz do espaço de soluções, com mecanismos que possibilitem "escapar" de máximos ou mínimos locais, devem ser avaliadas. Recentemente, houve um crescente interesse em aplicar técnicas elaboradas para a solução dos Problemas de Otimização Multiobjetivo com restrições (POMR) ao escalonamento de bateladas em processos químicos (ZHOU et al., 2018; LI et al., 2019; JANKAUSKAS; FARID, 2019;

FU et al., 2021; BEZDAN et al., 2022; ZHANG et al., 2022). Esses trabalhos exploram Algoritmos Genéticos Multiobjetivo (AGMO), abordagens híbridas com programação linear inteira mista e outras técnicas bioinspiradas. Muitas dessas abordagens ainda não foram aplicadas ao problema de escalonamento de bateladas na indústria farmacêutica, que apresenta objetivos conflitantes, restrições complexas e demandas incertas, apresentando potencial para otimizações de planejamento de processos e de estoque para redução de custo e receitas perdidas (CHEN et al., 2020; JANKAUSKAS; FARID, 2019). Além da eficácia em explorar espaços de busca complexos, algoritmos de otimização devem ser preferencialmente generalizáveis, ou seja, capazes de serem utilizados em diferentes problemas; e escaláveis, de modo que possam ser facilmente adaptados para resolver problemas com escopos maiores, sem comprometer sua eficiência. Métodos exatos podem ser utilizados, tais como: programação matemática (DREXL; KIMMS, 1997; OYEBOLU, 2019; MAJOZI; SEID; LEE, 2015; ZHANG et al., 2019) e programação linear (LIU et al., 2009; KOPANOS; MÉNDEZ; PUIGJANER, 2010; CASTRO; GROSSMANN; ROUSSEAU, 2011; MARCHETTI; MENDEZ; CERDA, 2012; VELEZ; MARAVELIAS, 2013; GEORGIADIS, 2021). Entretanto, esse tipo de abordagem costuma apresentar limitações de convergência, capacidade de generalização e escalabilidade de problemas, bem como não apresentar uma boa flexibilização para lidar com problemas multiobjetivo (MAJOZI; SEID; LEE, 2015; GEORGIADIS, 2021). Métodos baseados em heurísticas são específicos para determinados problemas (QUADT; KUHN, 2007) e não garantem a obtenção da solução ótima (SILVER, 1973; OYEBOLU, 2019; ZHANG et al., 2019); contudo buscam reduzir as limitações de linearidade e escalabilidade dos métodos exatos, podendo proporcionar soluções, pelo menos, próximas de soluções ótimas. Metaheurísticas são métodos mais gerais, capazes de adaptar a heurística da busca às características de um determinado problema de modo a aumentar seu desempenho (STÜTZLE, 1998). Dentre os tipos de metaheurísticas, os algoritmos baseados em conjunto de soluções, como os Algoritmos Genéticos (AG), são naturalmente flexíveis à otimização de múltiplos objetivos. AG são modelos bioinspirados que buscam explorar o espaço de busca do problema, evoluindo um conjunto de soluções (população de indivíduos ou cromossomos) por meio de operadores genéticos inspirados na teoria da seleção natural de Darwin (HOLLAND, 1992), como o cruzamento e mutação, os quais são responsáveis por guiar a busca na direção de regiões mais promissoras (MITCHELL, 1998). O conceito básico por trás dos métodos evolutivos é que os indivíduos mais aptos sobrevivam e se reproduzam, transmitindo sua carga genética ao longo das gerações (GOLDBERG; HOLLAND, 1988).

A evolução da população de soluções proporciona a exploração do espaço de busca e possibilita que o algoritmo supere máximos/mínimos locais, tornando-o adequado para a resolução de problemas de busca NP-completo ou difícil (ZITZLER; THIELE, 1998). AGMO lidam com dois ou mais objetivos, sendo capazes de manter na população indivíduos que priorizam, de forma diversificada, os objetivos considerados (FONSECA;

FLEMING, 1993; DEB, 1999; OYEBOLU, 2019; KÜSTER et al., 2021). Uma das etapas do AGMO é o cálculo das medidas de aptidão de uma solução candidata, que pode ser baseada em experimentos físicos ou simulações com alto custo computacional (alta fidelidade). Uma abordagem comum nesses estudos tem sido o uso de simulações de Monte Carlo para melhorar a qualidade das soluções. Neste trabalho é apresentado uma variação do NSGA-II (*Non-dominated Sorting Genetic Algorithm II*) (DEB et al., 2002), um MOEA bem conhecido e eficiente para problemas com 2 objetivos, adaptado para abordar o problema de otimização multiobjetivo de planejamento de capacidade de produção aplicado à indústria farmacêutica. Nossa abordagem foi concebida para tratar o problema apresentado em Jankauskas e Farid (2019), que busca lidar de forma eficiente com o escalonamento de bateladas em uma linha de produção de fármacos, encontrando um conjunto de soluções não dominadas robusto em relação à variabilidade de demanda do mercado. Nesse contexto, os objetivos consistem em maximizar a produção e reduzir o déficit entre a quantidade produzida e o estoque ideal de cada produto, enquanto a restrição envolve a necessidade de atender os pedidos de venda. Em outras palavras, trata-se de um problema multiobjetivo com restrição. A abordagem proposta consiste em uma variação do NSGA-II, que incorpora estratégias que visam melhorar a convergência do algoritmo, mais especificamente, o uso de uma heurística baseada em dados históricos para auxiliar na construção da população inicial, buscando gerar soluções mais promissoras, e de um operador de mutação baseado em busca local, a fim de proporcionar uma melhor exploração de uma localidade no espaço de soluções. Um dos objetivos dessas estratégias é reduzir a quantidade de gerações necessárias para obter um conjunto de soluções válidas e o número de avaliações (simulações de Monte Carlo) realizadas. Desta forma, diminui-se o tempo de execução do algoritmo, sem prejuízo significativo à qualidade do conjunto de soluções encontrado. Por outro lado, aumentar a convergência pode ocasionar muitos indivíduos similares, limitando a carga genética da população e, conseqüentemente, restringindo o alcance da exploração. Visando contornar esse problema, foi proposta uma estratégia de reinserção baseada na flexibilização da restrição, que visa assegurar a diversidade da população a partir da manutenção de soluções não dominadas inválidas (que não atendem a restrição do problema) como parte dos indivíduos da próxima geração. Essas estratégias foram implementadas de forma incremental, gerando quatro versões distintas do algoritmo: (i) ini-rand, que corresponde à implementação clássica do NSGA-II, considerando uma inicialização aleatória da população inicial e o uso de operadores genéticos tradicionais; (ii) ini-heu, que introduz a estratégia de inicialização baseada em heurística; (iii) d-pMG40, que incorpora a estratégia de reinserção proposta; e (iv) d-BLB75, que adiciona o novo operador de mutação baseado em busca local. Experimentos foram realizados a fim de avaliar a eficiência desses algoritmos em explorar de forma mais efetiva o espaço de soluções. Nossa abordagem também foi comparada ao algoritmo proposto em Jankauskas e Farid (2019), o qual foi denominado modelo de referência (ref). A qualidade dos conjun-

tos de soluções encontrados pelos algoritmos investigados foi determinada em função do número de soluções válidas encontradas (NS) e de diferentes métricas multiobjetivo, tais como: hipervolume (hv), taxa de erro (E) e distância geracional invertida (IGD+). O foco foi encontrar soluções de maior qualidade em relação às métricas avaliadas, abordando os desafios de convergência de soluções viáveis. Resultados experimentais mostram que todas as abordagens propostas apresentam melhoria significativa em relação ao modelo de referência. Nos experimentos com ambientes estáticos, onde um mesmo conjunto de demandas é usado na avaliação em todas as gerações do AG, o melhor desempenho foi do algoritmo d-pMG40, o qual alcançou uma melhora de 23,3% na taxa do erro, 82,5% no IGD+, 0,3% no hipervolume, embora apresente uma redução de 30,2% no NS. Nos experimentos em ambientes dinâmicos, onde simulações de Monte Carlo são realizadas a cada geração do AG, dois modelos propostos apresentaram bons resultados. O modelo din-pMG40 (equivalente ao modelo d-pMG40 avaliado no ambiente dinâmico) apresentou melhorias de 38,7% no E, 83,7% no IGD+, 86,6% no NS e 0,8% no hv. O modelo din-BLB (modelo d-BLB75 avaliado no ambiente dinâmico), em relação ao din-pMG40, mostrou uma redução de 10,5% no E, 0,1% no IGD+, 0,1% no hv e 10,7% no NS. Destacando a melhoria da busca local no número de bateladas em cenários dinâmicos.

1.1 Motivação

Um planejamento de produção eficaz deve ser capaz de flexibilizar a capacidade de produção disponível, de acordo com as flutuações de demanda do mercado, por meio de uma gestão eficiente de estoque e sequenciamento adequado dos produtos na linha de manufatura. No caso da indústria farmacêutica, existem riscos adicionais que dificultam esse planejamento: elevado investimento financeiro e de tempo, complexidade da pesquisa e dos processos regulatórios, bem como a alta volatilidade de demanda. Devido à necessidade de balancear a capacidade produtiva limitada e a demanda de mercado, há múltiplos objetivos frequentemente conflitantes, como o atendimento ao cliente dentro do prazo e a redução de estoques. Tais problemas possuem um espaço de busca mais complexo em relação a problemas com um único objetivo e requerem estratégias para guiar a busca para regiões próximas à fronteira de Pareto e manter a diversidade das soluções (DEB, 1999). Adicionalmente, em problemas reais, tipicamente há a presença de restrições que podem restringir ainda mais o espaço de busca, reduzindo a convergência e o uso de simulações com elevado custo computacional, limitando a aplicação em algumas situações. Majozí, Seid e Lee (2015) revelam que, apesar da complexidade do processo produtivo, há um baixo desenvolvimento de ferramentas de suporte à tomada de decisão para gestão de produção em processos biotecnológicos. Portanto, percebe-se a necessidade de estudos acerca de sistemas de planejamento e apoio à decisão para otimização da manufatura em indústrias farmacêuticas (MAJOZI; SEID; LEE, 2015), especialmente modelos de controle

implementáveis para otimização de startup, troca de linha de produção para distintos produtos e finalização (HONG et al., 2018). Além disso, há uma demanda por ferramentas aplicáveis em cenários reais, capazes de lidar com a otimização de múltiplos objetivos e resilientes a eventos estocásticos, como flutuações na demanda (JANKAUSKAS; PAPA-GEORGIOU; FARID, 2019).

Gao et al. (2019), em pesquisa bibliográfica no contexto do problema job-shop, concluem que modelos e estratégias para multiobjetivos ainda representam um desafio. Além disso, o desenvolvimento de estratégias de otimização mistas para aumentar a performance da busca combinatória permanece um ponto importante para algoritmos evolutivos.

1.2 Objetivos e Desafios da Pesquisa

O objetivo principal deste trabalho é investigar aprimoramentos nos operadores genéticos e novas estratégias de inicialização do AGMO, visando melhorar a convergência das soluções viáveis para o problema de escalonamento de bateladas na manufatura farmacêutica. Busca-se encontrar um conjunto de soluções não dominadas mais qualificado, abordando os desafios de convergência de soluções para soluções viáveis, de acordo com a restrição considerada. Durante o desenvolvimento desta pesquisa, pretendemos alcançar os seguintes objetivos específicos:

- ❑ Investigar a incorporação de heurísticas para melhorar a exploração do espaço de soluções restritas, buscando aprimorar a convergência do algoritmo.
- ❑ Avaliar o efeito de estratégias adaptativas de flexibilização das restrições na qualidade do conjunto de soluções válidas e não dominadas encontrado pelo AGMO.
- ❑ Investigar a integração de um método de busca local ao operador genético de mutação, com o objetivo de melhorar a exploração das soluções vizinhas aos indivíduos gerados pelo AGMO.
- ❑ Investigar e desenvolver um algoritmo genético multiobjetivo (AGMO) que seja capaz de encontrar diferentes planejamentos de produção para o problema de escalonamento de bateladas na manufatura farmacêutica, os quais representam diferentes estratégias de balanceamento entre produção e déficit e sejam robustas às variações de demanda do mercado. Avaliar a eficiência do algoritmo genético multiobjetivo proposto em encontrar o conjunto de soluções válidas e não dominadas, considerando as métricas Hipervolume (hv), Distância Geracional Invertida + (*Inverted Generational Distance Plus*) (IGD+), Taxa de erro (E), Cobertura entre dois Conjuntos (*Coverage of two Sets*) (CS) e Número de Soluções não dominadas válidas (NS).

- ❑ Analisar a qualidade das soluções encontradas pelos modelos desenvolvidos, considerando cenários estáveis (ambientes estáticos) e com maior variabilidade nas demandas de mercado (ambientes dinâmicos).

1.3 Hipóteses e Questões da Pesquisa

As hipóteses e questões de pesquisa investigadas durante o mestrado são apresentadas a seguir:

- ❑ A adoção de estratégias para a diversificação da população inicial permite aprimorar o conjunto de soluções encontrada pelo algoritmo com inicialização totalmente aleatória.
 - Aumentar a diversidade da inicialização da população permite que o AGMO alcance um conjunto de soluções mais qualificado, em relação às métricas de desempenho investigadas?
 - A diversificação da população inicial permite o uso de operadores genéticos de crossover e mutação mais simples que aqueles usados no modelo de referência (JANKAUSKAS; FARID, 2019), sem comprometer o desempenho do algoritmo?
- ❑ As operações do AG são mais efetivas ao considerar uma heurística que represente as características de produção.
 - A incorporação de heurística na inicialização da população pode aprimorar a convergência do algoritmo, resultando em um conjunto de soluções mais qualificado?
 - A incorporação de heurística na mutação pode aprimorar os resultados em comparação com a abordagem sem heurística?
- ❑ A utilização de técnicas de busca local pode melhorar a convergência do algoritmo e a qualidade do conjunto de soluções encontrado.
 - A incorporação de uma busca local baseada no número de bateladas como operador de mutação permite o refinamento da solução analisada, possibilitando uma evolução mais rápida do conjunto de soluções?
 - O ganho alcançado pela busca local compensa o custo computacional do processo?
- ❑ As estratégias propostas se mostraram efetivas e robustas, resultando nos melhores conjuntos de soluções em diferentes cenários de demandas.

- A busca aleatória é capaz de encontrar soluções similares ou próximas ao AGMO?
- A combinação das estratégias propostas proporcionaram uma maior efetividade do algoritmo multiobjetivo, tanto em relação ao modelo de referência, quanto ao algoritmo tradicional (ini-rand)?
- As abordagens desenvolvidas são mais robustas e eficientes tanto em ambientes estáticos, quanto dinâmicos, onde há maiores variações de demanda?

1.4 Contribuições

As principais contribuições do nosso trabalho foram:

- ❑ A proposta de uma estratégia de construção inicial baseada em dados históricos, que permite gerar indivíduos mais promissores, auxiliando na convergência do algoritmo e, conseqüentemente, na redução do número de gerações necessárias para alcançar um bom conjunto de soluções não dominadas válidas.
- ❑ A proposta de um operador de seleção com flexibilização da restrição permite o compartilhamento de informações genéticas dos melhores indivíduos válidos em relação aos objetivos. Essa estratégia promove uma maior diversidade de soluções e facilita a exploração de regiões válidas e inválidas do espaço de busca. Como resultado, aumenta-se a chance de encontrar soluções viáveis de alta qualidade, mesmo em um espaço de busca restrito e complexo.
- ❑ A incorporação da busca local de número de bateladas no operador de mutação melhora a exploração do espaço de busca, permitindo que o algoritmo refine as soluções candidatas de maneira mais precisa. Essa técnica resulta em soluções de maior qualidade ao identificar máximos locais de forma mais eficaz, melhorando a convergência do algoritmo em ambientes mais dinâmicos e, conseqüentemente, mais complexos.
- ❑ A proposta de uma abordagem evolutiva multiobjetivo capaz de lidar de forma eficiente e efetiva com as variações nas demandas de mercado.

1.5 Organização da Dissertação

Os demais capítulos deste trabalho estão estruturados da seguinte forma:

- ❑ **Capítulo 2:** Introduz os conceitos fundamentais relacionados aos algoritmos genéticos, problemas multiobjetivos e algoritmos evolutivos multiobjetivos, com destaque para o algoritmo NSGA-II, que será utilizado como referência.

- ❑ **Capítulo 3:** formaliza a tarefa de sequenciamento de bateladas (*scheduling*) para a produção de fármacos, na perspectiva de um problema de otimização multiobjetivo com restrição. Também é apresentada uma revisão dos trabalhos relacionados ao sequenciamento de bateladas, as estratégias evolutivas de manutenção de diversidade e a aplicação de algoritmos genéticos multiobjetivos nesse contexto específico. Adicionalmente, o modelo de referência usado na análise comparativa é apresentado, descrevendo como ele resolve o problema investigado.
- ❑ **Capítulo 4:** aborda as estratégias evolutivas propostas nesta pesquisa para aprimorar a eficiência do algoritmo genético multiobjetivo.
- ❑ **Capítulo 5:** descreve os principais experimentos realizados para avaliar o impacto das estratégias propostas no desempenho do algoritmo genético multiobjetivo e na qualidade do conjunto de soluções para o problema de escalonamento de bateladas. São descritas as quatro variações principais do AGMO: o modelo randômico (ini-rand), que incorpora inicialização aleatória, um operador de mutação aprimorado e um operador de crossover selecionado; o modelo heurístico (ini-heu), que incorpora uma inicialização de sequência de produtos baseada em heurística; e o modelo de seleção com particionamentos baseados em restrição (d-pMG40) e mutação com busca local (d-BLB75). O melhor modelo selecionado é comparado em relação à abordagens aleatórias. E em seguida realiza-se a comparação em relação à um ambiente de teste mais dinâmico. Os resultados experimentais são apresentados e analisados em função das métricas de desempenho multiobjetivo, adotadas para avaliar a qualidade do conjunto de soluções não dominadas.
- ❑ **Capítulo 6:** destaca as principais conclusões obtidas a partir dos experimentos realizados e das análises dos resultados. Também são discutidos possíveis desdobramentos dessa pesquisa que podem ser abordados em trabalhos futuros para aprofundar as investigações nesta área.

FUNDAMENTAÇÃO TEÓRICA

Este capítulo apresenta um embasamento teórico essencial para o entendimento da abordagem proposta neste trabalho. Primeiramente, são discutidos os conceitos da otimização evolutiva, com uma descrição dos algoritmos genéticos, incluindo uma breve descrição do funcionamento típico dos algoritmos genéticos. Em seguida, o problema de otimização multiobjetivo é introduzido, explicando os conceitos de dominância e fronteira de Pareto. Por fim, algoritmos evolutivos multiobjetivos são introduzidos, com ênfase no NSGA-II, utilizado em nossa pesquisa. Este capítulo apresenta um embasamento teórico acerca dos temas abordados neste trabalho, o qual é essencial para o entendimento da abordagem proposta. Inicialmente, são abordados os conceitos relativos à computação evolutiva, incluindo uma breve descrição do funcionamento típico dos algoritmos genéticos. O problema de otimização multiobjetivo também é abordado, introduzindo os conceitos de dominância e fronteira de Pareto e apresentando um algoritmo evolutivo multiobjetivo o NSGA-II que foi utilizado em nossa pesquisa.

2.1 Computação Evolutiva

A otimização de problemas pode ser abordada, de forma geral, por dois tipos de métodos: exatos e heurísticos (ROTHLAUF, 2011). Os métodos exatos são tipicamente aplicáveis quando a função objetivo é diferenciável, enquanto os métodos probabilísticos, que não requerem derivadas da função, são úteis em cenários onde o custo computacional é uma consideração secundária frente à necessidade de explorar espaços de busca complexos e não lineares. Dentro dos métodos heurísticos para problemas NP completo (ULLMAN, 1975), destaca-se a otimização evolutiva sendo adequada para ambientes dinâmicos e não lineares, repletos de discontinuidades e instabilidades estruturais (SCHWEFEL, 1995; JÁHN-ERDÖS; KÖVÁRI, 2022). A otimização evolutiva se inspira nos princípios da evolução biológica (DARWIN, 2023), como a seleção natural, mutação e recombinação genética. Como independem da diferenciabilidade da função objetivo, são capazes de ex-

plorar o espaço de busca de maneira robusta e adaptativa, embora não garantam uma solução exata. A computação evolutiva é uma técnica bio-inspirada que aborda a evolução como um processo de busca por soluções ótimas ou quase ótimas para problemas complexos. Nesse contexto, os algoritmos evolutivos podem ser vistos como procedimentos iterativos que simulam mecanismos evolutivos biológicos para resolver problemas tipicamente de busca ou otimização. Eles são fundamentados pela teoria sintática da evolução ou Neodarwinismo, que enfatiza a mutação, recombinação gênica e seleção natural como mecanismos evolutivos, contemplando as seguintes características essenciais (RIDLEY, 1996).

- ❑ População de indivíduos: Representa um conjunto de soluções candidatas (cromossomos ou indivíduos), onde cada indivíduo é um ponto no espaço de busca e possui um conjunto de características genéticas, que são utilizadas para a reprodução.
- ❑ Hereditariedade e variação genética: A troca da informação genética da população é utilizada para explorar o espaço de busca. Nesse sentido, novos indivíduos herdam características dos seus pais. Entretanto, para prover uma exploração mais ampla do espaço de busca, é necessária a inserção de novas informações genéticas na população por meio de mutações.
- ❑ Seleção natural: A cada geração selecionam-se os indivíduos com base em sua aptidão, de modo que aqueles mais aptos têm maiores chances de sobrevivência e reprodução, aumentando a probabilidade de propagação de sua carga genética. Dada a sua influência nos eventos de nascimento (seleção dos pais) e morte dos indivíduos, a seleção possibilita mudanças dinâmicas na população ao longo das gerações. Ela é baseada em uma função de aptidão, que define a qualidade ou capacidade de cada indivíduo em resolver o problema investigado.

Dentre os algoritmos evolutivos, destacam-se os Algoritmos Genéticos (AG) (HOLLAND, 1992). AG são metaheurísticas inspiradas nos princípios da evolução das espécies e seleção natural, comumente empregadas para resolver problemas de busca e otimização (MITCHELL, 1998). O processo de busca é bioinspirado no conceito de que os indivíduos mais aptos da população têm maior probabilidade de sobrevivência, bem como maiores chances de propagarem sua carga genética para as novas gerações (GOLBERG, 1989). Logo, busca evoluir iterativamente um conjunto de soluções candidatas, tipicamente denominados população de indivíduos. A cada iteração ou geração, indivíduos mais aptos da população têm maior probabilidade de serem selecionados para combinar-se com outro indivíduo e gerar um novo indivíduo potencialmente superior com a mistura de características (combinação de carga genética) de ambos os indivíduos pais, que seja potencialmente superior à ambos os pais, este procedimento é chamado de cruzamento. Este novo indivíduo pode passar pelo processo de mutação, em que ocorre uma modificação no indivíduo

(alteração da carga genética), desta forma durante a otimização pode ocorrer uma saída de um mínimo local. Ao final da geração apenas os indivíduos mais aptos são selecionados para continuar para a próxima evolução. E ao final a melhor solução da população é selecionada (GOLDBERG; DEB, 1991; MITCHELL, 1998). Assim, o algoritmo consegue explorar o espaço de soluções do problema, modificando a população por meio de operações genéticas de seleção, cruzamento e mutação. E por fim evoluindo através da seleção de apenas um subconjunto mais apto da população.

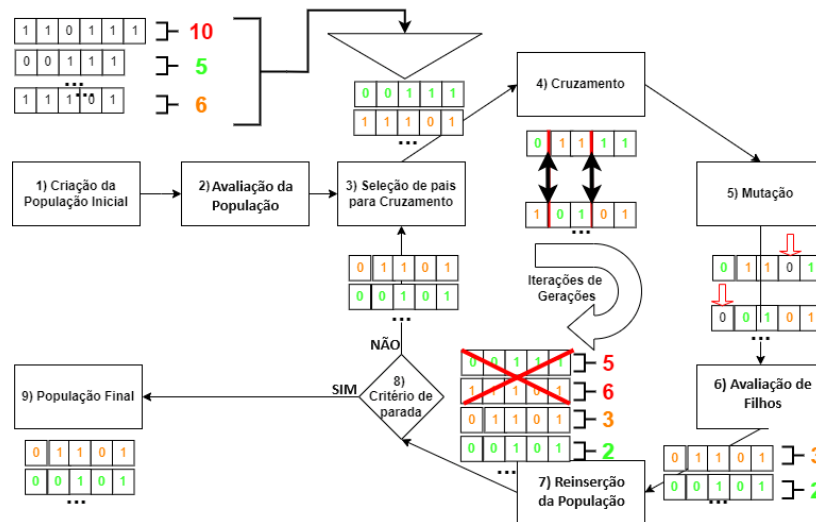


Figura 1 – Fluxograma representativo do processo geral de Algoritmos Genéticos.

O funcionamento típico de um algoritmo genético é ilustrado Figura 1.

- ❑ *1. Geração da população inicial:* O algoritmo genético começa com a geração da população inicial. Esse processo consiste em gerar um conjunto de soluções potenciais para o problema, as quais são denominadas indivíduos ou cromossomos. A quantidade de indivíduos na população é um parâmetro do AG, sendo que cada indivíduo é formado por um conjunto de genes. Portanto, essa etapa envolve a definição da quantidade de genes de cada indivíduo, quando é adotada uma representação que permite indivíduos de tamanho variável, e dos seus respectivos valores. De acordo com a característica do problema, pode-se empregar alguma estratégia para assegurar a criação de soluções válidas ou garantir a diversidade entre os indivíduos da população (ex: distribuição uniforme).
- ❑ *2. Avaliação dos indivíduos:* Nesta etapa calcula-se a aptidão de cada um dos indivíduos o que permite a quantificação e ordenação dos indivíduos mais apto, por exemplo em um problema de maximização uma função de aptidão é a soma da quantidade produzida por cada um dos planejamentos de produção.
- ❑ *3. Seleção de pais para o cruzamento:* Ocorre a seleção de indivíduos para participar da operação de cruzamento. Note que quanto mais apto, mais provável sua capaci-

dade de ser selecionado para cruzamento (MITCHELL, 1998). Um típico operador é o torneio de n indivíduos e o indivíduo com o melhor aptidão definido pela regra de prioridade é selecionado (MITCHELL, 1998).

- ❑ *4. Cruzamento*: Método de criação de soluções novas através de trocas de informações entre indivíduos. Desta forma altera-se o contexto de informações já disponíveis (BAECK, 1994). Um exemplo de cruzamento é o cruzamento de dois pontos, tal que seleciona-se aleatoriamente dois pontos de corte e alteram-se as seções advindas de pais copiando as para os filhos, gerando então dois filhos a partir de segmentos de dois pais.
- ❑ *5. Mutação*: Dentre os genes dos indivíduos filho realizam-se operações de modificação de atributos, desta forma gerando diversidade à população (BAECK, 1994) e explorando outras regiões do espaço de busca. Uma mutação comum é a seleção aleatória de determinados genes gerando uma alteração em seu valor.
- ❑ *6. Avaliação dos indivíduos filhos*: Calcula-se a aptidão para os novos indivíduos.
- ❑ *7. Reinserção de filhos e pais*: Define qual a estratégia da criação da nova população, ou seja, como será a inclusão de novos indivíduos à população corrente.
- ❑ *8. Critério de parada*: Os critérios de parada definem a estratégia utilizada para a parada das iterações de evolução, caso não seja atingido o processo retorna à etapa 2. Tipicamente utiliza-se um número pré-definido de gerações. Note que nas primeiras gerações é esperado ter soluções com objetivos baixos, porém ao longo das gerações o esperado é ter soluções com objetivos maiores.
- ❑ *9. População Final*: Caso o critério de parada seja satisfeito resultando em um conjunto de soluções ótimas ou satisfatórias seleciona-se então a melhor solução

2.1.1 Representação do Indivíduo

Um dos passos mais importantes é definir como será a representação (ou indivíduo) do problema para o qual se busca solução. Um indivíduo é uma abstração que contempla uma das possíveis soluções do problema. A representação da solução (indivíduo ou cromossomo) em AG influencia diretamente a capacidade de modelagem de problemas e o desempenho da otimização. Os indivíduos são implementados por meio de uma estrutura de dados, como: vetores de tamanho fixo ou variável, árvores ou grafos. Os tipos mais comuns de representações são: binária, inteira ou real (MITCHELL, 1998; HOLLAND, 1992). A representação de vetores binários de tamanho fixo é uma das formas mais populares, sendo amplamente utilizada devido à sua simplicidade de implementação e operação (HOLLAND, 1992). No entanto, essa representação pode não ser a mais adequada

para problemas numéricos de alta dimensionalidade que exigem alta precisão (MICHALEWICZ, 1996). Para esses casos, uma codificação de tamanho variável pode facilitar e agilizar a obtenção de resultados (GOLDBERG et al., 1993). Portanto, é essencial selecionar uma codificação compatível com a natureza do problema investigado para maximizar a eficiência do algoritmo.

Além das representações lineares, também é possível adotar estruturas mais complexas, como árvores ou grafos. Essas representações são particularmente interessantes para problemas em que as soluções devem representar algum tipo de organização espacial, como hierarquia ou vizinhança, sendo aplicáveis, por exemplo, no roteamento multicast (LA-FETA et al., 2018). A Figura 2 ilustra diferentes exemplos de representações.

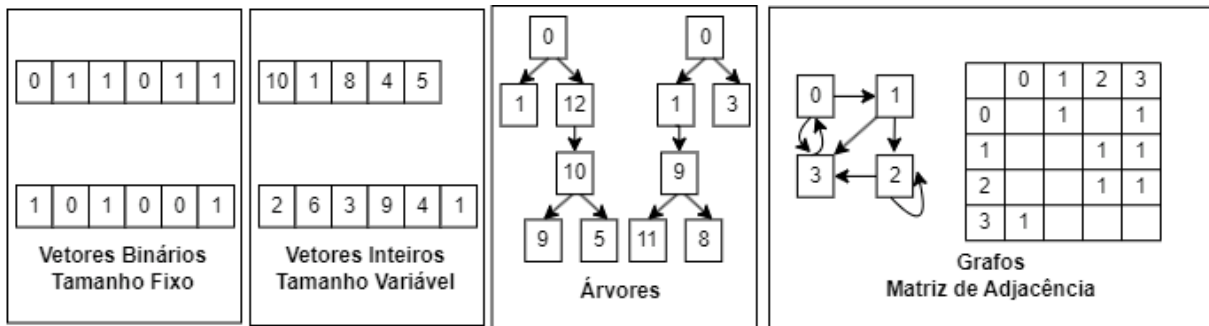


Figura 2 – Representações de indivíduos.

Essa diversidade de representações permite que os AGs sejam aplicáveis a uma ampla gama de problemas, proporcionando flexibilidade e eficiência na busca de soluções ótimas.

2.1.2 População

A população é o conjunto de indivíduos que sofrerão os processos evolutivos do Algoritmo Genético (AG). A população compõe o espaço de busca, e seu tamanho é um parâmetro do AG que deve ser ajustado conscientemente, pois afeta o desempenho, a eficiência e o tempo de execução do algoritmo. Na Figura 3, temos uma representação gráfica de uma população, composta por um conjunto de indivíduos, que por sua vez é formado por um conjunto de genes.

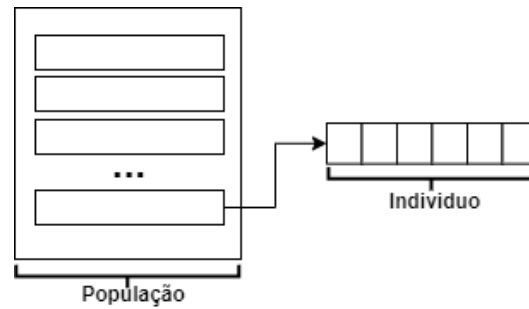


Figura 3 – População de indivíduos

A população inicial de indivíduos é gerada aleatoriamente. Deseja-se que essa população seja suficientemente diversa para permitir uma melhor exploração do espaço de busca, mesmo em problemas complexos, resultando em soluções de melhor qualidade e em uma convergência mais rápida (DIAZ-GOMEZ; HOUGEN, 2007). Diferentes estratégias para a sua inicialização podem ser utilizadas estratégias randômicas e/ou combinadas à heurísticas do problema para aumentar a diversidade e/ou qualidade da população inicial (FRAGA et al., 2021). Como por exemplo, dentre as estratégias têm-se priorização das operações com os tempos mínimos de processamento, utilização de buscas locais ou globais gulosas, priorização de operações que requerem maior tempo. Adicionalmente, é interessante evitar indivíduos duplicados para que seja possível explorar soluções diversas mais rapidamente, desta forma reduzindo o número de gerações necessárias. No caso de distintos problemas restritos, caso seja viável, também podem ser adicionadas heurísticas relacionadas, como não permitir soluções iniciais inválidas.

2.1.3 Função de Aptidão

A função de aptidão (*fitness*), em algoritmos genéticos, é definida como a habilidade de sobrevivência (GOLBERG, 1989). Inspirada nos processos genéticos dos organismos biológicos e no princípio da seleção natural, a função de fitness quantifica e compara a qualidade das soluções candidatas. No contexto de otimização, a aptidão de um indivíduo é o valor calculado da função a ser otimizada para esse indivíduo específico. Esta função deve ser recalculada a cada alteração do indivíduo. Um problema pode ter como característica mais de um objetivo para sua solução, sendo definido, neste caso, como otimização multiobjetivo.

2.1.4 Seleção dos Pais para Cruzamento

A seleção dos pais é responsável por escolher os indivíduos da população atual que participarão do cruzamento, gerando novos indivíduos. Este processo é fundamental para garantir uma maior probabilidade de melhores traços genéticos sejam combinados e propagados nas próximas gerações. Porém, também permitir que indivíduos com avaliações

piores sejam selecionados, assim garantem a diversidade da população. Os operadores de seleção geralmente são estocásticos e estão intimamente associados ao conceito de pressão seletiva. Quanto mais um operador se concentra em selecionar os indivíduos mais aptos da população, maior é a pressão seletiva aplicada. Esta alta pressão seletiva está diretamente ligada à convergência prematura, onde o algoritmo genético pode ficar preso em um mínimo local, impedindo a descoberta de soluções ótimas. Por outro lado, uma pressão seletiva menor permite uma maior diversidade na população, o que pode resultar em um tempo mais longo para a convergência ou até mesmo na dificuldade de encontrar uma solução satisfatória. Esse equilíbrio entre pressão seletiva e diversidade populacional é crucial para o desempenho eficaz dos algoritmos genéticos.

Alguns dos principais métodos de seleção dos pais são:

- ❑ **Roleta:** a ideia do processo de seleção se assemelha a uma roleta que é girada e para aleatoriamente em uma posição, onde os valores de *fitness* mais promissores ocupam um espaço maior com mais chance de serem selecionados. Logo o processo de seleção roleta, proposto por (HOLLAND, 1992), é baseado em probabilidades de seleção proporcionais à aptidão relativa de cada indivíduo. Indivíduos com maior aptidão têm mais chances de serem selecionados, mas indivíduos com menor aptidão também têm uma chance, garantindo diversidade (HOLLAND, 1992).
- ❑ **Torneio:** uma amostra de n indivíduos é selecionada aleatoriamente, ou de acordo com alguma heurística, e dentro dessa amostra os indivíduos competem pela seleção final de um indivíduo para participar da reprodução. O melhor indivíduo dentro desse grupo é selecionado. Este método pode ser ajustado variando o tamanho do torneio, controlando assim a pressão seletiva.
- ❑ **Torneio Estocástico:** Similar ao torneio, mas a seleção do vencedor do torneio é feita probabilisticamente, favorecendo indivíduos com maior aptidão, mas mantendo a possibilidade de escolha de indivíduos menos aptos para promover diversidade.

2.1.5 Operadores Genéticos

Os operadores genéticos atuam como mecanismos de transformação e diversificação da população de soluções. A escolha adequada desses operadores permite a exploração eficiente do espaço de busca e a manutenção das características adaptativas adquiridas ao longo das gerações (GOLBERG, 1989). Os dois operadores genéticos tipicamente usados em um algoritmo genético são:

- ❑ **Cruzamento** (*crossover*): é um processo pelo qual novos indivíduos (filhos) são gerados a partir da combinação do material genético de indivíduos pertencentes à população corrente (pais). Este operador é essencial para a recombinação genética,

permitindo que os descendentes herdem e potencialmente melhorem as características dos pais. A seleção dos pais para o cruzamento é geralmente realizada com base na aptidão dos indivíduos da população do AG, favorecendo a transmissão de materiais genéticos mais promissores à resolução do problema. Um exemplo típico é o cruzamento de um ponto, onde um ponto de corte é selecionado aleatoriamente ao longo dos cromossomos dos pais. Os segmentos dos pais são então trocados após esse ponto, criando dois novos filhos. Outro exemplo é o cruzamento de máscara, onde uma máscara binária é utilizada para determinar quais genes serão trocados entre os pais. A máscara é uma sequência de 0s e 1s do mesmo comprimento que os cromossomos. Um 1 na máscara indica que o gene correspondente do primeiro pai será trocado com o gene do segundo pai, e um 0 indica que o gene será mantido.

- ❑ **Mutação:** introduz variações genéticas aleatórias nos indivíduos, possibilitando o surgimento de novas características que podem melhorar para a adaptação da solução ao problema estudado. Esse tipo de operador é crucial para evitar a convergência prematura e incluir novas informações genéticas na população. Um exemplo típico de operador de mutação é a troca, onde dois genes selecionados aleatoriamente em um cromossomo trocam de posição. Outro exemplo é a mutação randômica, onde um gene é substituído por outro valor escolhido aleatoriamente dentro do seu domínio permitido.

2.1.6 Reinscrição da População

A renovação das populações em algoritmos genéticos é um processo iterativo onde cada geração, composta pela aplicação sequencial de operadores genéticos, busca aprimorar as soluções candidatas. Após a aplicação dos operadores genéticos, o número da população aumentou devido aos indivíduos filhos gerados. Portanto deve-se selecionar apenas um subconjunto da população que formarão a próxima geração. Esta seleção é denominada reinscrição e é crucial para manter a diversidade genética de indivíduos mais aptos, mas também permitir indivíduos menos aptos a fim de evitar a convergência prematura. Alguns métodos de reinscrição são:

- ❑ **Reinscrição pura:** Neste método, todos os indivíduos da nova geração substituem completamente os indivíduos da geração anterior. Embora simples, essa abordagem pode levar à perda de boas soluções encontradas nas gerações anteriores.
- ❑ **Elitismo:** O método elitista garante que os melhores indivíduos da geração anterior sejam mantidos na nova geração. Normalmente, um número fixo de indivíduos com a melhor aptidão é copiado diretamente para a nova população, preservando soluções de alta qualidade. Isso aumenta a pressão seletiva e pode melhorar a convergência do algoritmo.

- ❑ **Escolha dos melhores:** Este método combina indivíduos da geração anterior e novos indivíduos, selecionando os melhores com base na aptidão para formar a nova geração. Essa abordagem assegura que apenas as melhores soluções, independentemente de sua origem, sejam mantidas, promovendo uma população de alta qualidade. Contudo, pode levar a convergência prematura caso haja uma baixa diversidade de soluções.

2.1.7 Critério de Parada

O AG é um processo iterativo, portanto a definição do critério de parada é fundamental para determinar quando o algoritmo deve interromper sua execução. Este critério é importante para evitar que o algoritmo continue a executar indefinidamente e para garantir que os recursos computacionais sejam utilizados de maneira eficiente (SAFE et al., 2004). Estes são alguns exemplos típicos utilizados:

- ❑ **Limite Superior no Número de Gerações:** Este critério estabelece um número máximo de gerações (por exemplo 100 gerações) que o algoritmo pode executar. Uma vez atingido esse limite, o algoritmo para. Este método é simples de implementar e não requer conhecimento prévio sobre o problema específico, mas pode não garantir que uma solução ótima ou satisfatória tenha sido encontrada.
- ❑ **Limite Superior no Número de Avaliações da Função de Aptidão:** Similar ao critério de gerações, este critério estabelece um número máximo de vezes que a função de aptidão pode ser avaliada (Por exemplo 10000 avaliações da função de aptidão). Este método também é fácil de implementar, mas, assim como o critério anterior, não garante a qualidade da solução encontrada.
- ❑ **Baixa Probabilidade de Mudanças Significativas em Gerações Futuras:** Este critério é mais adaptativo e complexo. Ele para o algoritmo quando a probabilidade de obter melhorias significativas nas próximas gerações é muito baixa. Existem duas variantes principais:
 - **Critério Genotípico:** A interrupção ocorre quando a população atinge certos níveis de convergência em relação aos cromossomos. Por exemplo, se 90% da população possui um valor específico em um gene, diz-se que esse gene convergiu. Quando uma porcentagem predefinida dos genes na população converge, o algoritmo para. Por exemplo, o algoritmo para quando 80% dos genes na população convergirem.
 - **Critério Fenotípico:** A interrupção é baseada no progresso alcançado nas últimas n gerações. Se a melhora na aptidão média da população nas últimas n gerações estiver abaixo de um limite predefinido, o algoritmo é interrompido.

Por exemplo: O algoritmo para se a mudança na aptidão média nas últimas 50 gerações for menor que 0,01%.

2.2 Problemas Multiobjetivos

Na otimização com um único objetivo, a escolha de uma solução é relativamente simples e consiste na seleção daquela que alcançou o melhor valor para a função de avaliação. Entretanto, problemas reais costumam demandar a otimização de múltiplos critérios (FONSECA; FLEMING, 1998). Por exemplo, ao comprar um laptop, além do preço, o consumidor também considera fatores como desempenho do processador, capacidade de memória, entre outros. Dada a natureza conflitante entre os critérios analisados, pode existir mais de uma solução ótima onde tentar melhorar um objetivo pode prejudicar outros. No exemplo anterior, equipamentos com maior poder de processamento e memória, são mais caros, enquanto aqueles com menor preço, têm recursos mais limitados. Esse contexto caracteriza um problema otimização multiobjetivo (POM) e sua resolução consiste em adotar estratégias de busca mais elaboradas para encontrar um conjunto com as melhores soluções, considerando todos os critérios avaliados. Comparar soluções torna-se complexo na presença de múltiplos objetivos, pois uma solução pode ser superior em um objetivo enquanto uma segunda solução é superior em outro. Por exemplo, considere um problema de minimização bi-objetivo (f_1 e f_2) ilustrado na Figura 4. Cada possível solução do problema é representada por um ponto x_i ($1 \leq i \leq 7$). Por se tratar de um problema de minimização de ambos os objetivos, a linha pontilhada representa o conjunto das melhores soluções do problema, formado pelos pontos x_1, x_2, x_3, x_4 . Comparando x_1 e x_2 , temos que $f_1(x_1) < f_1(x_2)$, mas $f_2(x_1) > f_2(x_2)$. Isto é, não é possível determinar qual é melhor sem priorizar um dos objetivos (princípio de não dominância).

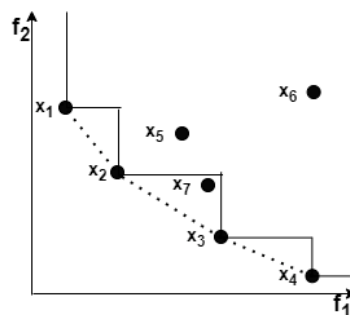


Figura 4 – Conjunto de soluções e fronteira de não dominadas para um problema de minimização com dois objetivos.

2.2.1 Definições

Lidar com vários objetivos demanda uma estratégia mais elaborada, sendo que quanto maior o número de funções objetivo a serem otimizadas, mais complexo é o processo de busca (DEB et al., 2014). Existem diferentes estratégias para a avaliação de soluções em problemas multiobjetivo. Uma alternativa é decompor o Problema de Otimização Multiobjetivo (POM) original em subproblemas mais simples, adotando algum tipo de função de escalarização (ex: ponderação) para transformar os vários critérios em um único valor escalar (FRANCCA; MARTINS; OLIVEIRA, 2018). Outras abordagens multiobjetivo costumam aplicar os conceitos de dominância e fronteira de Pareto para classificar o conjunto de soluções e identificar as melhores (ZITZLER; THIELE, 1998; DEB et al., 2002).

Os conceitos de dominância e fronteira de Pareto são comumente usados para classificar conjuntos de soluções e identificar as melhores soluções (ZITZLER; THIELE, 1998).

Dominância de Pareto

A dominância de Pareto estabelece quando uma solução é melhor que outra, considerando todos os objetivos analisados. Dadas duas soluções, x_1 e x_2 , pode-se dizer que x_1 domina x_2 se, e somente se:

□ x_1 não é pior que x_2 para todos os objetivos j , ou seja, $\neg(f_j(x_1) > f_j(x_2))$ para todos $j = 1, 2, \dots, J$ objetivos (considerando minimização, $<$ denota melhor e $>$ denota pior).

□ e x_1 deve ser melhor que x_2 em pelo menos um objetivo.

Ambas as condições devem ser verdadeiras para caracterizar que x_1 domina x_2 ($x_1 \prec x_2$), caso contrário, x_1 não domina x_2 ($x_1 \not\prec x_2$). Dado um conjunto F com todas as funções objetivo consideradas em um problema de minimização (quanto menor, melhor), o conceito de dominância pode ser formalizado matematicamente como (FRANÇA et al., 2018):

$$x_1 \prec x_2 \Leftrightarrow (\forall (f \in F) f(x_1) \leq f(x_2)) \wedge (\exists (f \in F) f(x_1) < f(x_2)) \quad (1)$$

Fronteira de Pareto

Um conjunto P de soluções não dominadas é chamado de Pareto ótimo (P^*) se não houver outro ponto x_i no espaço de busca que domine qualquer solução em P (DEB, 1999). Em problemas de otimização multiobjetivo, o objetivo é encontrar esse conjunto Pareto ótimo, ou seja, todas as soluções do espaço de busca que não são dominadas por nenhuma outra e cujo mapeamento para o espaço de objetivos forma a Fronteira de Pareto (em problemas mais complexos, alcançar uma boa convergência em direção a P^*). Em

outras palavras, quando não for possível encontrar o conjunto completo dos elementos pertencentes a P^* , os algoritmos multiobjetivos buscam encontrar um subconjunto de P^* e garantir que as demais soluções não dominadas encontradas estejam próximas a P^* .

2.2.2 Algoritmos Evolutivos Multiobjetivos

Os Algoritmos Evolutivos Multiobjetivos (AEMO) foram introduzidos por Schaffer (1985). Uma das principais características que tornam o AG adequados ao contexto multiobjetivo é sua capacidade de evoluir populações de indivíduos, possibilitando que múltiplas soluções não dominadas possam ser encontradas em uma única execução do algoritmo.

Um algoritmo genético multiobjetivo (AGMO) combina a estratégia evolutiva de um algoritmo genético com as técnicas de otimização multiobjetivo. Embora os princípios básicos dos AGs também estejam presentes nos AGMOs, esse último demanda estratégias mais elaboradas de avaliação e seleção de modo a encontrar o conjunto de soluções não dominadas, mais próximo possível da fronteira de Pareto. Nesse contexto, dois princípios básicos que devem ser assegurados pelos AGMOs: convergência e diversidade (DEB et al., 2002). A convergência consiste na capacidade do algoritmo em guiar a busca em direção a regiões promissoras do espaço de busca (Pareto ótimo ou próxima dele). A diversidade consiste em assegurar que soluções encontradas estejam espacialmente bem distribuídas pelas regiões da fronteira. Dentre os AGMOs presentes na literatura, o NSGA II (DEB et al., 2002) e o SPEA2 (ZITZLER et al., 2001) são algoritmos bem consolidados na literatura e que costumam alcançar bons desempenhos em problemas com até 3 objetivos (FRANCCA; MARTINS; OLIVEIRA, 2018), sendo que o primeiro foi utilizado como base para as abordagens implementadas neste trabalho. O Non-dominated Sorting Genetic Algorithm II (NSGA-II) (DEB et al., 2002) é um algoritmo evolutivo multiobjetivo bem estabelecido e bastante usado na literatura (FRANÇA et al., 2017). Ele é uma extensão do NSGA (SRINIVAS; DEB, 1994), que incorpora um processo menos complexo para a ordenação em fronteiras de dominância e um melhor método para a manutenção da diversidade, baseado no cálculo da distância de multidão (*crowding distance*). Com exceção desses dois mecanismos que são essenciais para as etapas que envolvem seleção (seleção de pais e reinserção), o restante do fluxo de execução é semelhante ao de um AG típico. O cálculo da aptidão inicia-se com o ranqueamento dos indivíduos da população em categorias (fronteiras), considerando o conceito de dominância, de modo que todas as soluções não dominadas (melhores indivíduos) são colocadas na primeira fronteira. Em seguida, verifica-se a dominância entre os indivíduos não classificados, sendo a segunda fronteira formada pelas soluções dominadas apenas pelos indivíduos da primeira fronteira. Esse processo se repete até que todos os indivíduos da população estejam em alguma fronteira de dominância, como ilustrado na Figura 5. Esse mecanismo garante que, quanto menor for a fronteira, melhor serão as suas soluções para o problema multiobjetivo.

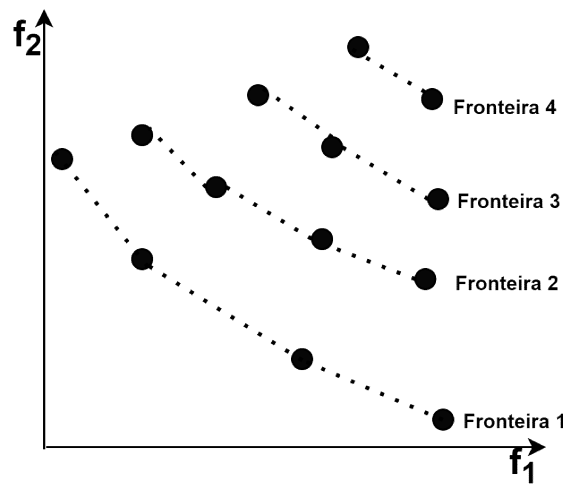


Figura 5 – Classificação dos indivíduos em fronteiras de dominância em um problema de minimização multiobjetivo. Fonte: Adaptado de (DEB et al., 2002).

A diferenciação entre os indivíduos de uma mesma fronteira é feita através do cálculo da *crowding distance*. Essa distância mede a densidade de soluções em torno de cada indivíduo e pode ser vista como uma aproximação do perímetro do cubóide gerado pelos vizinhos mais próximos daquela solução, considerando apenas indivíduos da fronteira. Para ilustrar essa ideia, a Figura 6 mostra o cubóide gerado a partir dos vizinhos do ponto x_i . O seu cálculo proporciona uma medida de distância entre os vizinhos de cada ponto da fronteira, visando atribuir a melhor avaliação para aos pontos mais esparsos, ou seja, ao indivíduo com maior distância em relação às demais soluções da fronteira, privilegiando a diversidade da população.

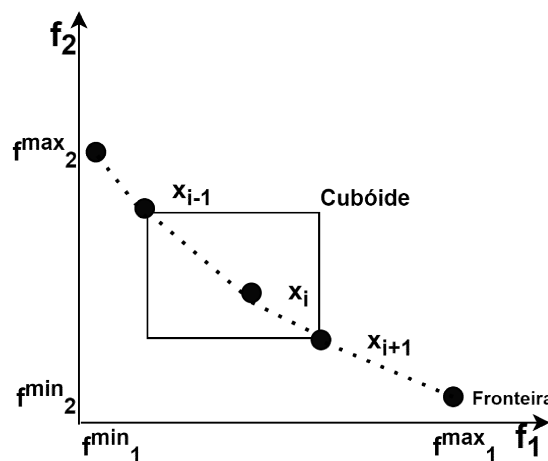


Figura 6 – *Crowding Distance*. Fonte: Adaptado de Deb et al. (2002).

Assim como em um AG típico, o NSGA-II inicia com a geração da população inicial. Cada indivíduo da população é então avaliado, calculando-se sua aptidão para todas as funções objetivo. Com base nesses valores, é realizado o ranqueamento da população em

fronteiras de dominância e o cálculo da crowding distance para os indivíduos de uma mesma fronteira. Em seguida, inicia-se o processo iterativo (gerações), o qual só termina quando o critério de parada é atingido. Os procedimentos realizados a cada geração também são os mesmos executados pelo AG mono-objetivo, envolvendo a seleção de pais, a aplicação dos operadores genéticos de cruzamento e mutações, e a seleção dos indivíduos sobreviventes (reinserção). As diferenças entre as versões mono e multiobjetivo estão nos métodos de seleção. Durante a seleção de pais para o cruzamento, ao comparar duas soluções, utiliza-se a classificação de fronteiras visando minimizar o número de fronteiras. Em caso de empate, é maximizada a *crowding distance*. Caso existam restrições, pode-se adotar uma abordagem que priorize a seleção de soluções de acordo com a fronteira e, posteriormente, pela *crowding distance*. Outro ponto distinto ocorre durante a reinserção, as soluções da primeira fronteira são adicionadas primeiro. Se houver empate, a ordenação é feita visando maximizar a *crowding distance*, garantindo assim uma maior diversidade na população.

ESCALONAMENTO DA PRODUÇÃO FARMACÊUTICA

Neste capítulo, é apresentada a formalização do problema de sequenciamento de bateladas de fármacos na perspectiva multiobjetivo. Em seguida, são apresentados os trabalhos relacionados ao planejamento da produção, com ênfase na aplicação de algoritmos evolutivos multiobjetivos em problemas de otimização com restrição. Por fim, é descrito o algoritmo da literatura que foi adotado como referencial para a análise comparativa das abordagens propostas.

3.1 Sequenciamento de Bateladas

O planejamento de produção no setor farmacêutico é uma tarefa complexa (MAJOZI; SEID; LEE, 2015), em especial devido às diferenças de características intrínsecas do ciclo de vida dos seus produtos, quando comparado com outros setores da indústria química (LAÍNEZ; SCHAEFER; REKLAITIS, 2012). A gestão de processos farmacêuticos, compreende planejamento e alocação de recursos em distintos projetos de fármacos, em períodos de tempo estratégicos (o ciclo de desenvolvimento de fármacos pode levar 7 a 15 anos e sujeito a eventos estocásticos, como a demanda de fármacos (BATTELLE, 2015; MAJOZI; SEID; LEE, 2015; NORMAN, 2016; LOLIC, 2016; HARRER et al., 2019), portanto horizontes de planejamento estratégicos são importantes (por exemplo três anos). O processo de produção de fármacos consiste na manufatura de terapias baseadas em proteínas e o processo pode ser dividido em duas fases principais (HONG et al., 2018):

- Processamento *Upstream* (USP) consiste nos processos necessários para a produção biotecnológica, englobando: desenvolvimento da linha celular, a produção do banco de células central e de trabalho, a inoculação, o cultivo no reator de semente (processo conhecido como semente) e o cultivo em escala de manufatura (JUNGBAUER, 2013).

- Processamento *Downstream* (DSP) envolve as etapas desde a purificação da cultura até produto final purificado, incluindo o crescimento celular, a quebra celular (se aplicável), a alteração de estrutura celular (se aplicável), a captura inicial, a purificação e o refino (JUNGBAUER, 2013).

O planejamento da capacidade produtiva consiste em um problema de otimização de programação de tarefas (*job shop*), tal que o objetivo geral é definir uma sequência de tarefas que podem possuir etapas paralelizáveis ou não, considerando a otimização de determinados objetivos (QUADT; KUHN, 2007). No contexto do problema investigado, espera-se obter uma sequência otimizada (considerando as tarefas de USP e DSP) das quantidades de produtos a serem manufaturados em um horizonte de planejamento de 3 anos, considerando eventos estocásticos de demanda a fim de aumentar a resiliência do planejamento e atender os pedidos dentro do prazo. Tipicamente, manufaturas biotecnológicas são processadas em lote ou bateladas. Uma batelada é definida como uma quantidade específica de produto, sendo a sua produção realizada através de uma ordem única de manufatura durante o mesmo ciclo produtivo (JUNGBAUER, 2013).

Usualmente, a produção é baseada em campanhas (conjuntos de sequências de batelada do mesmo produto), otimizando trocas de produto e reduzindo o risco de uma potencial contaminação entre eles (MAJOZI; SEID; LEE, 2015). Portanto, uma lista de campanhas sequenciais representa um planejamento de produção, conforme exemplificado na Figura 7. Cada retângulo representa a campanha de manufatura de um produto, indicando a quantidade do produto (em quilogramas) a ser produzida e o número de dias de produção. As cores diferenciam os produtos: verde para o produto A, lilás para o B, cinza para o C e vermelho para o D. Conforme ilustrado, a manufatura do produto D ocorre em três campanhas distintas ao longo do planejamento: 66 kg produzidas de janeiro a maio de 2017 (133 dias de produção); 99 kg de fevereiro a agosto de 2018 (175 dias de produção); e 165 kg de abril de 2019 a janeiro de 2020 (259 dias de produção). Os demais produtos seguem a mesma ideia, sendo as respectivas manufaturas organizadas em 2 campanhas (produtos A e C) ou em uma única campanha (produto B).

Note que os parâmetros do processo de manufatura podem variar de produto para produto (ex: rendimento e número mínimo de bateladas), bem como o tempo de paralisação da linha de produção, necessário para a troca entre diferentes produtos.

A Tabela 1 apresenta os parâmetros de produção empregados no escopo deste trabalho. Também são necessários processos de configuração e limpeza dos equipamentos durante a troca de produtos (LAKHDAR et al., 2005). O tempo para a execução desses processos é denominado tempo de troca ou de setup), e varia de acordo com o produto que está sendo fabricado e aquele que será produzido na sequência. A Tabela 2 descreve os tempos (em dias) para as trocas entre os produtos farmacêuticos processados no problema investigado. Em outras palavras, o tempo necessário para o ajuste da linha de produção para

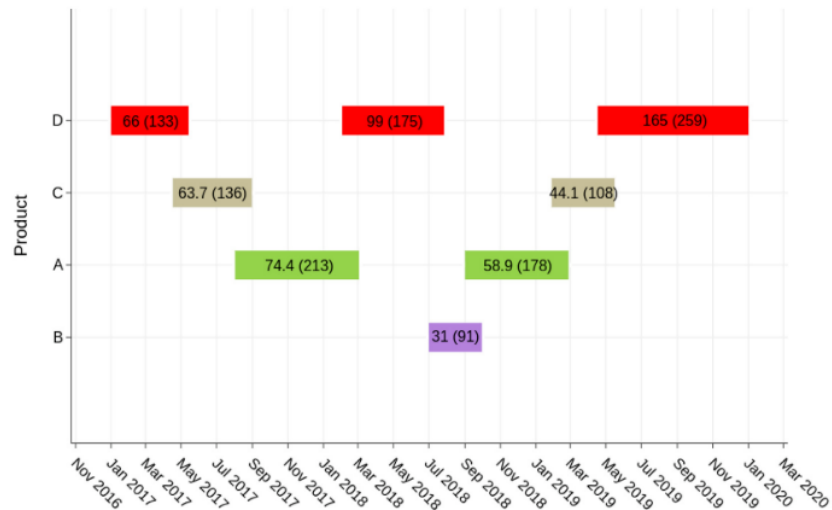


Figura 7 – Exemplo de planejamento de produção. Fonte: (JANKAUSKAS; FARID, 2019).

a fabricação de um novo produto diferente daquele que está sendo produzido atualmente. Os valores apresentados nas Tabelas 1 e 2 foram retirados de (JANKAUSKAS; FARID, 2019) e retratam um cenário real de produção farmacêutica.

Tabela 1 – Parâmetros de manufatura dos produtos farmacêuticos. Fonte: (JANKAUSKAS; FARID, 2019).

Dados de Processo	Produtos			
	A	B	C	D
USP (dias)	45	36	45	49
DSP (dias)	7	11	7	7
QC/QA (dias)	90	90	90	90
Rendimento (kg/batelada)	3,1	6,2	4,9	5,5
Estoque Inicial (kg)	18,6	0	19,6	33
Mínimo número de bateladas por gene	2	2	2	3
Máximo número de bateladas por gene	50	50	50	30
Produção de bateladas em múltiplos de	1	1	1	3
Inoculação	20	15	20	26
Semente	11	7	11	9
Produção	14	14	14	14

Tabela 2 – Tempo de troca (em dias) do produto p para p' . Fonte: (JANKAUSKAS; FARID, 2019).

$p \backslash p'$	A	B	C	D
A	0	10	16	20
B	16	0	16	20
C	16	10	0	20
D	18	10	18	0

3.2 Formulação do Problema Multiobjetivo

Fundamentalmente, o planejamento da produção (ou escalonamento de produção) consiste em estimar a capacidade de produção, o que permite o gerenciamento de recursos de forma distribuída no horizonte de planejamento, a fim de atender as demandas sem atraso. Em outras palavras, a equipe de planejamento busca definir planos de produção que possibilitem otimizar a quantidade produzida; manter níveis de estoque estratégicos e atender demandas dentro do prazo.

Contudo, a manufatura está exposta a eventos dinâmicos que podem interferir significativamente no planejamento da produção e na sua capacidade de atender a demanda dentro do prazo, como, por exemplo, a variabilidade na demanda de fármacos e distintas metas estratégicas de estoque por mês. Dentro desse contexto, algumas métricas relevantes à gestão de estoque são:

- **Produção:** que representa a quantidade total em kg de todos os produtos produzidos no mês.
- **Déficit da meta de estoque:** que representa a diferença entre a meta estratégica de estoque e o valor real.
- **Backlog:** que representa a quantidade de vendas não atendidas, em decorrência de uma demanda superior ao estoque disponível para venda em um determinado momento.

A Figura 8 ilustra como essas métricas são obtidas ao longo de dois meses. A cada mês, calcula-se a quantidade de produto disponível (barra cinza) a partir da soma do estoque já existente, ou seja, o passivo do mês anterior (barra verde) e da capacidade produtiva do mês atual (barra azul). Se o total disponível for superior à demanda mensal (barra amarela), todos os pedidos são atendidos/vendidos (barra roxa), não havendo *backlog* (barra salmão), sendo o restante da produção mantida como estoque para o mês seguinte. Caso contrário, todo o estoque disponível é vendido e as vendas perdidas são contabilizadas como *backlog*.

Como pode ser observado, uma meta estratégica para o estoque mensal é definida (linha vermelha no gráfico). Essa meta é determinada de acordo com o histórico de demandas do mercado e visa reduzir os riscos de falta de produto, sem elevar consideravelmente os custos com produção e armazenamento. A diferença entre essa meta e o estoque do produto no final do mês define a métrica déficit de estoque (barra vermelha).

Neste contexto, o planejamento de produção visa maximizar a produção, de modo a eliminar qualquer risco de *backlog* (restrição do problema), e minimizar o déficit de estoque. Como pode ser observado na Figura 8, percebe-se que os valores de demanda são necessários para o cálculo das métricas de déficit da meta de estoque e *backlog*. Porém,

usualmente essa demanda é baseada em projeções, as quais podem ser obtidas por meio de métodos estocásticos (ex: simulações de Monte Carlo) a fim de aumentar a resiliência de soluções (GAO et al., 2019; OYEBOLU, 2019).

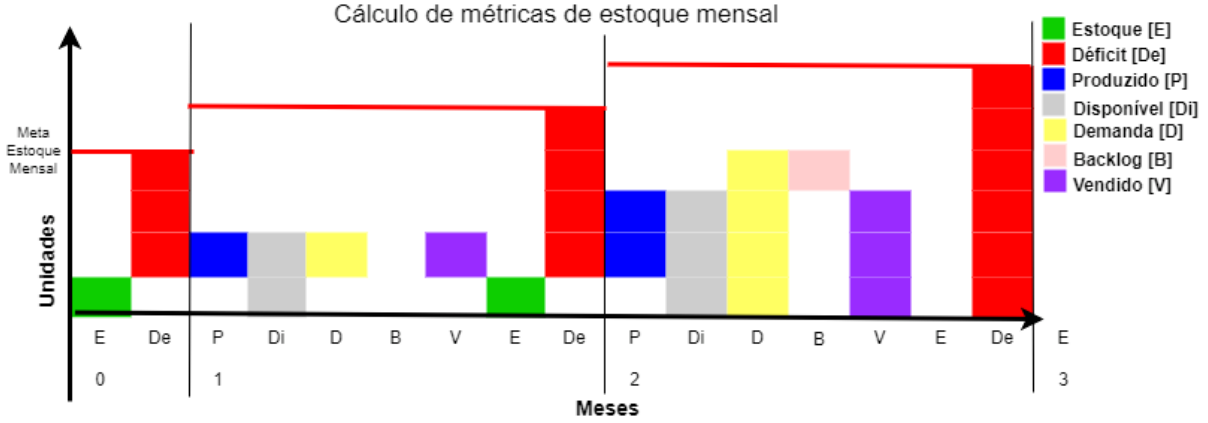


Figura 8 – Evolução das métricas de estoque para um produto, considerando as variações mensais em número de unidades produzidas[P], demanda[D] e *backlog* ao longo de 2 meses.

Investigamos aqui um problema de otimização multiobjetivo com restrições Problema de Otimização Multiobjetivo com Restrições (POMR) aplicado ao sequenciamento de bateladas de fármacos com demandas estocásticas de produtos apresentado em (JANKAUSKAS; FARID, 2019), o qual pode ser formalizado como:

$$\begin{aligned}
 &\text{maximizar } y_1 = f_1(x) \quad \text{e minimizar } y_2 = f_2(x) \\
 &\text{sujeito a } e(x) \leq 0 \\
 &\text{tal que } x = (x_1, x_2, \dots, x_n) \in X \text{ e } y = (y_1, y_2) \in Y
 \end{aligned} \tag{2}$$

Onde $f_i(x)$ representa a função objetivo i ; $e(x)$ é a função de restrição que determina o conjunto de soluções viáveis; x representa o plano de produção, com o sequenciamento de n produtos/bateladas; y é o vetor objetivo; X é o espaço de decisão e Y é o espaço objetivo. Em nossa investigação foram adotados os mesmos dados utilizados em (JANKAUSKAS; FARID, 2019), os quais são baseados em um problema real de manufatura farmacêutica. O primeiro objetivo $f_1(x)$ (Equação (3)) é a soma da quantidade produzida $prod_{(p,m)}$ em quilogramas (kg) para o produto p no mês m , a qual é ilustrada como Produzido na em Figura 8. P e M representam, respectivamente, o total de produtos que podem ser manufaturados na linha de produção ($P = 4$) e a duração, em meses, do planejamento ($M = 36$).

$$f_1(x) = \sum_{p,m}^{P,M} prod_{(p,m)} \tag{3}$$

A quantidade produzida é calculada decodificando o plano de produção considerando os parâmetros do processo de fabricação Tabela 1 e de custo de *setup* entre produtos na linha

de produção Tabela 2. O segundo objetivo $f_2(x, SD)$ (Equação (4)) é a mediana do déficit de estoque ($Def()$ sendo a função do déficit de estoque) avaliados para um determinado conjunto de S de previsões de demanda, obtidas a partir de simulações de Monte Carlo, ou seja, a mediana das diferenças acumuladas entre os déficits mensais (conforme ilustrado por De na Figura 8).

$$f_2(x, SD) = mediana\left\{ \sum_{p,m}^{P,M} Def(prod_{(p,m)}, SD_{(p,m)}^s) \right\}, \quad s \in S \quad (4)$$

Com $SD^s, s \leq S$ representando um determinado conjunto de S simulações de Monte Carlo geradas a partir dos dados gerados utilizando os dados de distribuição triangular (mínimo, moda e máximo) da demanda e metas de estoque mensais para os três anos do planejamento, os quais são apresentados na Tabela 3.

Tabela 3 – Distribuição triangular (mínimo; moda; máximo) para a demanda estocástica e meta de estoque mensal, considerando o período de 3 anos (2017 a 2019).
Fonte: Retirado de (JANKAUSKAS; FARID, 2019).

	Distribuição Triangular da Demanda				Meta de Estoque			
	Produtos				Produtos			
Data	A	B	C	D	A	B	C	D
01/01/2017	0,0	0,0	0,0	0,0	6,2	0,0	0,0	22,0
01/02/2017	0,0	0,0	0,0	(4,5; 5,5; 8,2)	6,2	0,0	4,9	27,5
01/03/2017	(2,1; 3,1; 4,6)	0,0	0,0	(4,5; 5,5; 8,2)	9,3	0,0	9,8	27,5
01/04/2017	0,0	0,0	0,0	0,0	9,3	0,0	9,8	27,5
01/05/2017	0,0	0,0	0,0	(4,5; 5,5; 8,2)	12,4	0,0	9,8	27,5
01/06/2017	(2,1; 3,1; 4,6)	0,0	0,0	(4,5; 5,5; 8,2)	12,4	0,0	9,8	33,0
01/07/2017	0,0	0,0	(3,9; 4,9; 7,3)	(4,5; 5,5; 8,2)	15,5	0,0	19,6	33,0
01/08/2017	(2,1; 3,1; 4,6)	0,0	(3,9; 4,9; 7,3)	(4,5; 5,5; 8,2)	21,7	0,0	19,6	27,5
01/09/2017	(2,1; 3,1; 4,6)	0,0	0,0	(4,5; 5,5; 8,2)	21,7	0,0	14,7	27,5
01/10/2017	(2,1; 3,1; 4,6)	0,0	0,0	0,0	24,8	0,0	19,6	27,5
01/11/2017	0,0	0,0	0	(10; 11; 16,5)	21,7	0,0	19,6	38,5
01/12/2017	(5,2; 6,2; 9,3)	0,0	(8,8; 9,8; 14,7)	(4,5; 5,5; 8,2)	24,8	0,0	19,6	33,0
01/01/2018	(5,2; 6,2; 9,3)	0,0	(3,9; 4,9; 7,3)	0,0	27,9	0,0	14,7	33,0
01/02/2018	(2,1; 3,1; 4,6)	0,0	0,0	(4,5; 5,5; 8,2)	21,7	0,0	19,6	33,0
01/03/2018	(5,2; 6,2; 9,3)	0,0	(3,9; 4,9; 7,3)	(4,5; 5,5; 8,2)	24,8	0,0	19,6	33,0
01/04/2018	0,0	0,0	0,0	(10; 11; 16,5)	24,8	0,0	14,7	33,0
01/05/2018	(2,1; 3,1; 4,6)	0,0	0,0	(4,5; 5,5; 8,2)	24,8	0,0	14,7	27,5
01/06/2018	(8,3; 9,3; 13,9)	0,0	(3,9; 4,9; 7,3)	(4,5; 5,5; 8,2)	27,9	6,2	19,6	33,0
01/07/2018	0,0	0,0	(8,8; 9,8; 14,7)	0,0	27,9	6,2	19,6	33,0
01/08/2018	(5,2; 6,2; 9,3)	0,0	0,0	(4,5; 5,5; 8,2)	27,9	6,2	9,8	33,0
01/09/2018	(5,2; 6,2; 9,3)	0,0	0,0	(4,5; 5,5; 8,2)	31,0	6,2	19,6	38,5
01/10/2018	0,0	0,0	0,0	(4,5; 5,5; 8,2)	31,0	6,2	19,6	33,0
01/11/2018	(5,2; 6,2; 9,3)	(5,2; 6,2; 9,3)	(3,9; 4,9; 7,3)	(10; 11; 16,5)	34,1	6,2	19,6	38,5
01/12/2018	(8,3; 9,3; 13,9)	0,0	(3,9; 4,9; 7,3)	(4,5; 5,5; 8,2)	34,1	6,2	19,6	33,0
01/01/2019	0,0	0,0	0,0	0,0	27,9	6,2	24,5	33,0
01/02/2019	(8,3; 9,3; 13,9)	0,0	(8,8; 9,8; 14,7)	(10; 11; 16,5)	27,9	6,2	34,3	33,0
01/03/2019	(5,2; 6,2; 9,3)	0,0	0,0	0,0	27,9	6,2	24,5	33,0
01/04/2019	(2,1; 3,1; 4,6)	0,0	0,0	(10; 11; 16,5)	27,9	6,2	29,4	44,0
01/05/2019	(5,2; 6,2; 9,3)	(5,2; 6,2; 9,3)	(3,9; 4,9; 7,3)	(4,5; 5,5; 8,2)	34,1	6,2	39,2	33,0
01/06/2019	(2,1; 3,1; 4,6)	0,0	(8,8; 9,8; 14,7)	(4,5; 5,5; 8,2)	34,1	6,2	39,2	33,0
01/07/2019	0,0	0,0	(8,8; 9,8; 14,7)	0,0	31,0	6,2	29,4	33,0
01/08/2019	(8,3; 9,3; 13,9)	0,0	0,0	(10; 11; 16,5)	31,0	6,2	19,6	33,0
01/09/2019	(5,2; 6,2; 9,3)	0,0	(3,9; 4,9; 7,3)	(10; 11; 16,5)	21,7	6,2	19,6	22,0
01/10/2019	(8,3; 9,3; 13,9)	0,0	(8,8; 9,8; 14,7)	0,0	15,5	6,2	14,7	11,0
01/11/2019	(5,2; 6,2; 9,3)	0,0	(3,9; 4,9; 7,3)	(4,5; 5,5; 8,2)	6,2	6,2	4,9	11,0
01/12/2019	0,0	(5,2; 6,2; 9,3)	0,0	(4,5; 5,5; 8,2)	0,0	6,2	0,0	5,5

Devido ao interesse em modelar problemas de otimização reais, técnicas para o tratamento de exceções são necessárias, provocando uma redução do espaço de busca válido.

Em especial, devido à necessidade de otimizar restrições juntamente com múltiplos objetivos, muitas vezes resultando em baixa convergência (DEB; ROY; HUSSEIN, 2020; TIAN et al., 2021). Usualmente algoritmos evolutivos que consideram as restrições utilizam técnicas baseadas no produto escalar de vetores de atribuição de pesos (FONSECA; FLEMING, 1998), tratando a restrição como um objetivo, flexibilizando limites de restrições (TIAN et al., 2021). Dependendo do tipo de problema de otimização, estratégias podem ser adotadas visando ajustar as soluções inválidas para atender as restrições. Entretanto, quando esse ajuste demandar um elevado custo computacional, pode ser mais viável adaptar os operadores genéticos de seleção para cruzamento e/ou reinserção para priorizar as soluções que atendam as restrições do problema. No contexto problema investigado, o ajuste das soluções inválidas é inviável, portanto, adotou-se estratégias de seleção que priorizam indivíduos válidos. Um planejamento eficiente consiste em gerar planos de produção que atendam às demandas durante o período analisado. Nesse contexto, a restrição do problema investigado consiste em garantir que a mediana do *backlog* total não retorne um valor positivo. O *backlog* de um produto (B) corresponde à quantidade de pedidos não atendidos, considerando um determinado conjunto de S Simulações de Monte Carlo (SMC) de demanda. O *backlog* pode ser acumulado ao longo do período planejado e, caso o estoque seja suficiente, a demanda represada pode ser atendida. O cálculo da função restrição $e(x, SD^s)$ é dado por:

$$e(x, SD) = mediana\left\{ \sum_{p,m}^{P,M} B(prod_{(p,m)}, SD_{(p,m)}^s) \right\}, \quad s \in S \quad (5)$$

3.3 Trabalhos Relacionados

Nesta seção, é apresentada uma revisão acerca dos trabalhos que investigam algoritmos evolutivos multiobjetivos na resolução de problemas com restrição, no contexto do escalonamento de produtos farmacêuticos com demandas estocásticas. Inicialmente, aborda-se o uso de algoritmos genéticos multiobjetivo em problemas de escalonamento com restrições. Em seguida, descreve-se as pesquisas que empregam heurísticas para aprimorar a eficiência dos algoritmos de otimização. Adicionalmente, será detalhado o modelo de referência utilizado como base de comparação nesta investigação.

3.3.1 AGMO Aplicados em Problemas Multiobjetivo de Escalonamento com Restrições

Várias pesquisas têm investigado o uso de algoritmos evolutivos multiobjetivo para resolver problemas de escalonamento com restrições (AHMADI et al., 2016; LI et al., 2020; SOUIER; DAHANE; MALIKI, 2019; HAN et al., 2020; AZADEH; GOLDANSAZ; ZAHEDI-ANARAKI, 2016).

Ahmadi et al. (2016) aplicaram o NSGA-II e o Algoritmo Genético de Classificação Não Dominada (denominado NRGa) ao problema de escalonamento de tarefas, considerando quebras de máquinas aleatórias com o objetivo de otimizar o tempo total e a estabilidade. A abordagem desenvolvida considera um crossover baseado em preservação de ordem e as mutações de troca de posições e de máquinas, garantindo que não sejam gerados indivíduos inválidos. Foram realizadas 10 simulações de quebras de máquinas que impactam as máquinas disponíveis. O NSGA-II obteve os melhores resultados na maioria das métricas avaliadas (espaçamento, média de distância ideal, número de soluções válidas na fronteira de não dominados).

Li et al. (2020) aplicaram o NSGA-II ao problema de escalonamento de tarefas considerando a chegada dinâmica de trabalhos e quebras de máquinas. O algoritmo emprega utilizando um cruzamento baseado na manutenção da ordem das operações e mutações de dois e um ponto. Esta metodologia utiliza dois tipos de heurísticas para ajuste de cronograma em resposta a eventos dinâmicos: uma que mantém a atribuição e sequência original das operações não processadas e outra que permite alterações. Os resultados dos testes demonstram que a metodologia é eficaz na adaptação a eventos dinâmicos, sendo a estratégia de alteração de atribuições mais eficiente em termos de economia de energia. O NSGA-II proposto supera os algoritmos MOGA e PSO em desempenho, ajudando na tomada de decisões dos gerentes de produção em relação à economia de energia.

Souier, Dahane e Maliki (2019) investigam o problema de escalonamento em de tarefas com restrições em manufatura re peças, sob cenários dinâmicos relacionadas à chegada aleatória de pedidos de peças e falhas de máquinas, considerando restrições de confiabilidade e manutenção. A abordagem utiliza o NSGA-II para tomar decisões em tempo real sobre a seleção de roteamento de peças, levando em conta a carga de trabalho, nível de utilização e confiabilidade das máquinas em uma estação de trabalho, visando minimizar bloqueios e maximizar a confiabilidade geral do sistema. Para simular os eventos dinâmicos, utiliza-se o modelo de distribuição estatística *Weibull* para gerar as quebras e chegada aleatória de pedidos. Utilizam-se um cruzamento de ponto simples e utiliza-se a mutação de troca de posição. Os resultados da simulação mostram que o algoritmo NSGA-II proposto oferece o melhor desempenho em termos de produtividade, taxa de utilização das máquinas, lucro total baseado em custos de produção e manutenção, e receitas geradas por produto.

Já Han et al. (2020) desenvolveram um AGMO, chamado DEMO, para o problema de escalonamento de tarefas considerando tempos de configuração e os objetivos de minimizar o tempo total e o consumo de energia. Eles utilizaram uma heurística para inicializar a população, operadores de evolução autoadaptativos para obter soluções de alta qualidade e uma estratégia de busca local para melhorar a capacidade de exploração do algoritmo. O operador de busca local é baseado no operador de inserção de um novo gene, gerando-se um determinado número de vizinhos distribuídos uniformemente em torno do indivíduo

selecionado. Por fim compara-se os vizinhos em função da distância em relação à um ponto de referência e seleciona-se o melhor. Os resultados das simulações demonstram que o DEMO supera três algoritmos do estado da arte em termos das métricas de hipervolume, cobertura entre dois conjuntos e distância. Adicionalmente, os autores apontam como pesquisa futura a consideração de cenários mais dinâmicos.

Azadeh, Goldansaz e Zahedi-Anaraki (2016) investigaram a incorporação de uma estratégia para a inicialização de população ao AGMO NSGA-II, que remove operações desnecessárias de cada um dos indivíduos. Utiliza o cruzamento de dois pontos e a operação de mutação de troca e inserção, desta forma permitindo a exploração de resultados. O modelo foi aplicado ao problema de escalonamento de tarefas aplicadas visando minimizar a soma dos tempos de conclusão (*makespan*) de todos os trabalhos e a carga de trabalho das máquinas. Os resultados demonstram que o NSGA-II proposto fornece soluções de alta qualidade em tempo computacional razoável, superando métodos tradicionais como o a restrição *epsilon*.

Estratégias de busca local são frequentemente utilizadas em problemas de otimização, em especial para os casos com avaliações rápidas.

Jia, Gao e Zhang (2020) desenvolveram um algoritmo evolutivo guiado por histórico baseado em decomposição e competição local aplicado ao escalonamento com três objetivos (tempo total, penalidade de atrasos e consumo energético). Foi proposta uma busca local competitiva que avalia a troca de posições de genes e a inserção de genes baseada em otimização de tempos, ajustando as posições das tarefas para acelerar a convergência. O método incorpora uma estratégia de reinserção baseada no elitismo, retendo metade da população e gerando as soluções restantes com base em informações históricas e características extraídas das melhores soluções, garantindo que apenas indivíduos válidos sejam gerados.

Zhang et al. (2022) propuseram um meta-algoritmo para selecionar os parâmetros do AGMO aplicado ao problema de escalonamento *flowshop*, com o objetivo de minimizar o *makespan* e o número total de bateladas. Eles também utilizaram características específicas do problema para codificação e decodificação das soluções, inicialização da população e mutação. A inicialização da população foi realizada utilizando métodos uniformes, aleatórios e mistos. O operador de mutação proposto inclui operações como permutação de genes, trocas de bateladas e distribuição de quantidades de bateladas em posições de genes distintos. Os resultados superam outros AGMO para o mesmo problema em relação às métricas multiobjetivo, como distância generacional invertida e cobertura entre conjuntos.

Fu et al. (2019) investigaram problemas de escalonamento em ambientes de *hybrid flow shop*, onde o tempo de processamento dos trabalhos é incerto e variável. O objetivo é minimizar o *makespan* e os atrasos totais. Para resolver efetivamente esse problema, os autores propõem um algoritmo de otimização multiobjetivo híbrido que mantém duas

populações: uma para busca global em todo o espaço de solução e outra para busca local em regiões promissoras. A busca local é baseada na recozimento simulado. A abordagem emprega uma estratégia coevolutiva que promove um mecanismo de compartilhamento de informações entre as populações, no qual os melhores indivíduos da busca global são incorporados na população da busca local. Segundo os autores, os resultados experimentais mostraram que o algoritmo proposto tem vantagens significativas na resolução do problema investigado.

3.3.2 Uso de Heurísticas no Aprimoramento das Otimizações

Diferentes heurísticas têm sido aplicadas para aprimorar os algoritmos evolutivos em problemas de escalonamento de tarefas, com destaque para as estratégias de inicialização que visam agilizar a convergência e o desempenho global da busca (GAO et al., 2019). Tipicamente, são utilizadas estratégias baseadas em heurísticas que visam a: priorização da alocação com a finalização mais rápida da tarefa, maior número de operações remanescentes (PEZZELLA; MORGANTI; CIASCETTI, 2008), minimização local e global do tempo das tarefas (BAGHERI et al., 2010), priorização do tempo mínimo do trabalho (GAO et al., 2016) e priorização com o maior número de trabalhos remanescentes (BAGHERI et al., 2010).

No caso dos trabalhos que investigam o escalonamento na fabricação de produtos farmacêuticos, heurísticas comuns incluem: tempo de processamento mais curto (priorizando trabalhos com o tempo de processamento mais curto), tempo de processamento mais curto ponderado (considerando a priorização de trabalhos e o tempo de processamento) e escalonamento baseado na data de entrega (CASTILLO, 1992; PACCIARELLI; MELONI; PRANZO, 2011).

Till et al. (2007) concentram-se em abordar restrições em um modelo de dois estágios para o escalonamento de lotes químicos utilizando AG híbridos combinados com Programação Inteira Mista (PIM). Para melhorar a convergência de soluções viáveis, é gerada uma população inicial usando programação de restrições, garantindo apenas soluções iniciais viáveis. Operadores de mutação específicos são projetados, e ajustes de viabilidade de aptidão são utilizados para evitar cálculos de aptidão computacionalmente caros para soluções inviáveis.

Sand et al. (2008) têm como objetivo melhorar o trabalho anterior, proporcionando uma melhor cobertura do espaço de busca viável, utilizando AG que exploram as estruturas específicas do problema. Eles aprimoram a representação do conjunto de soluções viáveis, aproveitando heurísticas específicas do problema e desenvolvendo um novo operador de mutação específico.

Estratégias que exploram localmente regiões do espaço de busca também costumam ser integradas aos algoritmos evolutivos para aprimorar o desempenho na resolução de problemas de otimização (KHOUKHI; BOUKACHOUR; ALAOUI, 2017). Especificamente,

em AGMO aplicados ao escalonamento de lotes, a busca local pode ser empregada para explorar o espaço de soluções de maneira mais eficiente, agilizando a condução a soluções com maior potencial. Wang et al. (2019) investigaram problemas de otimização com restrições e avaliações computacionalmente custosas, propondo uma busca local baseada no método do ponto interior. Em cada iteração, um novo ponto é gerado por meio do gradiente descendente, até que se atinja um critério de parada, resultando em soluções candidatas mais promissoras. Ademais, a busca local é integrada a um modelo surrogate para diminuir o tempo de avaliação, demonstrando melhorias significativas em comparação com benchmarks, inclusive em problemas com restrições não lineares de desigualdade. Habib et al. (2019) desenvolveram um algoritmo evolutivo para problemas restritos com alto custo de avaliação, onde um método de busca local é empregado para refinar indivíduos que infringem as restrições. Seleciona-se um número limitado n de soluções com as menores distâncias euclidianas em relação ao ponto de referência, sendo n definido dinamicamente conforme o número atual de soluções inválidas. Cai, Gao e Li (2020) propuseram estratégias de busca local em AG com uma região de confiança baseada em surrogate. O mecanismo de atualização baseia-se em uma estratégia de particionamento da região adjacente para preservar a diversidade populacional e avaliar a expectativa de melhoria. Múltiplos modelos preditivos conferem robustez à otimização, mostrando eficiência em problemas de alta dimensionalidade.

3.4 Modelo de Referência

O presente trabalho se baseou no problema apresentado em (JANKAUSKAS; FARID, 2019), que considera dados de um cenário real de produção de produtos farmacêuticos. Também é adotada algumas estratégias usadas no algoritmo multiobjetivo proposto pelos autores, principalmente em relação às formas de representação e avaliação de soluções potenciais (indivíduos da população). Ele é referenciado neste documento como modelo de referência e é baseado no algoritmo NSGA-II (DEB et al., 2002), com uma população de 100 indivíduos. Dada a complexidade do problema multiobjetivo com restrição, os autores aumentaram a convergência utilizando um número significativo de gerações (1000) e adotando como solução um conjunto de soluções não dominadas válidas, construído a partir da consolidação de 50 execuções do AGMO.

O indivíduo é um vetor de comprimento variável que representa um plano de produção com a sequência de produtos, sendo cada gene composto pelo produto e pela quantidade de bateladas a serem produzidas na linha de produção.

A Figura 9 ilustra a representação do cromossomo representando um plano de produção, contendo uma sequência de genes, cada gene contém o produto e o número de bateladas de produção. A decodificação do indivíduo gera uma sequência de bateladas a serem produzidas. A Figura 10 ilustra as principais etapas de uma execução do AGMO.

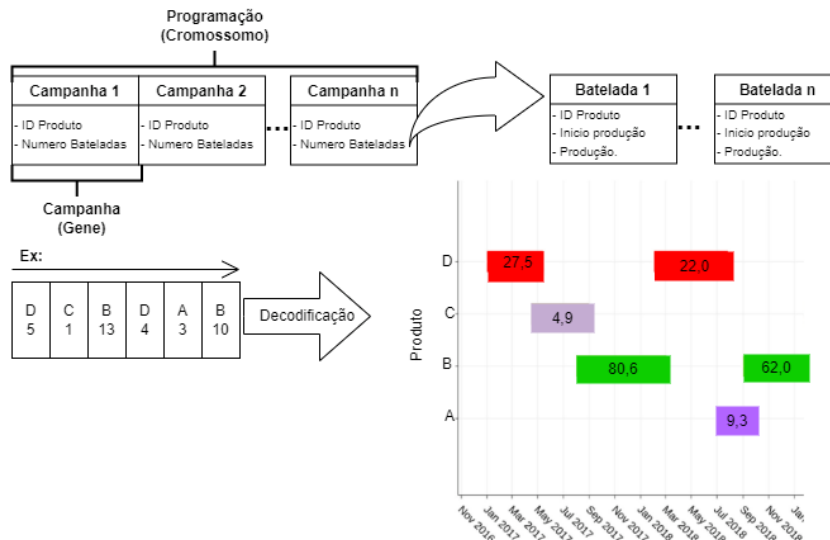


Figura 9 – Representação de um indivíduo e sua decodificação no plano de produção. Fonte: adaptado de (JANKAUSKAS; FARID, 2019).

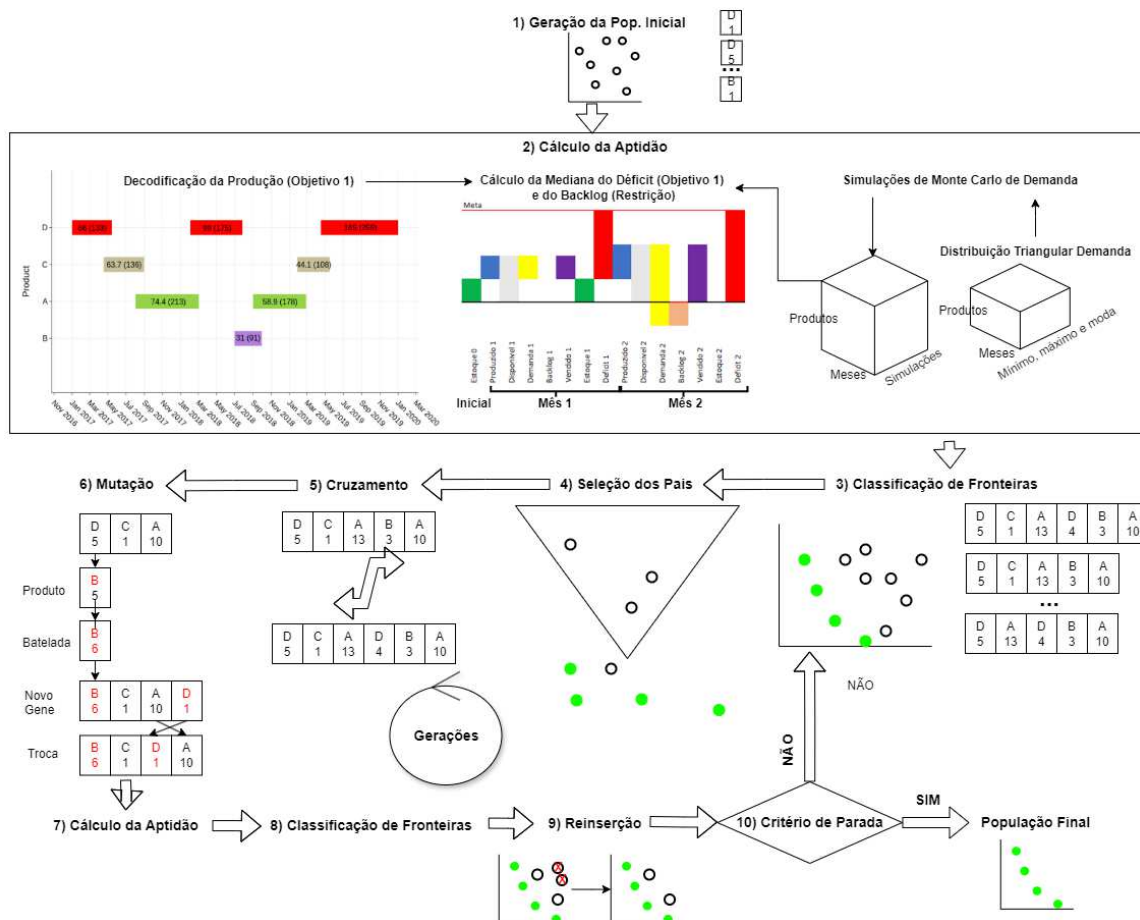


Figura 10 – Fluxo de execução do modelo de referência. Fonte: adaptado de (JANKAUSKAS; FARID, 2019).

A população é inicializada com indivíduos compostos por um único gene, representando a produção de uma batelada de um dos quatro produtos (A, B, C ou D) (etapa 1 da Figura 10). Destaca-se que essa estratégia de inicialização proporciona uma população redundante, sendo majoritariamente formada por indivíduos duplicados. Como já mencionado, o problema de escalonamento investigado visa maximizar a produção total e minimizar o déficit de estoque. Além disto, considera-se a restrição que as soluções geradas (planos de produção) atendam as demandas durante o período analisado ($backlog = 0$). A avaliação da aptidão de cada indivíduo da população (soluções candidatas) consiste no aferição dos objetivos e da restrição, descritos na Seção 3.2, conforme ilustrado na etapa 2 da Figura 10. O cálculo da produção total ($f_1(x)$) inicia-se com a decodificação do indivíduo avaliado em um plano de produção. Por exemplo, na decodificação mostrada na Figura 9, o indivíduo representa uma sequência de produção que começa com 5 bateladas do produto D, seguida, respectivamente, por 1 batelada de C, 13 bateladas de B, 4 bateladas de D, 3 bateladas de A e 10 bateladas de B. O plano de produção é então montado considerando os dados de manufatura descritos na Tabela 1 e os tempos de troca entre produtos da Tabela 2. Destaca-se, ainda, que as etapas de USP e DSP podem ocorrer concomitantemente, ou seja, assim que termina a fase de USP de uma batelada, já começa da produção da batelada seguinte. Da mesma forma, enquanto a última batelada de um produto estiver na etapa de DSP, a linha de produção de USP já pode ser modificada para o próximo produto, reduzindo o efeito do tempo de troca entre os produtos e possibilitando sobreposições das manufaturas dos produtos da sequência, como pode ser observado no plano de produção apresentado na Figura 9. Após a construção do plano de produção, a produção total (em quilogramas) é calculada pela Equação (3), considerando todos os produtos, ou seja, o plano de produção ilustrado na figura tem $f_1(x) = 206,3$ Kg.

No cálculo do segundo objetivo (déficit de estoque ou $f_2(x, SD)$) e da restrição ($backlog$ ou $e(x, SD)$), utiliza-se simulações de Monte Carlo, com base nas distribuições triangulares e metas de estoque descritas na Tabela 3, a fim de gerar um conjunto com 1000 cenários de demanda diferentes ($S=1000$). Então, calcula-se o déficit de estoque e o backlog para cada simulação e os respectivos valores medianos são atribuídos a $f_2(x, SD)$ e $e(x, SD)$, conforme estabelece as Equação (4) e Equação (5).

Após a avaliação, o algoritmo começa o processo iterativo de evolução da população (gerações). Inicialmente, um torneio binário é usado para seleção de pais para o cruzamento, priorizando sequencialmente o atendimento à restrição de backlog, em seguida, a classificação multiobjetivo, ou seja, as fronteiras de dominância e a *crowdingdistance* (DEB et al., 2002). Os indivíduos selecionados são, então, submetidos a um cruzamento uniforme adaptado para lidar com cromossomos de tamanhos diferentes (etapa 5), no qual ordena-se inicialmente os indivíduos em relação ao número de genes. Portanto, nas pri-

meiras gerações não são realizados cruzamentos, sendo a evolução da população limitada às operações de mutação, afetando a diversidade e, conseqüentemente, a exploração do espaço de busca. Os pais são selecionados par a par, onde apenas com 3 ou mais genes realizam o cruzamento. O funcionamento desse cruzamento (*crossover*) é mostrado na Figura 11, onde cada gene pode ser trocado entre pais com igual probabilidade. Caso um pai possua um número de genes maior que o outro, estes genes "excedentes" podem ser copiados, em ambos os filhos, conforme ilustrado na figura.

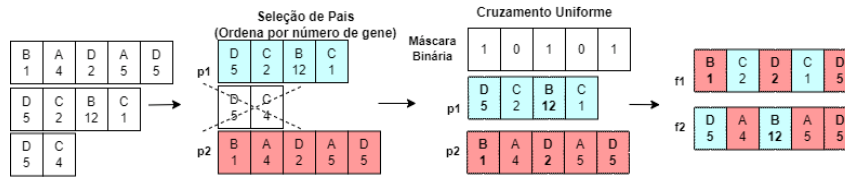


Figura 11 – Cruzamento uniforme proposto no modelo de referência. Fonte: adaptado de (JANKAUSKAS; FARID, 2019).

A mutação (etapa 6) consiste em cinco operações sucessivas, conforme descrito na Figura 12. Três operações são aplicadas por gene, envolvendo a troca do produto por outro escolhido aleatoriamente, o incremento e o decremento de uma batelada. Elas podem ser realizadas ou não, de acordo com as respectivas taxas de ocorrência (p_{MutP} , p_{PosB} e p_{NegB}). Também existem duas operações por cromossomo: trocar posições entre dois genes aleatórios e adicionar um novo gene à sequência. Enquanto a troca de posições ocorre segundo uma taxa p_{TrocaG} , a adição é sempre aplicada a todos os filhos. Por consequência, os indivíduos tendem a aumentar seus tamanhos ao longo das gerações.

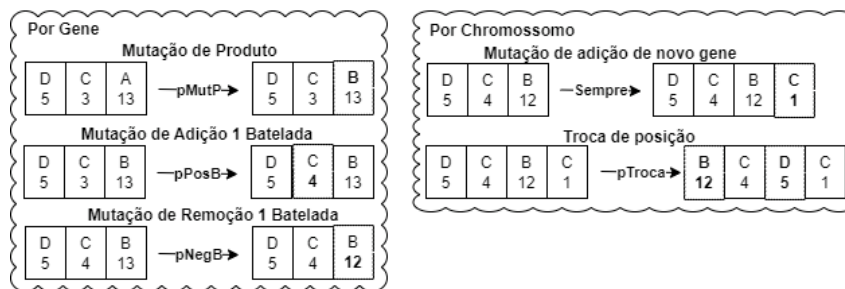


Figura 12 – Operadores de mutação utilizados no modelo de referência. Fonte: adaptado de (JANKAUSKAS; FARID, 2019).

Após a mutação, a população é novamente avaliada e classificada de acordo com as fronteiras de não dominância (etapas 7 e 8). Em seguida ocorre o processo de reinserção (etapa 9), que seleciona os melhores indivíduos com base nos mesmos critérios adotados na escolha dos pais para o cruzamento para formar a população da próxima geração. Isto é, os indivíduos com as menores restrições são priorizados. Em caso de empate, avalia-se a classificação das fronteiras dos indivíduos e a *crowding distance*. O algoritmo evolui a

população por 1000 gerações (critério de parada do processo iterativo) e, ao final, a fronteira de soluções não dominadas é filtrada, selecionando-se apenas indivíduos válidos, que são adicionados a um arquivo externo. Na abordagem apresentada em (JANKAUSKAS; FARID, 2019), dada a dificuldade em encontrar soluções válidas, o algoritmo é executado 50 vezes e ao final, o arquivo externo armazena as soluções não dominadas válidas de todas as execuções, e calcula-se a fronteira de não dominados considerando os indivíduos desse arquivo, a qual é definida como a solução final.

Após a reprodução e o estudo do modelo de referência, foram identificados alguns pontos a melhorar no algoritmo, visando aprimorar a sua convergência, são eles:

- ❑ A baixa diversidade na população inicial, dado que todos os indivíduos são formados por um único gene (produto), escolhidos dentre os quatro possíveis, e considera a fabricação de apenas uma batelada desse produto.
- ❑ O elevado número de gerações e execuções do AGMO demonstram sua dificuldade de convergência, seja pela dificuldade do algoritmo em explorar o espaço de busca de forma eficiente ou pela complexidade do problema em si, devido a restrição imposta.
- ❑ O cruzamento limita a troca de informações genéticas apenas para indivíduos com mais de três genes e entre indivíduos com número de genes semelhantes, o que também atrapalha a exploração eficiente do espaço de soluções.
- ❑ A cada filho gerado, sempre se adiciona um novo gene na mutação, não sendo permitido manter ou reduzir o tamanho de genes.
- ❑ Existe a possibilidade de ocorrerem operações de incremento e decremento da batelada em um mesmo gene, anulando o efeito operações.

As abordagens evolutivas propostas visam tratar essas questões, de modo a melhorar a eficiência da busca e aprimorar a qualidade do conjunto de soluções resultante.

DESENVOLVIMENTO

Neste trabalho, são investigadas novas abordagens evolutivas para lidar com o problema de otimização multiobjetivo com restrição apresentado em (JANKAUSKAS; FARID, 2019), que consiste no sequenciamento de bateladas de produtos farmacêuticos, considerando as demandas de mercado. Inicialmente, foi desenvolvido um algoritmo genético que incorpora métodos clássicos de inicialização e operadores genéticos ao modelo de referência, descrito na Seção 3.4, a fim de melhorar a eficiência e eficácia da busca. Em seguida, é investigada uma nova estratégia de inicialização da população baseada em heurísticas, que visa começar a exploração a partir de uma população de soluções mais promissoras e com maior diversidade. Também são propostos métodos de seleção e operadores genéticos de cruzamento e mutação específicos para o problema estudado, com o objetivo de aprimorar a qualidade do conjunto de soluções não dominadas encontradas. Uma descrição mais detalhada dessas estratégias é apresentada a seguir.

4.1 Algoritmo Baseado em Métodos Clássicos

Inicialmente, foi desenvolvido um algoritmo multiobjetivo, baseado no NSGA-II, que integra métodos e operadores comumente adotados em abordagens evolutivas, com algumas características/estratégias observadas no modelo de referência, como a forma de representação, o método de avaliação dos indivíduos baseado em simulações de Monte Carlo e a estratégia de construção do conjunto de soluções final. Esse algoritmo, denominado ini-rand, é a abordagem base da nossa investigação e foi concebido para tratar algumas limitações identificadas no modelo de referência (descritas na Seção 3.4), principalmente no que se refere à construção de uma população inicial diversificada e formada por soluções, bem como aprimorar a exploração do espaço de busca.

Diferentemente do modelo de referência, cujo a população inicial representa planos de fabricação de uma batelada de um único produto, as estratégias de inicialização desenvolvidas criam soluções que representam sequências de um número variado de produtos e considerando diferentes quantidades de bateladas. A forma como essas sequências com

produtos são definidas possibilita eliminar ou, pelo menos, reduzir drasticamente a ocorrência de indivíduos duplicados na população inicial, aumentando sua diversidade e favorecendo a convergência do AGMO. A inicialização aleatória é um método comumente utilizado na criação de indivíduos em algoritmos evolutivos e visa assegurar a diversidade da população inicial. Nesta abordagem, o número de produtos e suas respectivas bateladas são definidos aleatoriamente para cada indivíduo na população inicial. O tamanho do cromossomo, que corresponde ao número de produtos na sequência, é selecionado a partir de uma distribuição uniforme dentro de intervalos específicos (por exemplo, de 1 a 5 produtos por indivíduo), o produto atribuído a cada gene, bem como o número de bateladas do produto são escolhidos aleatoriamente. No caso do número de bateladas, os valores devem respeitar os limites mínimos e máximos especificados para cada produto (Tabela 1).

No planejamento de produção, as estratégias de cruzamento buscam combinar características promissoras de ambos os pais. Inicialmente, avaliou-se diferentes formas de cruzamento bem estabelecidas na literatura. O cruzamento do modelo de referência estava limitado a casos em que ambos os pais tinham, no mínimo, três genes e possibilita a troca genética apenas entre os genes dentro do tamanho do menor indivíduo. Para aumentar a diversidade da população, exploramos vários métodos de cruzamento que facilitassem a troca de material genético entre pais com diferentes tamanhos e também permitissem o cruzamento independentemente da quantidade de genes do cromossomo, viabilizando o uso de indivíduos menores (1 ou 2 genes). Adicionalmente, avaliou-se a sensibilidade das probabilidades de cruzamento no contexto do cruzamento uniforme, a qual também foi mantida para os métodos de cruzamento baseados em pontos de corte simples ou duplo investigados. No cruzamento de ponto simples é selecionado um único ponto de corte para cada pai, sendo os filhos gerados a partir da combinação genética das partes distintas dos pais. No cruzamento de ponto duplo, são selecionados dois pontos de corte para cada pai e os filhos são criados a partir da troca genética da parte central (genes localizados entre os pontos de corte) dos cromossomos pais. A Figura 13 apresenta as cinco operações de cruzamento baseadas em ponto de corte avaliadas. No cruzamento ponto simples (CPS), o mesmo ponto de corte é selecionado e usado em ambos os pais, enquanto que no cruzamento ponto simples variável (CPSV), um ponto de corte é escolhido especificamente para cada pai. Ambas estratégias de escolha do ponto de corte também são aplicadas nos cruzamentos de ponto duplo. O cruzamento de ponto duplo fixo (CPDF) adota pontos comuns para ambos os pais, enquanto o cruzamento de ponto duplo variável (CPDV) emprega a seleção individualizada por pai. Também foi testada uma variação denominada cruzamento de ponto duplo fixo-variável (CPDFV), que seleciona um ponto fixo e comum para ambos os pais, sendo o segundo ponto específico para cada pai. Todos esses métodos foram implementados e comparados entre si e com o cruzamento uniforme, sendo que o melhor desempenho do algoritmo foi obtido com o cruzamento CPS, usando uma taxa de

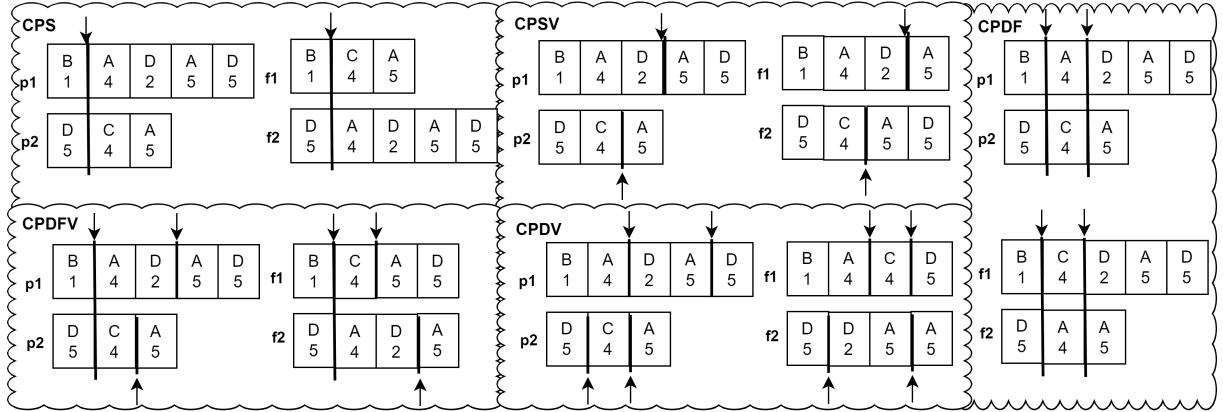


Figura 13 – Operadores de cruzamento baseados em pontos de corte avaliados.

crossover de 30%. Portanto, esse foi o operador de cruzamento adotado na abordagem ini-rand.

As mutações aplicadas no modelo de referência permitem a execução simultânea de operações conflitantes, que podem se cancelar mutuamente. Além disso, elas favorecem o aumento dos cromossomos, dado que realizam a adição compulsória de um gene, sem qualquer opção de remoção. Essa estratégia faz sentido ao considerar que a população inicial gerada pelo algoritmo é formada por soluções incompletas, constituídas por um único produto, e as mutações são usadas para construir incrementalmente as sequências, ao longo da evolução do AG. Entretanto, nossa abordagem base adota uma estratégia de inicialização diferente, gerando uma população com soluções completas. Portanto, foi implementado um conjunto de operadores mais adequados a esse cenário, que combinam mutações de produto e batelada a fim de possibilitar uma maior variação no cromossomo. Esse conjunto é formado por quatro operações sucessivas (Figura 14). Duas delas são executadas por gene: alteração do produto com probabilidade $pMutP$; e variação da batelada com $pMutB$, que possibilita o incremento ou decremento de uma batelada daquele produto. A escolha da ação a ser executada é definida por sorteio aleatório, sendo que o incremento ocorre quando o valor sorteado for menor ou igual a $pAddB$ e o decremento, caso contrário. As demais operações de mutação são aplicadas por cromossomo e incluem: mutação no número de genes com $pMutG$, adicionando um novo gene aleatoriamente, se o número sorteado for até $pAddG$, ou removendo um gene do cromossomo, caso contrário; e a troca de posição entre dois genes, a qual ocorre com probabilidade $pSwapG$. Os valores desses parâmetros de mutação foram escolhidos empiricamente.

4.2 Inicialização Baseada em Heurísticas

Uma das contribuições deste trabalho é a proposta de estratégias de inicialização da população que empregam heurísticas baseadas no tempo de troca entre produtos, sendo

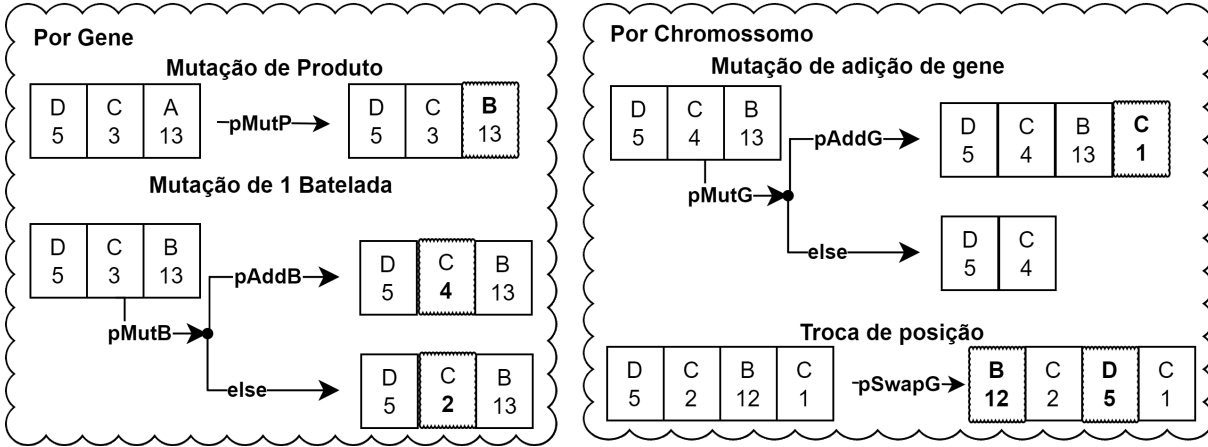


Figura 14 – Operadores de muta  o aplicados nas abordagens evolutivas propostas.

investigadas tr s abordagens: sele  o de produtos (h-iPr), defini  o do n mero de genes inicial (h-iG) e defini  o da quantidade de bateladas (h-iB).

A heur stica de sele  o de produtos prioriza minimizar o tempo gasto na configura  o (*setup*) da linha de produ  o para as trocas entre os produtos (h-iPr). Nessa abordagem, a quantidade de produtos em cada solu  o inicial tamb m   selecionado a partir de uma distribui  o uniforme entre 1 e 5. Entretanto, a sequ ncia de produtos em cada solu  o   determinada usando probabilidades para otimizar os tempos m dios de *setup*. A partir de qualquer aloca  o de um produto p no primeiro gene, o pr ximo produto na sequ ncia do plano (aquele que ser  atribuído ao pr ximo gene)   escolhido de acordo com probabilidades relacionadas   troca de p por qualquer poss vel $p' \in P$. Por exemplo, se p for o produto A, os poss veis produtos no pr ximo gene s o $P = \{B, C, D\}$. Esse processo se repete, considerando o  ltimo gene alocado, at  que se atinja a quantidade de produtos desejado.

Formalmente, a probabilidade de selecionar o produto p' para ser produzido (pr ximo gene) ap s o produto p (gene atual)   denotada por $\Phi_{p \rightarrow p'}$ e seu c lculo   dado por:

$$\Phi_{p \rightarrow p'} = \frac{\Gamma_{p'} - s(p, p')}{\sum_{p' \in P} (\Gamma_{p'} - s(p, p'))} \quad (6)$$

Sendo, $s(p, p')$ o tempo de *setup* necess rio para trocar o produto p por p' na linha de produ  o; e $\Gamma_{p'}$ a soma de todos os tempos de *setup* para qualquer poss vel $p' \in P$, dada por:

$$\Gamma_{p'} = \sum_{p \in P} s(p, p') \quad (7)$$

As probabilidades dos produtos $p' \in P$ s o usadas para escolher sequencialmente os produtos a partir do primeiro produto na sequ ncia, considerando todos os tempos de *setup* do estudo de caso investigado (Tabela 2). Ao utilizar essa heur stica, os produtos com menor tempo de *setup* s o favorecidos em uma escolha probabil stica.

Avaliou-se também a combinação de h-iPr com uma heurística de número de genes, denominada h-iG, buscando melhorar a qualidade dos indivíduos da população inicial e facilitar sua convergência para soluções válidas. A ideia é simples e consiste em utilizar uma distribuição uniforme do número de genes, com um intervalo selecionado a partir de informações históricas de indivíduos válidos. Com base no conjunto histórico criado a partir de experimentos com o modelo de referência e o ini-rand, definiu-se um intervalo de genes entre 7 e 17. Espera-se que indivíduos com tamanhos de genes semelhantes aos indivíduos válidos já encontrados sejam mais promissores e, portanto, devem ser favorecidos.

Similar à abordagem anterior, a heurística para a escolha da quantidade de bateladas (h-iB) também adota um intervalo selecionado a partir de um conjunto histórico de indivíduos válidos, favorecendo a adoção de bateladas semelhantes às aquelas adotadas em sequências já conhecidas. Com base no conjunto criado, é selecionado um número de bateladas entre 6 e 19 unidades para cada produto do cromossomo.

4.3 Cruzamento Produto-Batelada

Considerando as particularidades do problema investigado, um novo operador de cruzamento foi proposto e avaliado durante o desenvolvimento das abordagens evolutivas, com o intuito de aumentar a variabilidade dos indivíduos resultantes (filhos).

Cada gene do indivíduo é composto por um par de informações: produto e número de bateladas. Todos os operadores de cruzamento usados tanto no modelo de referência, quanto na abordagem base (ini-rand), realizam o cruzamento de pares, trocando ambas as informações simultaneamente. Diferentemente dessas abordagens, o operador de cruzamento proposto proporciona o tratamento independente entre produtos e bateladas, permitindo a troca genética dessas informações de forma não simultânea, a fim de aumentar a diferenciação entre os indivíduos gerados e seus pais e, conseqüentemente, a diversidade na população, bem como uma maior exploração do espaço de busca. Para isso, o cruzamento uniforme usado no ini-rand foi adaptado, de modo que, para cada gene tratado no cruzamento, há uma determinada probabilidade de cruzamento de produto (pCP) e uma determinada probabilidade de cruzamento de bateladas (pCB). Assim, para cada operação de cruzamento (produto ou batelada), um número aleatório é sorteado e se ele for menor ou igual a probabilidade correspondente, a troca é realizada. Destaca-se que essa estratégia possibilita, ainda, que não ocorra a troca genética de um determinado gene, quando os números sorteados forem superiores às respectivas probabilidades, ou que todo o gene seja trocado, quando os números sorteados atendem ambas as probabilidades de troca. Em outras palavras, pode ocorrer o cruzamento de produto sem trocar o número de bateladas; o cruzamento do número de bateladas sem alterar os produtos; o cruzamento de ambos os dados (produto e batelada); ou não ocorrer o cruzamento daquele

gene, mantendo a carga genética original.

Adicionalmente, dado que indivíduos com diferentes números de genes podem ocorrer, o gene do indivíduo maior é transferido para o indivíduo menor. Isso possibilita o aumento do número de genes do indivíduo menor e a redução no número de genes do indivíduo maior, contribuindo para a maior diferenciação dos indivíduos gerados.

4.4 Operadores de Mutação

Operações de mutação foram propostas com o objetivo de explorar regiões promissoras do espaço de busca, visando encontrar novas soluções e aumentar a quantidade e qualidade do conjunto de soluções válidas encontradas. Nesse contexto, duas estratégias são investigadas: a mutação heurística e a mutação baseada em busca local, as quais são descritas a seguir.

4.4.1 Mutação Heurística

Durante nossa investigação, foi proposto um novo operador de mutação baseado em heurística, que busca melhorar a convergência do algoritmo. A ideia é substituir as seleções aleatórias por métodos que consideram o conhecimento prévio sobre o problema. Por exemplo, em vez da mutação de produtos ser totalmente aleatória, é adotada a mesma heurística usada na inicialização (h-iPr) para selecionar o produto. A nova estratégia de mutação é formada por operações sucessivas, ilustradas na Figura 15, que podem ser realizadas por gene e por cromossomo. As operações que podem ocorrer por gene são:

- ❑ **Produto:** a mudança do produto ocorre com probabilidade $pMutP$. Caso a mutação ocorra, um novo produto p deve ser selecionado. Na nova estratégia de mutação, existe uma probabilidade $pHeu$ de selecionar o produto utilizando a heurística baseada na otimização dos custos de troca entre os produtos (h-iPr). Esse processo é representado por `produtoSetupOtimizado()` na Figura 15. Quando a probabilidade não é atendida ($1-pHeu$), adota-se a mutação original, onde o produto é selecionado aleatoriamente ilustrado por `produtoRandomico()` na Figura 15.
- ❑ **Número de bateladas:** a mutação do número de batelada ocorre com $pMutB$. Caso ocorra, deve-se selecionar um número n de bateladas para adicionar ou remover. Com probabilidade $pHeu$, seleciona-se um número de bateladas baseado na distribuição gaussiana. Caso contrário, adota-se $n = 1$. Após a escolha do número de bateladas, essa quantidade é adicionada ao gene com probabilidade $pAddB$ ou removida, caso contrário.

As mutações que podem ser realizadas por cromossomo são:

- ❑ **Número de genes:** essa mutação consiste na mudança da quantidade de genes do cromossomo, que pode ocorrer com probabilidade $pMutG$. Um novo gene pode ser adicionado ao final do cromossomo, considerando uma probabilidade $pAddG$ de ocorrência, o qual representa a fabricação de n bateladas do produto p . Caso contrário ($1 - pAddG$), o último gene do cromossomo é removido. Quando ocorre a adição de um novo gene, tanto o produto p , quanto o número de bateladas n , devem ser selecionados. A escolha do produto emprega a mesma estratégia da mutação de produto, ou seja, usa a heurística h-iPr com probabilidade $pHeu$, ou a seleção totalmente aleatória, caso contrário. A escolha das bateladas também é baseada na mutação do número de bateladas, exceto pela seleção aleatória, onde um valor é escolhido entre as quantidades mínima e máxima de bateladas do produto (Tabela 1). Essa seleção aleatória é representada por `randomico(min,max)` na Figura 15.
- ❑ **Troca de genes** Ocorre troca de posições entre dois genes com probabilidade $pSwapG$. Essa é a única mutação que não foi modificada em relação ao conjunto de operações aplicadas em `ini-rand`.

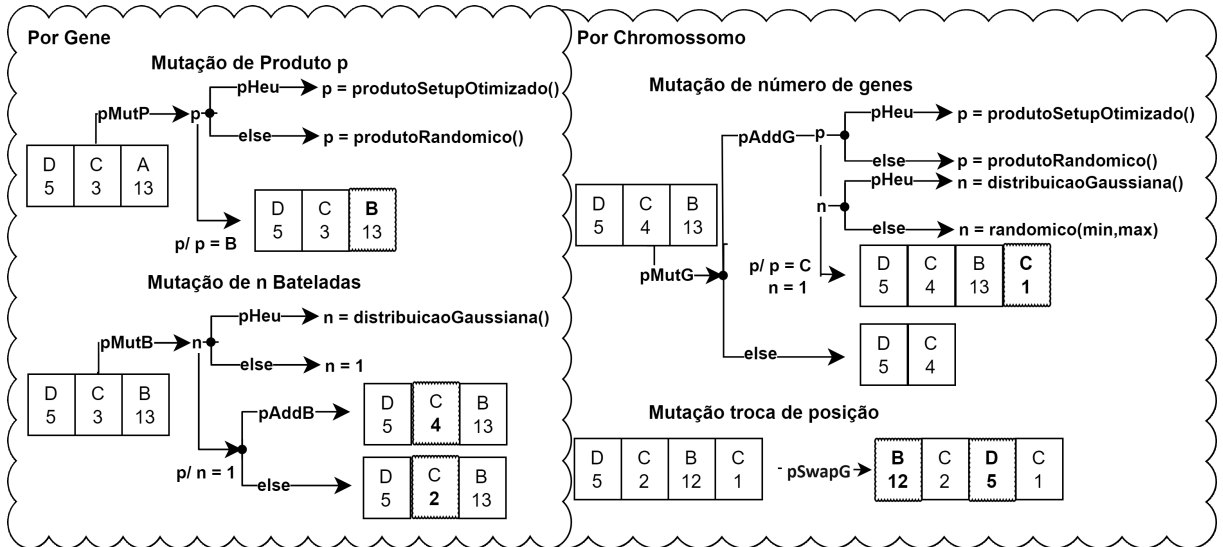


Figura 15 – Operações de mutação baseadas em heurística.

Os valores das distribuições gaussianas para a seleção do número de bateladas a adicionar ou remover, foram definidos em função dos dados de meta de estoque e distribuição de demanda (Tabela 3). Na distribuição gaussiana aplicada na mutação por gene de número de bateladas, utilizou-se $1,142 \pm 1,245$; $1,000 \pm 0,232$; $1,000 \pm 0,877$ e $1,000 \pm 1,044$; para os produtos A,B,C e D, respectivamente. Ao adicionar um novo gene na mutação por cromossomo de número de genes, são adotadas as distribuições gaussianas $12,429 \pm 2,730$; $0,767 \pm 0,409$; $6,343 \pm 1,790$ e $9,483 \pm 0,908$, respectivamente, para os produtos A,B,C e D. Destaca-se, ainda, que o número selecionado deve ser arredondado para um valor

inteiro, uma vez que as bateladas são discretas. Diferentes valores para a probabilidade de utilização de heurística $pHeu$ foram avaliados nos experimentos, a fim de encontrar aquele que maximize a eficiência do AGMO.

4.4.2 Mutação com Busca Local

O problema apresenta uma elevada quantidade de combinações possíveis, decorrente da quantidade de combinações de produtos e, principalmente, da quantidade de bateladas que pode ser produzida para cada produto. A exploração e o aumento do tamanho do indivíduo, através das operações de mutação de adição de novo gene e de alteração no número de bateladas por produto, demonstram-se importantes para a convergência de soluções válidas. Na tentativa de aprimorar ainda mais essa convergência, foi proposta um novo operador de mutação que integra uma estratégia de busca local em substituição à distribuição gaussiana, para selecionar o número de bateladas para os novos genes adicionados. Esse novo operador é ilustrado na Figura 16.

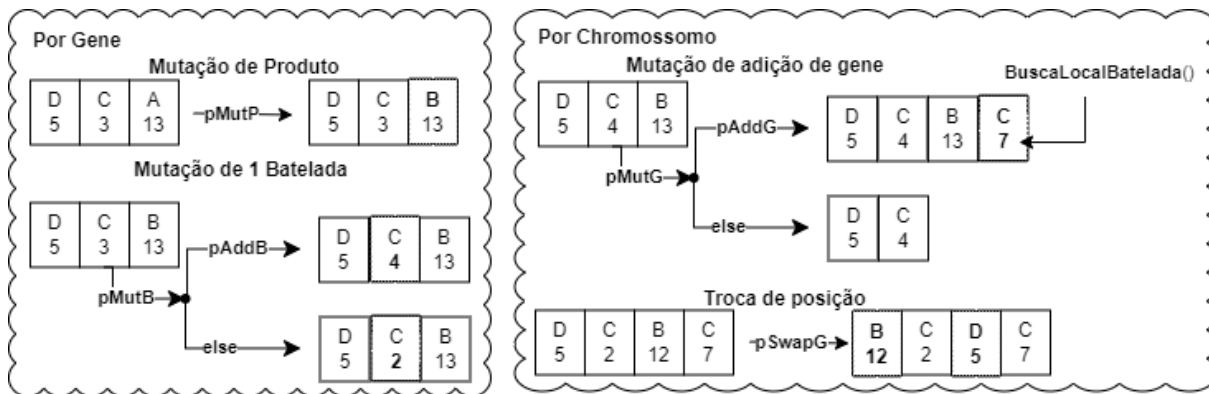


Figura 16 – Operações de mutação aplicadas na abordagem com busca local do número de bateladas.

Como pode ser observado na figura, a busca local do número de bateladas ocorre apenas para o novo gene adicionado ao cromossomo (`BuscaLocalBatelada()` na Figura 16). O funcionamento básico dessa busca local é descrito no Algorithm 1.

Caso ocorra a adição de um novo gene, o produto é selecionado aleatoriamente, enquanto o número de bateladas é escolhido através de uma busca local gulosa. Avalia-se o número médio de bateladas da população atual para o produto selecionado, e calcula-se um percentual $pBM\%$ em relação a esse valor médio para a avaliação do indivíduo no caso de um incremento e decremento de número de bateladas. O indivíduo com o número original de bateladas é então comparado com a variação proposta de incremento e decremento de bateladas considerando o critério de restrição e divisão. Primeiro, avalia-se a menor restrição; no caso de empate avalia-se o resultado da divisão dos objetivos $\frac{f_2(x) * (-1)}{f_1(x)}$ com o objetivo de minimização.

Saída: *Individuo*

Entrada: *individuo*, *pAddGen* (Probabilidade adição de gene), *mediaNumeroBateladas* (Média do número de bateladas), *deltaBateladas* (passo para número de bateladas), *MI* (Número máximo de iterações), *pDM* (Percentual da média do número de bateladas), *minNumBateladas* (Mínimo número de bateladas), *maxNumBateladas* (Máximo número de bateladas)

```
Function buscaBateladas(individuo, mediaNumBateladas, pDM): deltaBateladas = mediaNumBateladas * pDM;
  individuo1 ← copia(individuo); individuo2 ← copia(individuo);
  individuo1.getUltimoGene().setNumBateladas(mediaNumBateladas + deltaBateladas)
  individuo2.getUltimoGene().setNumBateladas(mediaNumBateladas - deltaBateladas)
  if individuo1 < individuo then
    | deltaBateladas = deltaBateladas*-1; individuo = individuo1;
  end
  if individuo2 < individuo then
    | individuo = individuo2;
  else
    | return individuo
  end
  for i ← 1 to MI do
    novoNumBateladas ← novoNumBateladas + deltaBateladas;
    if novoNumBateladas < minNumBateladas OR novoNumBateladas > maxNumBateladas then
      | break;
    else
      | i
    end
    individuo1.getUltimoGene().setNumBateladas(mediaNumBateladas + deltaBateladas) if
      individuo1 < individuo then
        | individuo = individuo1;
      else
        | break;
      end
  end
end
return individuo
```

Algoritmo 1: Procedimento para busca local de número de bateladas no último gene.

Caso as variações no número de bateladas não resultem em uma melhoria em relação ao número médio, a busca é encerrada, mantendo-se o número de bateladas anterior. Este procedimento é repetido com a mesma operação de incremento ou decremento selecionada na primeira iteração até que se atinja o número máximo de iterações (*MI*) ou os limites estabelecidos para o número de bateladas.

4.5 Seleção Particionada

Nesta seção, avaliou-se a partição de conjuntos de indivíduos baseada em critérios de avaliação distintos, visando aprimorar a qualidade das métricas. O critério de seleção original prioriza a ordenação por restrição; em caso de empate, aplica-se o critério do NSGA-II, considerando primeiro a fronteira e, em seguida, a *crowding distance*. Contudo, essa abordagem pode restringir a diversidade e a capacidade exploratória do espaço de busca. Propõe-se, portanto, o particionamento do conjunto selecionado em dois subconjuntos: o primeiro, considerando o critério original (restritivo), e o segundo, aplicando o critério do NSGA-II (não restritivo). Espera-se que isso aumente a diversidade genética da população e potencialize a exploração do espaço de busca, mitigando o risco de convergência prematura. Avaliou-se também a incorporação da seleção particionada baseada em Restrições na seleção de pais e na reinserção.

4.5.1 Reinserção

A estratégia de reinserção particionada proposta visa construir a população P , de tamanho N , que será mantida (sobreviverá) para a próxima geração a partir da junção de duas subpopulações (P_1 e P_2). Os indivíduos de P_1 são selecionados com base no critério original (restrito), que prioriza a restrição do problema (soluções com o menor backlog) em detrimento aos critérios multiobjetivos (fronteiras de não dominância e crowding distance), enquanto os indivíduos de P_2 são selecionados com base no critério do NSGA-II (não restrito). Durante o processo de reinserção, calculam-se as fronteiras de não dominância para todas as soluções, independentemente das restrições. O conjunto P_1 é preenchido selecionando as $N \times pRe$ soluções de acordo com a ordenação pelo original. As soluções remanescentes são reordenadas, utilizando apenas os critérios multiobjetivos, e as $N \times (1 - pRe)$ melhores são selecionadas para formar o conjunto P_2 . Por exemplo, dada uma população com $N = 100$ e um percentual de particionamento $pRe = 60\%$, a cada geração, as 60 soluções com os melhores resultados para a restrição (menor *backlog*) são selecionadas, sendo que, em caso de empate, utiliza-se as fronteiras de não dominância e, em um novo empate, a crowding distance. Os 40 indivíduos restantes são escolhidos considerando exclusivamente os critérios multiobjetivos.

As abordagens evolutivas multiobjetivo resultantes da implantação dos métodos e operadores genéticos propostos nesta pesquisa foram avaliados e comparados, considerando diferentes métricas de desempenho multiobjetivo. Esses resultados experimentais são apresentados e analisados no próximo capítulo.

4.5.2 Seleção de Pais

Similar ao processo de reinserção descrito anteriormente (Subseção 4.5.1), essa estratégia consiste em incorporar a seleção particionada, usada na reinserção, também na seleção de pais para o processo de cruzamento. Assim, durante a realização do torneio, com uma probabilidade pSC , o pai é escolhido entre os indivíduos selecionados com base no critério mais restritivo, o qual prioriza a restrição do problema, empregando os critérios multiobjetivos apenas em caso de empate. Caso contrário, o indivíduo pai é selecionado utilizando apenas os critérios multiobjetivos, sem considerar a restrição. Após experimentos preliminares, adotou-se uma probabilidade de 70% para a seleção mais restritiva e, conseqüentemente, 30% para a seleção mais flexível.

EXPERIMENTOS E RESULTADOS

Neste capítulo é descrito o ambiente de experimentos e as principais investigações das estratégias investigadas. Realizamos vários experimentos utilizando quatro abordagens evolutivas multiobjetivo: o modelo de referência (referido simplesmente como **ref-s1**), bem como as variações propostas e que foram nomeadas como (i) Modelo com inicialização aleatória (**ini-rand**), (ii) o Modelo Heurístico (**ini-heu**), (iii) o modelo **d-pMG40** e (iv) o modelo **d-BLB75**.

O modelo **ini-rand** é avaliado na Seção 5.2 e utiliza a melhor configuração de parâmetros para a inicialização aleatória combinada com a melhor configuração do operador de mutação otimizado e o melhor operador de cruzamento conforme investigado em (Seção 4.1). O modelo **ini-heu** (desenvolvido em Seção 5.3) aprimora o modelo **ini-rand** através da investigação da incorporação de estratégias heurísticas de inicialização heurística (Seção 4.2) e o operador de mutação heurístico (Subseção 4.4.1). Na Seção 5.4 investigou-se os modelos **d-pMG40** e **d-BLB75**. O modelo **d-pMG40** (selecionado em Subseção 5.4.3) incorpora um aumento da probabilidade de cruzamento em contrapartida com uma redução do número de gerações total (conforme Subseção 5.4.1); e a incorporação da reinserção particionada Subseção 4.5.1. O modelo **d-BLB75** (investigado na Subseção 5.4.5) incorpora a busca local de número de bateladas ao operador de mutação conforme descrito em Subseção 4.4.2. Em seguida realizou-se a comparação com um modelo de amostragem aleatória (Seção 5.5). Por fim, avaliou se os principais modelos em um cenário de teste diferente para investigar a sensibilidade do modelo (Seção 5.6).

5.1 Ambiente de Experimentos

Nesta seção, será apresentada a metodologia para o cálculo das métricas, o teste estatístico para comparação de modelos e os ambientes de testes de dados de demanda.

5.1.1 Metodologia para Cálculo das Métricas

Nos experimentos realizados durante a pesquisa, foram consideradas duas categorias de métricas: multiobjetivo e específica ao problema.

Definimos P como o conjunto não dominado obtido do arquivo externo final após uma execução total do algoritmo (composta por 50 rodadas do AGMO, tal que cada rodada é composta por um determinado número de gerações). Devido à natureza estocástica do AGMO, cada configuração de modelo foi executada 30 vezes. Cinco métricas foram avaliadas para fornecer uma análise abrangente do relacionamento de dominância e proximidade com um conjunto de referência P^* : **hv**, **NS**, **E**, **IGD+** e **CS**. Como o conjunto de Pareto ótimo P^* é desconhecido, este foi aproximado experimentalmente. A métrica de hipervolume (hv) mede a distribuição do conjunto de soluções não dominadas em relação ao espaço de busca total. Representa o quanto um conjunto específico de soluções abrange um determinado espaço de soluções. No caso de um problema com dois objetivos, representa a área dominada pelo conjunto P em relação à área total do espaço de busca. A fórmula é a seguinte: A métrica de hipervolume (hv) quantifica a distribuição do conjunto de soluções não dominadas em relação ao espaço de busca total. Para um problema biobjetivo, corresponde à área dominada pelo conjunto P . A fórmula é dada por:

$$hv = \sum_{i=1}^P \sum_{j=1}^M |f_j(s_i) - f_j(w)| \quad (8)$$

onde w é a pior solução encontrada em P , e $f_j(s_i)$ se refere ao valor do objetivo j da solução não dominada s_i . O Número de Soluções viáveis (NS) é total de soluções não dominadas s_i encontradas em P que satisfaz a restrição do problema ($B = 0$). Ou seja, é o número de soluções viáveis pertencentes ao conjunto retornado pelo algoritmo. Esta métrica é está relacionada ao problema e é utilizada para entender a complexidade do problema em questão. Não é uma métrica multiobjetivo, uma vez que não proporciona informações a respeito da qualidade das métricas encontradas. Uma vez que uma boa solução pode dominar outras reduzindo o NS. A taxa de erro (E) representa o percentual de soluções em P que pertencem a um conjunto não dominado ou ótimo de Pareto P^* , formulada como:

$$E = \frac{\sum_{i=1}^{|P|} e_i}{|P|} \quad (9)$$

onde $e_i = 1$ se a solução $i \in P$ é dominada por qualquer solução em P^* , caso contrário, $e_i = 0$. É definida como a porcentagem de elementos no conjunto P^* que não estão contidos em P .

A métrica Cobertura entre Dois Conjuntos CS representa o percentual de elementos no conjunto B que são dominados pelo conjunto A de soluções não dominadas. Por-

tanto, $CS(A,B)=1$ indica que todas as soluções em B são dominadas por A , enquanto $CS(A,B)=0$ indica que nenhum dos elementos em P' é dominado por P (ZITZLER; THIELE, 1999). Calculada como:

$$CS(A, B) = \frac{|\{p \in P | \exists p' \in P' : p' \geq p\}|}{|P|} \quad (10)$$

Além disso, ambas as direções devem ser analisadas, uma vez que $CS(A,B)$ não é necessariamente igual a $1 - CS(B, A)$.

A métrica $IGD+$ (*Inverted Generational Distance Plus*) é uma melhoria das métricas *Inverted Generational Distance* (IGD) e *Generational Distance* (GD), que quantifica e qualifica o conjunto não dominado alcançado usando o AGMO (P) em relação a um conjunto de referência. Tipicamente, P é comparado ao conjunto de Pareto ótimo (P^*).

O $IGD+$ considera as relações de dominância entre a solução e o ponto de referência (P^*). Se uma solução é dominada por um ponto de referência, a distância euclidiana é usada para o cálculo da distância. Além disso, é compatível com o Pareto. No caso em que eles não são dominados um pelo outro, é usada a distância mínima do ponto de referência à região dominada pela solução. Como resultado, a distância é a quantidade de inferioridade da solução em relação ao ponto de referência (ISHIBUCHI et al., 2015). Essa métrica é calculada como:

$$IGD + (A, B) = \frac{1}{|A|} \sum_{s \in A} \min_{r \in B} \left(\sqrt{\sum_{i=1}^M (\max\{s_i - r_i, 0\})^2} \right) \quad (11)$$

Os algoritmos foram implementados utilizando a linguagem Java e executados em máquinas com a seguinte configuração: processador Intel Core i7, 8GB de RAM e sistema operacional Windows 10. Para a análise de dados utilizou-se a linguagem Python, para o cálculo da métrica do hv utilizou-se a biblioteca DEAP (FORTIN et al., 2012) utilizada pelo artigo de referência.

5.1.2 Ambientes de Teste

Nesta seção, são apresentados os ambientes de teste utilizados nos experimentos. Para o cálculo do objetivo $f_2(x, SD^s)$ e da restrição $e(x, SD^s)$ utilizam-se conjuntos de simulações de demanda SD^s . Foram utilizados três ambientes com abordagens diferentes em relação ao conjunto de SD^s utilizadas ao longo do AGMO: Perfis de Conjuntos de Demanda Aleatórios (PCDA), Perfis de Conjuntos de Demanda Estáticos (PCDE) e Perfis de Conjuntos de Demanda Geracionais Aleatórios (PCDGA).

Para o PCDA os conjuntos de SD^s utilizados durante o processo evolutivo do AGMO são gerados aleatoriamente, com um número elevado de 1000 simulações. Adotou-se a mesma quantidade de simulações empregada no artigo original (JANKAUSKAS; FARID, 2019) para assegurar sua reprodução no modelo de referência e facilitar a análise compa-

rativa com as abordagens propostas neste trabalho. Entretanto, uma análise mais aprofundada do número de simulações pode ser realizada, a fim de identificar a quantidade que proporciona o melhor equilíbrio entre a qualidade do planejamento e o custo computacional do algoritmo. Contudo devido à utilização de métricas paramétricas o IGD+ e o E dependentes de uma fronteira P^* , optou-se pelo uso dos ambientes estáticos PCDE para facilitar a análise comparativa entre as abordagens investigadas. Nesse ambiente, todos os conjuntos de SD^s utilizados ao longo do AGMO são estáticos e serão referidos como **s-1**. Dado que a fronteira P^* é desconhecida, avaliamos experimentalmente o P_{PCDE}^* a partir de 42726 rodadas do AGMO, com um total de apenas 32 indivíduos. Isso evidencia a complexidade do problema e a dificuldade em alcançar soluções viáveis. Durante os experimentos iniciais Seção 5.2, utilizamos o ambiente PCDA. No entanto, para manter a comparabilidade na análise, recalculamos todas as métricas paramétricas dos resultados finais da (Seção 5.2) utilizando o P_{PCDE}^* . Além disso, na Subseção 5.2.4, reavaliamos os melhores modelos utilizando o ambiente PCDE durante o processo evolutivo. Nos experimentos das Seção 5.3, Seção 5.4 e Seção 5.5, também utilizamos o ambiente PCDE. Na Seção 5.6, exploramos o ambiente PCDGA, que altera os conjuntos de SD^s a cada 10 gerações. Durante essas 10 gerações, as demandas permanecem estáticas. Essa abordagem busca avaliar a sensibilidade dos modelos em relação a cenários mais dinâmicos. O P_{PCDGR}^* para esse cenário é aproximado a partir de cerca de 2882 rodadas, com um total de apenas 31 indivíduos.

5.1.3 Teste de Hipótese

Os resultados obtidos através de AGMOs são estocásticos, portanto, para realizar comparações de desempenho entre as estratégias investigadas, utilizou-se o comportamento médio de uma população de 30 amostras, sendo que cada amostra representa uma execução independente do algoritmo. Em nossa pesquisa, utilizou-se o teste de hipóteses bilateral para comparar duas estratégias distintas (1 e 2). Nesse contexto, é avaliado se a média de uma métrica qualquer μ (por exemplo, o hipervolume) na população gerada pela configuração 1 (μ_1) pode ser considerada, em média, igual ao μ_2 , obtido a partir da configuração 2. Portanto, as hipóteses são:

$$\square H_0 : \mu_{h1} = \mu_{h2}$$

$$\square H_1 : \mu_{h1} \neq \mu_{h2}$$

No caso em que a H_1 é aceita, realiza-se então o teste de hipóteses para avaliar a inferioridade. Para avaliar a significância de superioridade entre as abordagens para cada métrica, foi utilizado o teste de hipótese com confiança de 5%. Para as distribuições que se aproximam da distribuição normal, foi utilizado o teste T (JOHNSON, 1949), caso con-

trário, foi utilizado o teste Mann-Whitney U (MANN; WHITNEY, 1947). Desta forma compara-se os resultados entre dois testes.

5.2 Modelo com Inicialização Aleatória (ini-rand)

Esta seção descreve os principais experimentos realizados para a criação do modelo ini-rand. Inicialmente, conduzimos experimentos utilizando o ambiente de teste PCDA. Definimos a melhor configuração de parâmetros para a inicialização aleatória (Seção 3.2). Além disso, investigamos a incorporação da melhor configuração do operador de mutação otimizado e o melhor operador de cruzamento (Seção 4.1). Após selecionar a melhor configuração, repetimos o experimento utilizando o ambiente de teste PCDE para comparar o melhor modelo em relação à referência.

5.2.1 Inicialização de Indivíduos Aleatória

Nesta seção, investigamos a inicialização aleatória da população (Seção 4.1). Avaliamos a sensibilidade em relação ao intervalo de número de genes, incorporando a inicialização aleatória ao modelo de referência (ref-s1) (JANKAUSKAS; FARID, 2019). Exploramos três faixas de intervalos de número de genes: 1-5 (**p-ng5**), 1-10 (**p-ng10**) e 1-15 (**p-ng15**). Os resultados são apresentados na Tabela 4. Em relação ao hipervolume (hv), os testes apresentaram médias semelhantes de 0,898. Em relação ao NS, houve uma pequena variação entre os testes, mas a abordagem p-ng15 obteve a melhor média, com uma melhora de 17,5% em relação ao modelo de referência (ref-s1). No entanto, ao avaliar a métrica IGD+, todos os testes apresentam diferenças significativas em relação ao modelo ref-s1, variando de 48,0% a 50,8%. Isso indica que, apesar dos valores de hv e NS semelhantes, o aumento de indivíduos diversos melhora a capacidade do modelo de encontrar soluções mais próximas a P^* . O melhor desempenho foi obtido pelo modelo p-ng10, apresentou uma melhora de 50,8% em relação à referência. Ao analisar os valores de E, observamos valores elevados variando de 95,5% a 99,4%, o que representa a dificuldade de encontrar soluções pertencentes a P^* no espaço de busca. Houve uma redução na taxa de erro com a modificação, sendo que o melhor resultado ocorre para p-ng10, com uma diferença de 3,9% em relação à ref, seguido por p-ng15. Portanto, apesar dos valores semelhantes de hv e NS, houve uma melhoria significativa em relação ao IGD+ para todos os testes, destacando p-ng10 e p-ng5.

Nesta seção, avaliamos a dominância entre as soluções encontradas. A Tabela 5 apresenta os valores da métrica CS(A,B), que quantifica a cobertura de dominância entre dois conjuntos de soluções. As linhas representam os testes A, enquanto as colunas representam os testes B. Para cada combinação, é medida a porcentagem de elementos do conjunto de B que são dominados por A.

Tabela 4 – Métricas avaliação para os testes de inicialização aleatória.

		p-ng5	p-ng10	p-ng15	ref
hv	média	0,898±0,001	0,898±0,0	0,898±0,001	0,897±0,001
	min	0,896	0,897	0,897	0,894
	max	0,900	0,899	0,900	0,898
NS	média	19,879±3,276	18,556±2,751	20,250±2,634	17,233±3,126
	min	15	14	15	11
	max	30	24	28	25
IGD+	média	3,932±0,797	3,718±0,862	3,923±0,656	7,568±1,867
	min	2,502	2,372	2,971	4,765
	max	5,641	6,288	6,176	11,509
E	média	0,973±0,035	0,955±0,032	0,963±0,031	0,994±0,017
	min	0,844	0,875	0,875	0,938
	max	1,000	1,000	1,000	1,000

Tabela 5 – Métrica CS das estratégias de inicialização aleatória.

		p-ng10	p-ng15	p-ng5	ref
A\B					
p-ng10	-	0,439	0,415	0,715	
p-ng15	0,390	-	0,385	0,698	
p-ng5	0,419	0,446	-	0,734	
ref	0,154	0,151	0,143	-	

Avaliando os resultados, observamos que todas as variações p-ng avaliadas apresentam maior dominância em relação à referência. Comparando p-ng5 e ref-s1, notamos que 73,4% dos elementos do conjunto não dominado do modelo de referência são dominados por p-ng5, enquanto apenas 14,3% dos elementos do conjunto não dominado de p-ng5 são dominados por soluções de ref. Em relação às outras variações p-ng5 e p-ng15, a análise é similar, embora os resultados de p-ng10 sejam mais expressivos. Ao comparar p-ng15 e ref, observamos que 69,8% dos elementos do conjunto não dominado de ref são dominados por p-ng15, enquanto apenas 15,1% dos elementos do conjunto não dominado de p-ng15 são dominados por ref. Além disso, ao comparar p-ng5 com p-ng10 e p-ng15, observamos uma melhor dominância de p-ng5. Nesta seção, apresentamos os resultados do teste de hipótese aplicado com um nível de confiança de 95% para avaliar a significância dos resultados obtidos pela melhor variação do modelo com inicialização aleatória em comparação com outros modelos (p-ng5) em comparação com os demais modelos. Os resultados são apresentados na Tabela 6, onde cada linha representa uma métrica (hv e NS para maximização, e IGD+ e E para minimização), e cada coluna representa a comparação do p-ng5 com um modelo diferente. As células verdes indicam que o p-ng5 apresenta melhorias estatisticamente significativas ao outro modelo em termos da métrica correspondente. O teste confirma a superioridade do p-ng5 em relação ao modelo ref em

todas as métricas. Contudo, comparando com o p-ng15, apresenta resultados semelhantes. Ao comparar com o p-ng10, apresenta inferioridade em relação à métrica E.

Tabela 6 – Teste de hipóteses comparando p-ng5 em relação às estratégias de inicialização aleatória.

	p-ng10	p-ng15	ref
hv	=	=	>
NSV	=	=	>
IGD+(P,P*)	=	=	<
E	>	=	<

A partir do teste de hipóteses, observamos que o p-ng5 apresenta evidências de superioridade em relação ao modelo ref e demonstrou uma melhor relação de dominância entre seus pares. Apesar de evidências de inferioridade em relação ao E para o modelo p-ng10, o modelo p-ng5 apresenta resultados semelhantes em relação às outras métricas. Como o modelo original utiliza um número de genes inicial fixo e igual a 1, no modelo de inicialização aleatória será utilizado um número de genes inicial mais próximo do artigo de referência, podendo ser sorteado um número entre 1 e 5 genes para compor o tamanho de cada indivíduo da população inicial. Portanto, a configuração selecionada será o p-ng5 e será utilizada nas seções seguintes.

5.2.2 Mutação

Nesta seção serão apresentados os resultados da investigação do operador de mutação otimizado (conforme Seção 4.4), incorporado ao modelo com inicialização aleatória com a melhor configuração para o sorteio do número de genes obtida na Subseção 5.2.1. Avaliou-se sequencialmente os melhores parâmetros em relação as duas operações que formam o novo método de mutação explicado na seção Seção 4.4: (i) mutação de adição de gene, envolvendo duas probabilidades ($pMutG$ e $pAddG$), (ii) mutação de batelada, envolvendo duas probabilidades ($pMutBat$ e $pAddBat$). Esse ajuste das probabilidades foi feito de forma incremental, ou seja, primeiro definiu-se a melhor configuração para $pMutG$, incorporando esta nova configuração selecionou-se a melhor configuração em relação à $pAddG$, incorporando a melhor configuração de $pMutG$, $pAddG$ e $pMutBat$ assim subsequentemente.

Ajuste do parâmetro $pMutG$

A partir do modelo selecionado em Subseção 5.2.1, avaliou se o parâmetro $pMutGen$ utilizando 5 probabilidades 0% (**p-pMG0**), 7% (**p-pMG7**), 15% (**p-pMG15**), 30% (**p-pMG30**), 50% (**p-pMG50**), 70% (**p-pMG70**); Os resultados obtidos foram comparados com os obtidos pela melhor configuração da Subseção 5.2.1 (ng5).

Tabela 7 – Métricas avaliação para os testes de sensibilidade de pMutG.

	hv média	NS média	IGD+(P,P*) média	E média
p-mMG30	0,895±0,002	14,233±2,542	9,033±2,956	0,997±0,01
p-mMG50	0,885±0,003	9,400±2,027	19,08±6,43	1,0±0,0
p-mMG70	0,861±0,009	6,400±2,581	33,424±9,422	1,0±0,0
p-mMG0	0,0±0,0	0,000±0,0	nan±nan	1,0±0,0
p-mMG7	0,898±0,001	18,900±2,657	5,145±1,468	0,985±0,026
p-mMG15	0,898±0,001	18,333±2,551	5,724±2,009	0,985±0,024
p-ng5	0,898±0,001	19,879±3,276	3,932±0,797	0,973±0,035

Os resultados são apresentados na Tabela 7. Considerando-se as 4 métricas avaliadas (hv, NS, IGD+ e E), dentre os valores de probabilidade investigados (0%, 7%, 15%, 30%, 50% e 70%), o melhor desempenho foi apresentado pelo algoritmo p-pMG7, nas 4 métricas. Entretanto, na comparação com a versão da seção anterior (p-ng5), o desempenho do p-pMG7 foi inferior em 3 métricas (NS, IGD+ e E) e equivalente na métrica hv. A Tabela 8 apresenta os resultados da métrica CS. Ao comparar p-pMG7 e p-ng5, observamos que 27,4% dos elementos do conjunto de p-ng5 são dominados por p-pMG7, enquanto 59,7% das soluções de p-pMG7 são dominadas por p-ng5. Esses resultados indicam que o algoritmo p-ng5 domina mais soluções de p-pMG7, do que o contrário. Por outro lado,

Tabela 8 – Métrica CS das estratégias de sensibilidade de pMutG.

A\B	p-mMG0	p-mMG15	p-mMG30	p-mMG50	p-mMG7	p-mMG70	p-ng5
p-mMG0	-	0,000	0,000	0,000	0,000	0,000	0,000
p-mMG15	1,000	-	0,800	0,975	0,345	1,000	0,232
p-mMG30	1,000	0,124	-	0,912	0,086	0,998	0,055
p-mMG50	1,000	0,008	0,040	-	0,007	0,936	0,004
p-mMG7	1,000	0,510	0,834	0,975	-	1,000	0,274
p-mMG70	1,000	0,000	0,000	0,021	0,000	-	0,000
p-ng5	1,000	0,677	0,905	0,983	0,597	1,000	-

o algoritmo p-pMG com 7% de probabilidade de mutação apresenta maior dominância em relação às outras variações avaliadas. O teste de hipóteses apresentado na Tabela 9 confirma que p-mMG7 apresenta resultados significativamente superiores em relação aos outros testes, exceto para p-pMG15 e p-ng5. Em relação à métrica CS (Tabela 8), o p-mMG7 demonstra melhor dominância em relação ao p-pMG15. Portanto, p-mMG7 foi selecionado como a configuração ideal.

Tabela 9 – Teste de hipóteses para as métricas comparando p-mMG7 em relação à outros

	p-mMG30	p-mMG50	p-mMG70	p-mMG0	p-mMG15	p-ng5
hv	>	>	>	>	=	<
NSV	>	>	>	>	=	=
IGD+(P,P*)	<	<	<		=	>
E	<	<	<	<	=	=

Ajuste do parâmetro pAddGen

Utilizando a configuração selecionada em Subseção 5.2.2, avaliou-se então a sensibilidade em relação à pAddGen com a probabilidade pAddGen variando de 35% (**p-mMG7AG35**), 65% (**p-mMG7AG65**) e 100% (**p-mMG7AG100**). Os resultados

Tabela 10 – Métricas avaliação para os testes de sensibilidade de pAddGen.

		p-mMG7AG35	p-mMG7AG65	p-mMG7AG100	p-ng5
hv	média	0,897±0,001	0,898±0,001	0,898±0,001	0,898±0,001
	min	0,895	0,895	0,898	0,896
	max	0,899	0,900	0,900	0,900
NS	média	18,588±2,787	19,733±2,477	21,433±2,812	19,879±3,276
	min	12	15	15	15
	max	25	25	26	30
IGD+(P,P*)	média	6,146±1,332	4,495±1,456	4,242±1,511	3,932±0,797
	min	3,612	2,425	2,115	2,502
	max	9,171	8,165	7,440	5,641
E	média	0,987±0,026	0,96±0,031	0,935±0,042	0,973±0,035
	min	0,906	0,906	0,844	0,844
	max	1,000	1,000	1,000	1,000

são apresentados em Tabela 10. Em relação ao (hv), todos os testes apresentaram resultados semelhantes, com médias variando de 0,897 a 0,898. Em relação ao (NS), a abordagem p-mMG7AG100 alcançou a maior média, sendo 7,82% maior em comparação com a abordagem p-ng5. Em relação ao IGD+, a abordagem p-ng5 ainda apresenta a melhor média, seguido pela abordagem p-mMG7AG100. Em relação à métrica E, p-mMG7AG100 apresentou o melhor resultado com uma diferença de 3,8% em relação à estratégia p-ng5. Em resumo, a abordagem p-mMG7AG100 se destacou em relação aos outros testes, apresentando melhores resultados nas métricas NS e E quando comparado à abordagem p-ng5.

A análise de dominância é então realizada, conforme resultados na Tabela 11. Ao comparar p-mMG7AG100 e p-ng5, observa-se que 54,5% das soluções de p-ng5 são dominadas por p-mMG7AG100, enquanto 32,6% das soluções de p-mMG7AG100 são dominadas por p-ng5. Logo, p-mMG7AG100 apresenta melhor desempenho em termos de dominância. Adicionalmente, ao comparar p-mMG7AG65 e p-mMG7AG100, verificamos que

p-mMG7AG100 continua a apresentar superioridade em dominância. Especificamente, 25,8% das soluções de p-mMG7AG100 são dominadas por p-mMG7AG65, enquanto 55,3% das soluções de p-mMG7AG65 são dominadas por p-mMG7AG100.

Tabela 11 – Métrica CS das estratégias de sensibilidade de pAddGen.

A\B	p-mMG7AG100	p-mMG7AG35	p-mMG7AG65	p-ng5
p-mMG7AG100	-	0,788	0,553	0,545
p-mMG7AG35	0,099	-	0,188	0,187
p-mMG7AG65	0,258	0,662	-	0,386
p-ng5	0,326	0,712	0,473	-

Avaliando os resultados dos testes de hipóteses em Tabela 12, têm-se que a abordagem p-mMG7AG100 apresenta performance melhor ou igual em relação aos outros testes inclusive o p-ng5. Portanto o p-mMG7AG100 foi selecionado e será utilizado para os próximos testes.

Tabela 12 – Teste de hipóteses para as métricas comparando p-mMG7AG100 em relação à outros

	p-mMG7AG35	p-mMG7AG65	p-ng5
hv	>	>	>
NSV	>	>	>
IGD+(P,P*)	<	=	=
E	<	<	<

Ajuste do parâmetro pMutBat

Utilizando a configuração selecionada em Subseção 5.2.2, avaliou-se então a sensibilidade em relação à pMutB com a probabilidade pMutB para os valores de 0% (**p-mMB0**), 25%(**p-mMB25**), 50%(**p-mMG7AG100**), 75%(**p-mMB75**), 100%(**p-mMB100**). A

Tabela 13 – Métricas avaliação de sensibilidade de pMutBat.

	hv média	NS média	IGD+(P,P*) média	E média
p-mMG7AG100	0,898±0,001	21,433 ±2,812	4,242±1,511	0,935±0,042
p-mMB0	0,895±0,003	11,367±2,632	6,548±2,282	0,998±0,008
p-mMB25	0,899 ±0,001	20,500±2,316	1,994 ±0,427	0,823 ±0,046
p-mMB75	0,897±0,001	17,433±2,849	6,217±1,562	0,979±0,022
p-mMB100	0,894±0,002	13,367±2,748	9,968±2,027	0,998±0,008
p-ng5	0,898±0,001	19,879±3,276	3,932±0,797	0,973±0,035

Tabela 13 demonstra em relação ao hv os testes apresentam hv semelhantes entre 0,894 e

0,899, tal que a abordagem p-mMB25 obteve a maior média de hv. Em relação ao NS, a abordagem p-mMG7AG100 obteve a maior média, seguido por p-mMB25; tal que ambos apresentaram NS superiores 7,8% e 3,1% respectivamente à p-ng5. Avaliando a métrica IGD+, a abordagem p-mMB25 se destacou, apresentando a melhor média de IGD+ com uma melhoria de 49,3% em relação à p-ng5. Comparando a métrica E, a abordagem p-mMB25 também se sobressaiu, apresentando a menor média de E, com uma diferença de 15,0% em relação à p-ng5.

A análise de dominância é realizada com base nos resultados apresentados na Tabela 14. Têm-se que comparando p-mMB25 em relação a p-ng5, logo 81,3% das soluções de p-ng5 são fracamente dominadas por p-mMB25 e 5,6% das soluções de p-mMB25 são fracamente dominadas por p-ng5. Comparando p-mMG7AG100 e p-mMB25, logo 10,2% das soluções de p-mMB25 são dominadas por p-mMG7AG100 e 69,2% das soluções de p-mMG7AG100 são dominadas por p-mMB25. Portanto p-mMB25 apresenta os melhores resultados em relação à p-ng5 e p-mMG7AG100.

Tabela 14 – Métrica CS das estratégias de sensibilidade de pMutBat.

A\B	p-mMB0	p-mMB100	p-mMB25	p-mMB75	p-mMG7AG100	p-ng5
p-mMB0	-	0,632	0,015	0,334	0,098	0,137
p-mMB100	0,228	-	0,003	0,110	0,012	0,023
p-mMB25	0,930	0,982	-	0,902	0,692	0,813
p-mMB75	0,488	0,813	0,024	-	0,090	0,153
p-mMG7AG100	0,791	0,961	0,102	0,800	-	0,545
p-ng5	0,749	0,942	0,056	0,716	0,326	-

Conforme Tabela 15, têm-se que p-mMB25 apresenta evidências de melhor ou igual performance em relação às métricas, exceto em relação à NS. Portanto a configuração p-mMB25 será utilizada.

Tabela 15 – Teste de hipóteses para as métricas comparando p-mMB25 em relação à outros

	p-mMG7AG100	p-mMB0	p-mMB75	p-mMB100	p-ng5
hv	>	>	>	>	>
NSV	=	>	>	>	=
IGD+(P,P*)	<	<	<	<	<
E	<	<	<	<	<

Ajuste do parâmetro pAddBat

A partir do modelo selecionado em Subseção 5.2.2, avaliou-se então a sensibilidade da probabilidade pAddB para os valores: 25%(p-mMAddB25), 50%(p-mMB25) e 75%(p-mMAddB75). Os resultados são mostrados em Tabela 16. Avaliando o hv, ocorreu uma

pequena variação de 0,898 a 0,900, sendo que p-MAddB25 obteve a maior média de hv. Em relação ao NS, a abordagem p-MAddB25 novamente apresentou a maior média, sendo 12,1% superior em relação ao p-ng5. Analisando a métrica IGD+, os testes p-mMB25

Tabela 16 – Métricas avaliação de sensibilidade de pAddBat

		p-mMB25	p-MAddB25	p-MAddB75	p-ng5
hv	média	0,899±0,001	0,900±0,001	0,898±0,0	0,898±0,001
	min	0,898	0,898	0,898	0,896
	max	0,901	0,902	0,899	0,900
NS	média	20,500±2,316	22,286±2,76	19,933±2,449	19,879±3,276
	min	17	14	14	15
	max	26	27	25	30
IGD+(P,P*)	média	1,994±0,427	2,039±0,681	2,178±0,654	3,932±0,797
	min	1,442	1,140	1,142	2,502
	max	3,426	4,663	3,949	5,641
E	média	0,823±0,046	0,815±0,052	0,835±0,051	0,973±0,035
	min	0,750	0,688	0,719	0,844
	max	0,906	0,906	0,938	1,000

obtiveram os melhores resultados, com uma redução de 49,3% e 48,1%, respectivamente, em relação ao p-ng5. Avaliando a métrica E, p-MAddB25 obteve a menor média de E, com uma diferença de 15,8 em relação ao p-ng5. A análise de dominância foi realizada com base nos resultados apresentados na Tabela 16. Ao comparar p-MAddB25 e p-ng5, constatou-se que 81,2% das soluções de p-ng5 são dominadas por p-MAddB25, enquanto 0,6% das soluções de p-MAddB25 são dominadas por p-ng5. Esses resultados indicam uma superioridade significativa de p-MAddB25 em termos de dominância sobre p-ng5. Além disso, ao comparar p-mMB25 e p-MAddB25, verificou-se que 27,8% das soluções de p-MAddB25 são dominadas por p-mMB25, enquanto 30,1% das soluções de p-mMB25 são dominadas por p-MAddB25. Esses resultados sugerem uma vantagem competitiva de p-MAddB25 em relação a p-mMB25 em termos de dominância.

A Tabela 17 demonstra os resultados dos testes de hipóteses para as métricas, indicando que p-MAddB25 apresenta resultados superiores ou iguais aos outros testes. Apresentando melhorias significativas em relação à todas as métricas de p-ng5 e apresentando melhoria significativa em relação ao NS comparando com o p-mMB25.

Tabela 17 – Teste de hipóteses para as métricas comparando p-MAddB25 em relação à outros

	p-mMB25	p-MAddB75	p-ng5
hv	=	>	>
NSV	>	>	>
IGD+(P,P*)	=	=	<
E	=	=	<

Nesta seção apresentou-se os resultados da análise de sensibilidade dos parâmetros da operação de mutação. Portanto a melhor configuração de mutação selecionada foi de pMutG de 7%, pAddG de 100%, pMutB de 25% e pAddBat 25%. Conforme a Tabela 16 a configuração apresentou uma melhora em relação ao p-ng5 de 0,2% no hv, 12,1% para NS, 48,1% para o IGD+ e uma diferença de 15,8% para E.

5.2.3 Cruzamento

Nesta seção será investigados os diferentes tipos de cruzamento e a sensibilidade em relação à probabilidade de cruzamento também será avaliada (conforme Seção 4.1).

Probabilidade de cruzamento

Nesta seção, a partir do modelo de referência (ref), avaliou-se a sensibilidade das probabilidades de cruzamento de 30% (ref), 50% (p-cC50) e 70% (p-cC70). Os resultados estão na Tabela 18. Em relação ao hv médio, as três taxas apresentaram valores semelhantes, variando entre 0,897 e 0,899. No que tange ao (NS), p-cC70 obteve a melhor média, seguido por p-cC50, com uma média 38,4% e 29,4% superior à ref, respectivamente. Analisando o IGD+, a abordagem p-cC70 alcançou a melhor média, seguido por p-cC50, com reduções de 34,1% e 25,4%, respectivamente, em comparação ao teste ref. As melhores performances em relação ao E também foram apresentadas por p-cC70 e p-cC50, com diferenças de 10,9% e 5,2% em relação à ref, respectivamente. Ao avaliar a probabilidade de cruzamento, foi considerado também o tempo de execução devido ao aumento significativo. Analisando o tempo de execução $t(\min)$, a abordagem ref obteve a melhor média, enquanto p-cC70 e p-cC50 apresentaram resultados 182,2% e 114,2% superiores à ref, respectivamente. A análise de dominância foi realizada com base nos resultados apresentados na Tabela 19. Ao comparar p-cC50 e ref, constata-se que 79,4% das soluções de ref são dominadas por p-cC50, enquanto 9,2% das soluções de p-cC50 são dominadas por ref. Isso indica que p-cC50 apresenta uma superioridade significativa em termos de dominância em relação a ref. Da mesma forma, ao comparar p-cC70 e ref, observa-se que 87,8% das soluções de ref são dominadas por p-cC70, enquanto apenas 4,1% das soluções de p-cC70 são dominadas por ref. Novamente, isso evidencia uma melhora significativa de p-cC70 em termos de dominância sobre ref. Por fim, ao comparar p-cC50 e p-cC70 entre si, verifica-se que 22,8% das soluções de p-cC70 são dominadas por p-cC50, enquanto 57,1% das soluções de p-cC50 são dominadas por p-cC70, indicando a superioridade de p-cC50 em relação a p-cC70 em termos de dominância.

A Tabela 20 confirma os resultados esperados para o teste de hipóteses, tal que p-cC70 apresenta melhora significativa em comparação à ref; e em relação à p-cC50 apresenta melhora significativa em re diferenças significativas em relação às métricas comparando

Tabela 18 – Métricas avaliação para a sensibilidade de probabilidade de cruzamento.

		p-cC50	p-cC70	ref
hv	média	0,899±0,001	0,899±0,001	0,897±0,001
	min	0,898	0,898	0,894
	max	0,900	0,903	0,898
NS	média	22,300±2,641	23,857±2,962	17,233±3,126
	min	16	19	11
	max	27	30	25
IGD+(P,P*)	média	5,641±2,295	4,981±1,812	7,568±1,867
	min	2,574	2,375	4,765
	max	12,192	10,221	11,509
E	média	0,942±0,032	0,885±0,052	0,994±0,017
	min	0,875	0,750	0,938
	max	1,000	0,969	1,000
t(min)	média	168,761±9,653	222,326±19,202	78,778±0,238
	min	157,683	197,400	78,517
	max	179,400	249,533	79,567

Tabela 19 – Métrica CS das estratégias ref p-cC50 p-cC70

A\B	p-cC50	p-cC70	ref
p-cC50	-	0,228	0,794
p-cC70	0,571	-	0,878
ref	0,092	0,041	-

com ref. Contudo, apresentam um aumento substancial de 114,2% e 182,2% para o tempo em relação ao modelo ref.

Tabela 20 – Teste de hipóteses para as métricas comparando p-cC70 em relação à outros

	p-cC50	ref-s1
hv	>	>
NS	>	>
IGD+(P,P*)	=	<
E	<	<
t(min)	>	>

Em resumo, os testes p-cC70 e p-cC50 obtiveram melhor desempenho superior em relação a variação ref, apresentando uma melhora percentual para a métrica IGD+ de respectivamente de 34,1% e 25,4%. No entanto, a abordagem ref se destacou no aspecto de tempo de execução, tal que os outros testes obtiveram respectivamente um aumento substancial de 182,2% e 114,2% em relação ao ref. Portanto, para o modelo randômico a configuração selecionada para os testes subsequentes é a manutenção da taxa de cruzamento no valor de 30%.

Tipos de cruzamento com melhor mutação

Nesta etapa, avaliamos diferentes tipos de cruzamento utilizando a taxa de cruzamento selecionada em Subseção 5.2.3 e incorporando ao modelo com a melhor mutação selecionada em Subseção 5.2.2 (teste de referência **p-mMG7AG100**). Avaliamos cinco novas variações de cruzamento: 1) Cruzamento de um ponto fixo (**p-cC1PF**), 2) Cruzamento de um ponto variável (**p-cC1PV**), 3) Cruzamento de dois pontos fixos (**p-cC2PF**), 4) Cruzamento de dois pontos fixos e variáveis (**p-cC2PFV**), e 5) Cruzamento de dois pontos variáveis (**p-cC2PV**). Comparamos as variações com o cruzamento do modelo de referência (**p-MAddB25**).

Tabela 21 – Métricas avaliação para os testes

	hv média	NSV média	IGD+(P,P*) média	E média
p-C1PV	0,897±0,001	15,064±2,407	5,81±2,533	0,991±0,02
p-C1PF	0,899±0,001	20,867±2,543	1,804±0,431	0,806±0,047
p-C2PF	0,899±0,001	21,933±2,69	2,315±1,063	0,83±0,067
p-C2PFV	0,897±0,001	14,433±2,176	7,915±2,699	0,999±0,006
p-C2PV	0,887±0,005	10,133±2,886	18,537±5,109	1,0±0,0
p-MAddB25	0,900±0,001	22,286±2,76	2,039±0,681	0,815±0,052

. Conforme Tabela 21 em relação ao hv, os cruzamentos apresentaram resultados semelhantes, variando de 0,887 a 0,900; tal que a abordagem p-MAddB25 apresentou a melhor média. Quanto ao NS, a abordagem p-MAddB25 também obteve a maior média, seguido pelos testes p-C1PF e p-C2PF. Avaliando o IGD+, a abordagem p-C1PF obteve a melhor média de IGD+, com uma redução de 11,5% em comparação com p-MAddB25. Avaliando E, a abordagem p-C1PF obteve a menor média de E com uma diferença de 0,8%. Portanto, em relação às métricas, o cruzamento p-C1PF apresenta os melhores resultados. A análise de dominância é realizada com base nos resultados apresentados na Tabela 22. Os valores da métrica CS(A,B) indicam a porcentagem de soluções de A que dominam as soluções de B.

Ao comparar p-C1PF e p-MAddB25, temos que 30,9% das soluções de p-MAddB25 são dominadas por p-C1PF, enquanto 27,5% das soluções de p-C1PF são dominadas por p-MAddB25. Portanto, p-C1PF apresenta melhor dominância em relação a p-MAddB25.

Ao comparar p-C2PF e p-MAddB25, observamos que 34,2% das soluções de p-MAddB25 são dominadas por p-C2PF, enquanto apenas 29,9% das soluções de p-C2PF são dominadas por p-MAddB25. Assim, p-C2PF também apresenta melhor dominância em relação a p-MAddB25. Ao comparar p-C2PF e p-C1PF, temos que 28,4% das soluções de p-C1PF são dominadas por p-C2PF, enquanto 36,3% das soluções de p-C2PF são dominadas por p-C1PF. Portanto, p-C1PF apresenta melhor dominância em relação a p-C2PF. Ao ana-

lisar a tabela, podemos concluir que p-C1PF possui melhor dominância em relação a p-MAddB25 e p-C2PF.

Tabela 22 – Métrica CS das estratégias p-MAddB25 p-C1PV p-C1PF p-C2PF p-C2PFV p-C2PV

A\B	p-C1PF	p-C1PV	p-C2PF	p-C2PFV	p-C2PV	p-MAddB25
p-C1PF	-	0,880	0,363	0,971	0,997	0,309
p-C1PV	0,034	-	0,052	0,589	0,946	0,041
p-C2PF	0,284	0,865	-	0,959	0,996	0,299
p-C2PFV	0,007	0,280	0,014	-	0,922	0,008
p-C2PV	0,000	0,020	0,001	0,026	-	0,000
p-MAddB25	0,275	0,872	0,342	0,968	0,998	-

Avaliando a Tabela 23, têm se que p-C1PF apresenta diferença significativa em relação à maior parte dos testes. Em relação à p-C2PF, não apresenta diferença significativa contudo em relação à análise de dominância apresenta melhores resultados. Adicionalmente em comparação com p-MAddB25, a abordagem p-C1PF apresenta resultados significativamente inferiores em relação ao hv e NS. Contudo, avaliando a análise de dominância p-C1PF apresenta os melhores resultados; Portanto o cruzamento da abordagem p-C1PF será utilizado.

Tabela 23 – Teste de hipóteses para as métricas comparando p-C1PF em relação à outros

	p-C1PV	p-C2PF	p-C2PFV	p-C2PV	p-MAddB25
hv	>	=	>	>	<
NSV	>	=	>	>	<
IGD+(P,P*)	<	=	<	<	=
E	<	=	<	<	=

5.2.4 Melhor Modelo Randômico (ini-rand)

Nesta seção, reavaliamos, utilizando o ambiente de experimentos PCDE, o modelo de referência, denominado **ref-s1**, e o melhor modelo encontrado na seção anterior, denominado **ini-rand**. O modelo ini-rand utiliza as mesmas configurações encontradas anteriormente para a inicialização da população randômica com distribuição uniforme, com número de genes variando de 1 a 5; o operador de mutação otimizado com os parâmetros pMutG de 7%, pAddG de 100%, pMutB de 25% e pAddBat de 25% (Subseção 5.2.2); e o cruzamento C1PF com probabilidade de 30%. A Tabela 24 apresenta os resultados comparativos das métricas de desempenho. Em relação ao hv, o modelo ini-rand demonstrou um desempenho ligeiramente superior. Na análise da métrica NS, o modelo ini-rand apresentou uma melhoria de 24,9% em relação ao ref-s1. Para a métrica IGD+,

ini-rand mostrou uma melhoria significativa, com uma redução de 71,9%. Finalmente, na métrica E, ini-rand obteve um desempenho superior com uma diferença percentual de aproximadamente 11,9%. . Conforme Tabela 25 têm se que avaliando ini-rand e ref, têm

Tabela 24 – Métricas avaliação para os testes ref-s1 ini-rand

		ini-rand	ref-s1
hv	média	0,900 $\pm$ 0,001	0,898 $\pm$ 0,001
	min	0,900	0,895
	max	0,903	0,901
NSV	média	24,812 $\pm$ 2,657	19,867 $\pm$ 3,181
	min	19	14
	max	30	26
IGD+(P,P*)	média	2,134 $\pm$ 1,187	7,612 $\pm$ 2,475
	min	1,044	3,497
	max	5,581	12,265
E	média	0,867 $\pm$ 0,068	0,986 $\pm$ 0,026
	min	0,719	0,875
	max	0,969	1,000

se que 87,2% dos elementos de ref são fracamente dominados por ini-rand, enquanto 6,5% dos elementos de ini-rand são fracamente dominados por ref. Confirmando que ini-rand apresenta superioridade em relação à dominância .

Tabela 25 – Desempenho do ini-rand em relação ao modelo de referência.

	ini-rand	ref-s1
A\B		
ini-rand	-	0,872
ref-s1	0,065	-

Conforme a Tabela 26, têm se que há significância estatística da melhoria de ini-rand em relação às métricas. Portanto a combinação das estratégias de inicialização, mutação e cruzamennto propostas proporcionaram uma melhora na qualidade das soluções de 71,9% em relação ao IGD+, 24,97% em relação à NS, diferença de 11,6% em E. .

Tabela 26 – Análise comparativa entre o ini-rand e o modelo de referência.

	ref-s1
hv	>
NSV	>
IGD+(P,P*)	<
E	<

Análise de gerações

Nesta seção avaliou-se a evolução das métricas ao longo das gerações comparando ref-s1 com o ini-rand. As Figuras 17 e 18 apresentam a evolução das métricas hv, E, IGD+ e NS por geração.

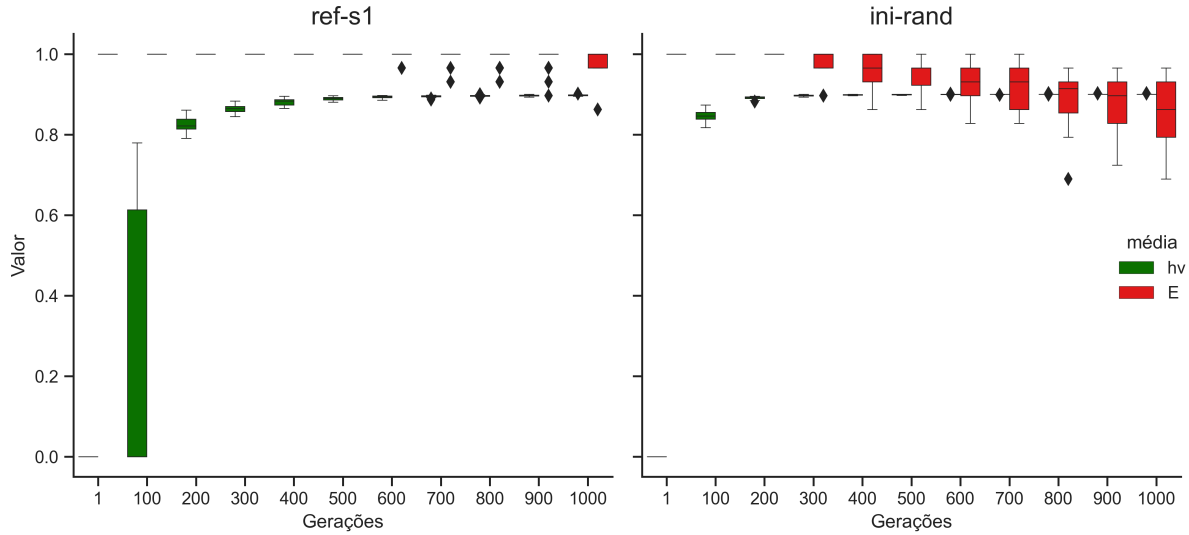


Figura 17 – Diagrama de caixas para as métricas hv e E.

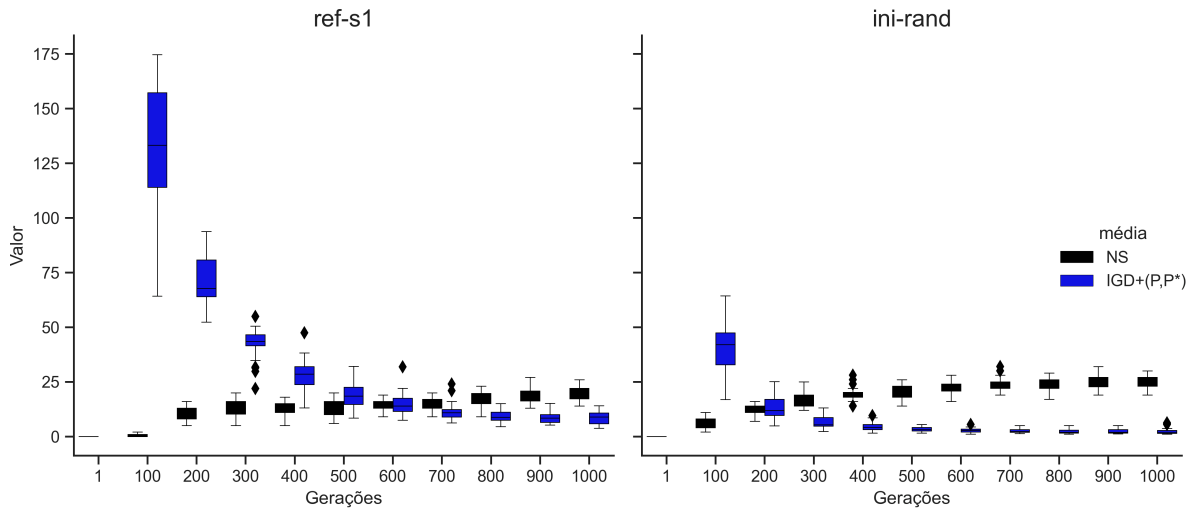


Figura 18 – Diagrama de caixas para as métricas IGD+ e NS.

Os valores numéricos são apresentados na Tabela 27. Para a métrica hv, ini-rand apresenta uma melhora significativa na geração 100, com um aumento de 287,9% em relação ao ref-s1 (melhora da média de 0,218 para 0,847); adicionalmente ini-rand encontra uma média superior à ref-s1 na geração 1000 logo na geração 400. Avaliando a métrica E, observa-se que a partir da geração 300, ini-rand apresenta uma média equivalente ao

erro do ref-s1 na geração 1000. Em relação ao IGD+, na geração 100 ini-rand mostra uma redução de 69,7% comparado ao ref-s1. A partir da geração 300, ini-rand consegue alcançar valores de IGD+ inferiores aos obtidos pelo ref-s1 na geração 1000. Por fim, ao avaliar a métrica NS, ini-rand demonstra uma melhoria substancial de 1384,3%, elevando a média de 0,400 para 5,938. A partir da geração 500, a média de NS de ini-rand supera a do modelo ref-s1 na geração 1000. Portanto, ini-rand apresenta uma convergência mais rápida em termos de número de gerações, indicando que a melhoria do modelo permite reduzir o número de gerações para 600 sem comprometer a qualidade dos resultados em relação ao ref-s1, além de reduzir o tempo de execução.

Tabela 27 – Métricas avaliação para os testes

testName	gen média	100	200	300	400	500	600	700	800	900	1000
ini-rand	E medio	1,000	1,000	0,985	0,960	0,947	0,921	0,923	0,895	0,878	0,853
	IGD+ medio	40,259	13,230	6,573	4,574	3,340	2,764	2,567	2,311	2,296	2,357
	NS medio	5,938	12,281	16,781	19,438	20,344	22,344	24,094	23,844	24,938	24,812
	hv medio	0,847	0,891	0,897	0,899	0,899	0,900	0,900	0,900	0,900	0,900
ref-s1	E medio	1,000	1,000	1,000	1,000	1,000	0,999	0,994	0,992	0,991	0,985
	IGD+ medio	133,034	70,976	42,410	28,075	19,026	14,689	11,381	9,311	8,847	8,557
	NS medio	0,400	10,167	13,200	12,600	12,900	14,467	14,933	17,433	18,700	19,933
	hv medio	0,218	0,824	0,864	0,880	0,889	0,893	0,895	0,896	0,897	0,898

5.3 Modelo Heurístico (ini-heu)

Nesta seção investigamos a criação do Modelo Heurístico (ini-heu), que busca aprimorar o ini-rand através da incorporação de estratégias heurísticas. Investigou-se a incorporação em **ini-rand** das estratégias de inicialização baseada em heurística (Seção 4.2) e o operador de mutação heurístico (Subseção 4.4.1). Utilizou-se o ambiente de teste PCDE.

5.3.1 Estratégias de Inicialização da População

Nesta seção avaliou-se estratégias de inicialização da população. Avaliou-se três alterações incrementais sequenciais: 1) ini-heu incorpora à ini-rand a utilização de probabilidades de seleção de trocas entre produto baseados em menores custos de trocas (h-iPr Seção 4.2); 2) ini-h-pg adiciona o número de genes inicial variando entre determinados intervalos (7-17)(h-iG Seção 4.2); 3) ini-h-pb adiciona também o número de bateladas dentro dos intervalos (6 e 19 unidades)(h-iB Seção 4.2).

Tabela 28 – Métricas avaliação para os testes de inicialização de população heurísticas

	hv média	NSV média	IGD+(P,P*) média	E média
ini-heu	0,9±0,0	26,800±3,498	2,021±0,825	0,852±0,065
ini-h-pg	0,9±0,0	25,460±2,512	2,166±1,01	0,876±0,053
ini-h-pb	0,9±0,0	25,971±3,252	1,947±0,729	0,858±0,063
ini-h-pgb	0,901±0,001	27,032±2,983	1,888±0,794	0,859±0,055
ini-rand	0,9±0,001	24,812±2,657	2,134±1,187	0,867±0,068

A Tabela 28 apresenta os resultados das métricas avaliadas. Em relação à métrica hv, todos os modelos apresentaram médias próximas de 0,9. Para a métrica NS, o modelo ini-h-pgb apresentou um aumento de aproximadamente 8,9% em relação ao modelo ini-rand, seguido por ini-heu com um aumento de 8,0%. Na métrica IGD+, o modelo ini-h-pgb obteve o melhor desempenho, com uma redução de aproximadamente 11,5% em relação ao modelo ini-rand. Seguido por ini-h-pb e ini-heu, com reduções de 8,7% e 5,2%, respectivamente. Em relação à métrica E, o modelo ini-heu apresentou uma redução de 1,5% em relação ao modelo ini-rand. Analisamos as relações de dominância entre as abordagens diferentes de inicialização utilizando a métrica CS. Dentre os algoritmos que apresentaram os melhores resultados das métricas destacam-se ini-h-pgb, ini-heu e ini-h-pg.

Tabela 29 – Métrica CS das estratégias de inicialização de população.

A\B	ini-heu	ini-h-pb	ini-h-pg	ini-h-pgb	ini-rand
ini-heu	-	0,409	0,412	0,351	0,413
ini-h-pb	0,327	-	0,364	0,315	0,367
ini-h-pg	0,334	0,372	-	0,318	0,378
ini-h-pgb	0,381	0,407	0,417	-	0,415
ini-rand	0,346	0,385	0,391	0,336	-

A Tabela 29 apresenta os valores da métrica CS, sendo que ini-h-pgb supera levemente ini-heu, sendo capaz de dominar um maior número de soluções. Ao comparar ini-h-pgb com ini-rand, observamos que um número maior de soluções de ini-rand são dominados por ini-h-pgb. De forma similar, observamos que um número maior de soluções de ini-rand são dominados por ini-heu e 34,6% de ini-heu são dominados por ini-rand. Portanto em comparação à ini-rand ambos os modelos apresentam dominância semelhante, contudo ini-h-pgb apresenta resultados levemente superiores que ini-heu. Comparando ini-h-pg com ini-rand, observamos que embora um número maior de soluções de ini-rand sejam dominados por ini-h-pg, os valores são mais próximos que os analisados anteriormente (ini-h-pgb e ini-h-p). Ao comparar diretamente ini-h-pg e ini-heu, o último apresenta uma maior dominância.

A Tabela 30 apresenta os resultados do teste de hipótese aplicado com um nível de confiança de 95% para avaliar a significância do modelo ini-heu em comparação aos outros modelos. Comparando em relação à ini-h-pgb têm-se que não há diferença estatisticamente significativa, contudo há uma melhoria em relação à ini-rand em relação ao NS e em relação ao ini-h-pb para o hv. Finalmente, observamos que ini-h-pg e ini-h-pgb não apresentam diferença estatística significativa em relação a ini-heu. Em resumo, observa-se que

Tabela 30 – Teste de hipóteses para as métricas comparando ini-h-p em relação às outras estratégias de inicialização.

	ini-h-pg	ini-h-pb	ini-h-pgb	ini-rand
hv	=	>	=	=
NSV	=	=	=	>
IGD+(P,P*)	=	=	=	=
E	=	=	=	=

as estratégias ini-heu, ini-h-pgb e ini-h-pg superam ini-rand em dominância, contudo, não apresentam diferença estatística significativa entre si. Avaliando a dominância entre os modelos, conclui-se que ini-h-pgb demonstra um desempenho levemente superior em relação a ini-heu. No entanto, ini-h-pgb utiliza informações baseadas nos dados da população válida encontrada experimentalmente, diferentemente de ini-heu, que utiliza informações de dados de entrada do processo, sendo, portanto, mais generalizável a outros problemas. Portanto, ini-heu foi selecionada e será utilizada nas seções subsequentes. Em comparação com o modelo de referência, observa-se uma melhoria de 73,4% em relação ao IGD+, 34,9% em relação ao NS e uma diferença de 13,4% em relação ao E. Comparando com ini-rand, verifica-se uma melhoria de 8,0% em relação ao NS, 5,2% em relação ao IGD+ e uma diferença de 1,5% em relação ao E.

5.3.2 Estratégias de Mutação

Nesta seção será investigada a incorporação do operador de mutação baseado em heurísticas ao modelo ini-rand, será avaliado a sensibilidade das estratégias de mutação heurísticas criação de heurísticas para a mutação. Avaliou-se três probabilidades de utilização de heurísticas $pHeu = 25\%$ (mut-h-pH25), $pHeu = 50\%$ (mut-h-pH50) e $pHeu = 75\%$ (mut-h-pH75).

Tabela 31 – Desempenho do algoritmo de acordo com a estratégia de mutação utilizada.

		mut-h-pH25	mut-h-pH50	mut-h-pH75	ini-rand
hv	média	0,9±0,001	0,9±0,001	0,9±0,0	0,900±0,001
	min	0,898	0,898	0,900	0,900
	max	0,900	0,902	0,902	0,903
NSV	média	23,933±3,331	24,767±3,39	24,200±2,483	24,812±2,657
	min	17	17	18	19
	max	32	34	29	30
IGD+(P,P*)	média	3,19±1,183	3,452±1,966	3,54±1,96	2,134±1,187
	min	1,699	1,729	1,411	1,044
	max	5,927	10,545	10,820	5,581
E	média	0,895±0,053	0,912±0,049	0,902±0,051	0,867±0,068
	min	0,781	0,812	0,750	0,719
	max	1,000	1,000	1,000	0,969

A Tabela 31 apresenta os resultados de desempenho. Em relação à métrica hv, observamos que todos os testes apresentam médias próximas de 0,9. Portanto, não há uma melhoria expressiva em relação ao modelo de referência. Para a métrica NS, ini-rand apresenta o melhor resultado com as estratégias mut-h-pH50 e mut-h-pH75 apresentando médias ligeiramente inferiores de cerca de 0,2% e 2,4%, respectivamente em comparação com ini-rand. Em relação à métrica IGD+, observamos que ini-rand apresenta o melhor valor médio no valor de 2,134; seguido pela estratégia mut-h-pH25 que obteve o valor médio de 3,190. Na métrica E, ini-rand também demonstra o melhor resultado. Em seguida têm-se que mut-h-pH75 demonstra médias ligeiramente inferiores em comparação com ini-rand, com uma diferença de 2,8%. Nesta seção analisaremos a análise de dominância das fronteiras das estratégias de mutação. A Tabela 32 apresenta os resultados da métrica CS entre as diferentes estratégias.

Tabela 32 – Métrica CS das estratégias ini-rand mut-h-pH25 mut-h-pH50 mut-h-pH75

	ini-rand	mut-h-pH25	mut-h-pH50	mut-h-pH75
A\B				
ini-rand	-	0,518	0,553	0,526
mut-h-pH25	0,290	-	0,431	0,409
mut-h-pH50	0,263	0,381	-	0,384
mut-h-pH75	0,284	0,387	0,414	-

. Comparando ini-rand com as outras estratégias, observamos que ini-rand apresenta melhor relação de dominância em comparação às estratégias avaliadas. Dentre as estratégias de mutação avaliadas, destacam-se mut-h-pH25 e mut-h-pH50; comparando as duas estratégias, têm-se, que 38,1% das soluções de mut-h-pH25 são dominadas por mut-h-pH50. E que 43,1% das soluções de mut-h-pH50 são dominadas por mut-h-pH25.

Portanto dentre as estratégias de mutação avaliadas mut-h-pH25 apresenta maior dominância.

Os resultados dos testes de hipóteses pode ser visualizados em Tabela 33

Tabela 33 – Teste de hipóteses comparando mut-h-pH25 em relação às demais estratégias de mutação investigadas

	mut-h-pH50	mut-h-pH75	ini-rand
hv	=	=	<
NSV	=	=	=
IGD+(P,P*)	=	=	>
E	=	=	=

Analisando o teste de hipóteses têm se que mut-h-pH25 não apresenta diferença significativa em relação à mut-h-pH50 e mut-h-pH75. E apresenta piora em relação ao ini-rand, nas métricas hv e ao IGD. Ao avaliar o teste de hipóteses, observa-se que não há melhoria significativa entre mut-h-pH25, mut-h-pH50 e mut-h-pH75. No entanto, mut-h-pH25, quando comparada às outras operações de mutação heurística avaliadas. Portanto, será examinada a combinação de mut-h-pH25 com a estratégia de inicialização heurística.

5.3.3 Integração da Heurística de Inicialização da População e Mutação

Nesta seção, serão avaliados os resultados da combinação da heurística de inicialização de população selecionada (ini-heu) e da melhor configuração do operador de mutação baseado em heurísticas (mut-h-pH25). O algoritmo resultante dessa integração é referenciado como comb-h. A Tabela 34 mostra os resultados obtidos, em relação à métrica

Tabela 34 – Métricas avaliação para os testes de inicialização de população heurísticas

		ini-heu	mut-h-pH25	comb-h	ini-rand
hv	média	0,900 ±0,0	0,9±0,001	0,9±0,001	0,9±0,001
	min	0,900	0,898	0,898	0,900
	max	0,901	0,900	0,903	0,903
NSV	média	26,800 ±3,498	23,933±3,331	24,667±3,241	24,812±2,657
	min	20	17	21	19
	max	34	32	33	30
IGD+(P,P*)	média	2,021 ±0,825	3,19±1,183	2,868±1,237	2,134±1,187
	min	1,093	1,699	1,295	1,044
	max	4,933	5,927	6,603	5,581
E	média	0,852 ±0,065	0,895±0,053	0,905±0,045	0,867±0,068
	min	0,719	0,781	0,812	0,719
	max	0,969	1,000	1,000	0,969

hipervolume, observamos o desempenho semelhante nessa métrica. Para a métrica NS, o modelo ini-heu manteve o melhor resultado com uma melhora de 8,0% em relação a ini-rand. Analisando a métrica IGD+, novamente o ini-heu apresentou os melhores resultados com uma redução de 5,2% em relação a ini-rand. O modelo ini-heu também apresentou o melhor valor em relação à métrica E, com uma diferença de 1,5%. Essa análise comparativa confirma que a incorporação da mutação heurística não apresenta melhorias ao modelo ini-heu. Nesta seção serão avaliados os resultados da análise de dominância entre os testes. Nossa investigação partiu do AGMO publicado em (JANKAUSKAS; FARID, 2019), que foi proposto para uma instância do escalonamento de bateladas aplicados em dados reais de uma produção farmacêutica, sujeito a restrições de demanda. A partir da reprodução desse algoritmo, que denominamos modelo de referência ini-rand, conduzimos uma série de investigações para aumentar o desempenho dos AGMO nesse problema, que se mostrou complexo e com baixa convergência para soluções viáveis.

Tabela 35 – Métrica CS das estratégias de inicialização.

	comb-h	ini-heu	ini-rand	mut-h-pH25
A\B				
comb-h	-	0,239	0,275	0,395
ini-heu	0,549	-	0,413	0,557
ini-rand	0,520	0,346	-	0,518
mut-h-pH25	0,400	0,252	0,290	-

A Tabela 35 apresenta os resultados da métrica CS. Comparando os testes ini-heu e ini-rand, têm-se que 41,3% das soluções de ini-rand são dominadas por ini-heu, enquanto apenas 34,6% das soluções de ini-heu são dominadas por ini-rand. Portanto, ini-heu apresenta um desempenho superior em relação à dominância em comparação com ini-rand. Comparando comb-h e ini-rand, têm-se que 27,5% das soluções de ini-rand são dominadas por comb-h, enquanto apenas 52,0% das soluções de comb-h são fracamente dominadas por ini-rand. Portanto, ini-rand apresenta um desempenho superior em relação à dominância em comparação com comb-h. Essa análise comparativa revela as diferenças no desempenho dos modelos em relação à dominância. Os resultados demonstram que ini-heu é capaz de dominar uma maior porcentagem de soluções em comparação com comb-h e ini-rand.

Avaliando o teste de hipóteses apresentado em Tabela 36, comparando o teste ini-heu em relação aos testes mut-h-pH25 e comb-h têm-se que ini-heu apresenta melhorias significativas. Contudo ao comparar em relação à ini-rand há diferenças estatisticamente significativas apenas em relação à métrica NS. Isto confirma que a operação de mutação heurística não apresenta melhorias ao ser incorporada juntamente com a inicialização da população heurística.

Nesta seção avaliou-se os modelos com a combinação de inicialização da população e

Tabela 36 – Teste de hipóteses para as métricas comparando ini-heu em relação à outras estratégias de inicialização

	mut-h-pH25	comb-h	ini-rand
hv	>	>	=
NSV	>	>	>
IGD+(P,P*)	<	<	=
E	<	<	=

estratégia de mutação heurística. Avaliando os resultados dos testes de hipóteses, têm-se que o melhor modelo permanece sendo o modelo heurístico com inicialização da população ini-heu, que será denominado ini-heu.

5.3.4 Análise de Gerações

Nesta seção avaliou-se a evolução das métricas ao longo das gerações comparando ini-heu com o ini-rand. As Figuras 19 e 20 apresentam a evolução das métricas hv, E, IGD+ e NS por geração.

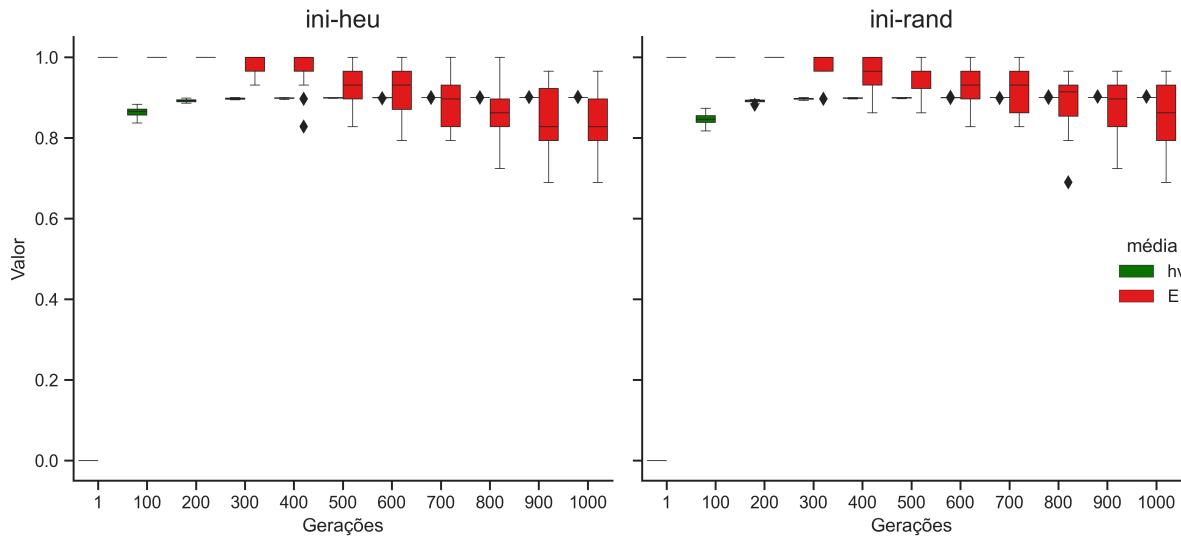


Figura 19 – Diagrama de caixas para as métricas hv e E.

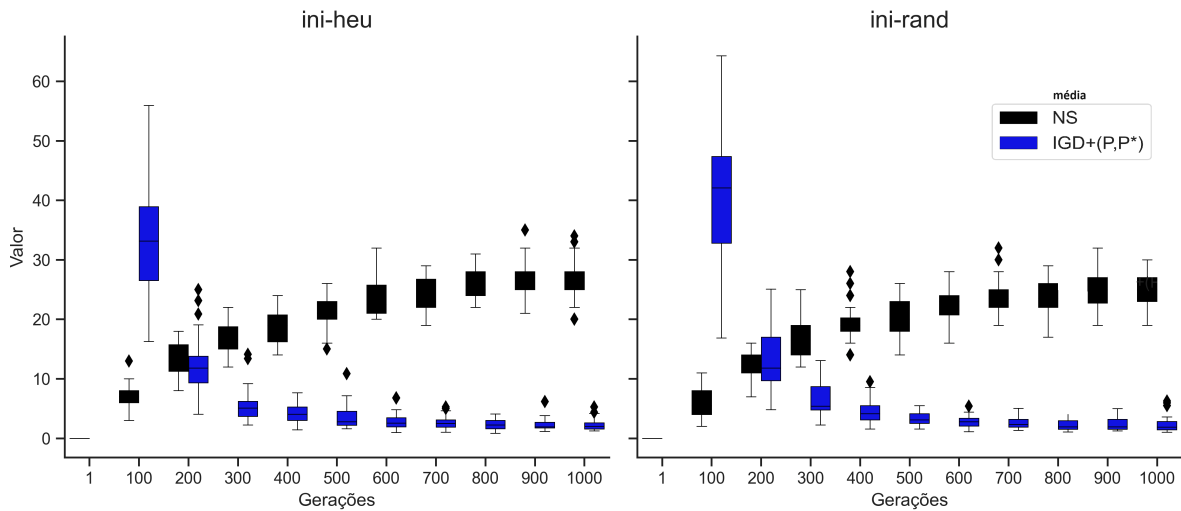


Figura 20 – Diagrama de caixas para as métricas IGD+ e NS.

Conforme Tabela 37, analisando a métrica hv, ini-heu apresenta uma melhora significativa na geração 100, com um aumento de 296,1% em relação à ref-s1. Adicionalmente, ini-rand encontra uma média superior à ref-s1 na geração 1000, a partir da geração 500.

Analisando o NS, observa-se uma melhora substancial de ini-heu em relação à ref-s1 de 1708,3%, aumentando a média de 0,400 para 7,233. Avaliando a métrica E, verifica-se que a partir da geração 400, ini-rand apresenta uma média menor de erro em comparação com ref-s1 para 1000 gerações. Avaliando o IGD+, observa-se que ini-heu, em relação à ref-s1, apresenta uma redução de 74,6% já na geração 100. Adicionalmente, a partir da geração 300, ini-heu é capaz de encontrar valores de IGD+ inferiores aos encontrados na geração 1000 de ref-s1. Portanto, ini-heu apresenta uma convergência mais rápida em

Tabela 37 – Métricas avaliação para os testes

testName	gen média	100	200	300	400	500	600	700	800	900	1000
ini-heu	E	1,000	1,000	0,986	0,964	0,925	0,908	0,886	0,862	0,846	0,837
	IGD+	33,721	12,524	5,644	4,185	3,654	2,861	2,631	2,321	2,288	2,287
	NS	7,233	13,333	16,567	18,800	21,333	23,267	23,567	25,467	26,600	26,800
	hv	0,865	0,892	0,897	0,898	0,899	0,900	0,900	0,900	0,900	0,900
ini-rand	E	1,000	1,000	0,985	0,960	0,947	0,921	0,923	0,895	0,878	0,853
	IGD+	40,259	13,230	6,573	4,574	3,340	2,764	2,567	2,311	2,296	2,357
	NS	5,938	12,281	16,781	19,438	20,344	22,344	24,094	23,844	24,938	24,812
	hv	0,847	0,891	0,897	0,899	0,899	0,900	0,900	0,900	0,900	0,900
ref-s1	E	1,000	1,000	1,000	1,000	1,000	0,999	0,994	0,992	0,991	0,985
	IGD+	133,034	70,976	42,410	28,075	19,026	14,689	11,381	9,311	8,847	8,557
	NS	0,400	10,167	13,200	12,600	12,900	14,467	14,933	17,433	18,700	19,933
	hv	0,218	0,824	0,864	0,880	0,889	0,893	0,895	0,896	0,897	0,898

relação ao número de gerações comparado ao ref-s1, indicando que o modelo, semelhante ao ini-rand, permite reduzir o número de gerações para 500 sem comprometer a qualidade dos resultados em relação ao ref-s1 e diminuindo o tempo de execução.

Ao comparar ini-heu em relação à ini-rand na Tabela 37, observa-se que o comportamento do HV ao longo das gerações é semelhante. Avaliando o NS, ini-heu apresenta uma melhora de 21,1% em relação à ini-rand na geração 100. Avaliando o IGD+, o ini-heu apresenta resultados 33,7% melhores em relação ao ini-rand na geração 100. Adicionalmente, ini-heu encontra um IGD+ na geração 800 inferior ao IGD+ encontrado pelo ini-rand na geração 1000. Analisando o E, há uma redução de 2,2% do E do ini-heu em relação ao ini-rand a partir da geração 500, sendo que ini-heu encontra uma taxa de Erro (E) na geração 900 inferior ao ini-rand mesmo para uma geração 1000. Portanto, há indícios de que ini-heu apresenta uma convergência mais rápida em relação à ini-rand.

5.3.5 Análise Gráfica ini-rand e ini-heu

Nesta seção será avaliada a análise gráfica para as melhores fronteiras para cada uma das métricas (hv, NS, IGD+ e E) para o modelo de referencia (**ref-s1**), melhor modelo randômico encontrado (**ini-rand**) e o melhor modelo heurístico (**ini-heu**).

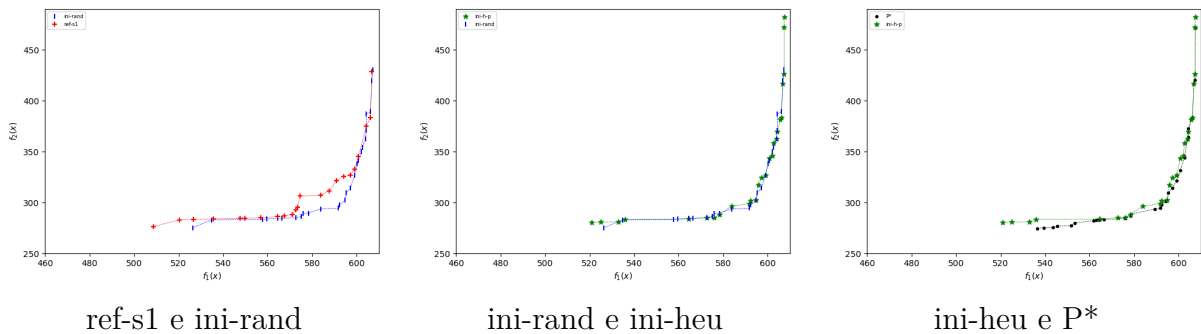


Figura 21 – Melhores fronteiras de não dominados em relação ao hv.

A Figura 21 apresenta as comparações entre diferentes modelos: ref-s1 vs. ini-rand,

ini-rand vs. ini-heu e ini-heu vs. P*. Para a métrica hv, ini-rand tem melhor desempenho em relação a ref-s1, particularmente para $570 < f_1(x) < 600$. O modelo ini-heu pode encontrar soluções para valores elevados de $f_1(x)$ em comparação com ini-rand. No entanto, ini-heu não consegue encontrar soluções para valores mais baixos de $f_1(x)$ tão próximas de P*.

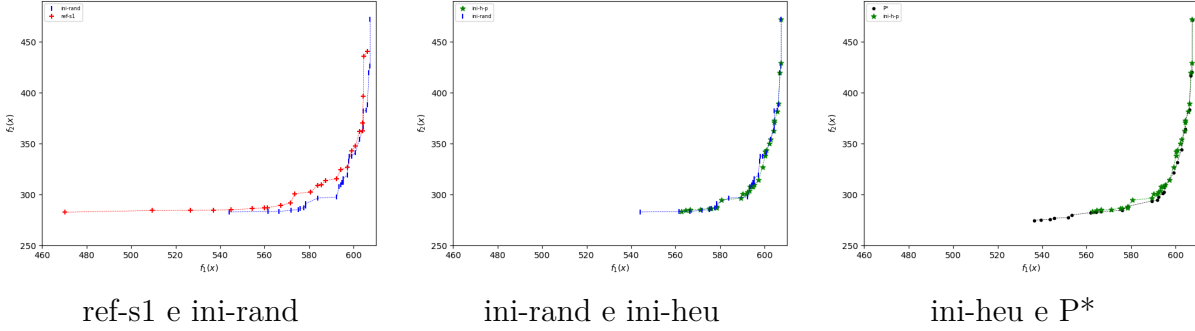


Figura 22 – Melhores fronteiras de não dominados em relação ao NS.

Analisando a métrica NS, na Figura 22, áreas de alta concentração de soluções, como próximas a $f_1 = 595$ para ini-rand, indicam o potencial de técnicas para melhorar a diversidade de soluções. Ao comparar ini-heu com P*, ini-heu não consegue encontrar soluções próximas a P* para $f_1 \geq 560$. Isso sugere que a heurística para otimização do custo de configuração na população inicial pode guiar a busca para valores de produção mais altos.

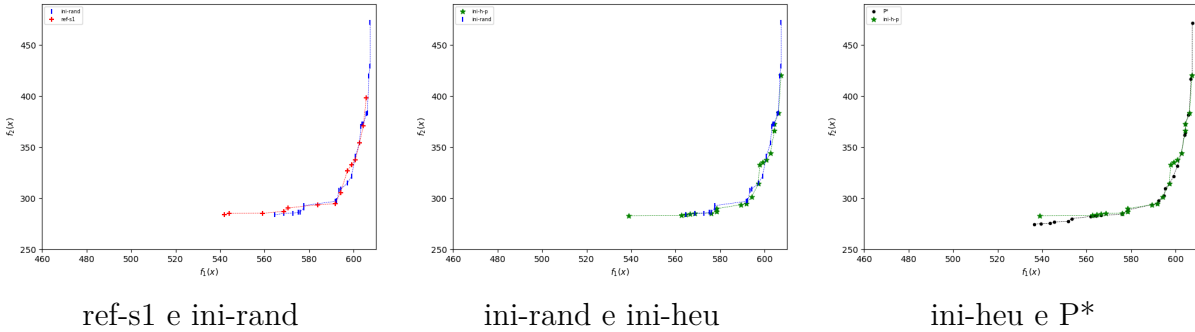


Figura 23 – Melhores fronteiras de não dominados em relação ao IGD+.

Para a métrica IGD+, na Figura 23, a fronteira ini-rand apresenta melhor desempenho em relação a ref para a maioria das soluções, especialmente para valores mais altos de $f_1(x)$. Comparando ini-heu e ini-rand, ini-heu domina ini-rand para a maioria dos valores em $580 \leq f_1(x) \leq 610$. Ao comparar ini-heu e P*, ini-heu demonstra comportamento semelhante a P* para a maioria dos intervalos, exceto para valores mais baixos de $f_1(x) < 540$ e para $f_1(x) = 600$. Em relação à métrica E (Figura 24), tanto ini-rand quanto ini-heu superam ref para a maioria dos valores. O ini-heu consegue encontrar soluções para valores altos de $f_2(x)$, porém apresenta dificuldades em encontrar soluções para $f_1(x) < 560$, de forma semelhante à métrica NS (Figura 24).

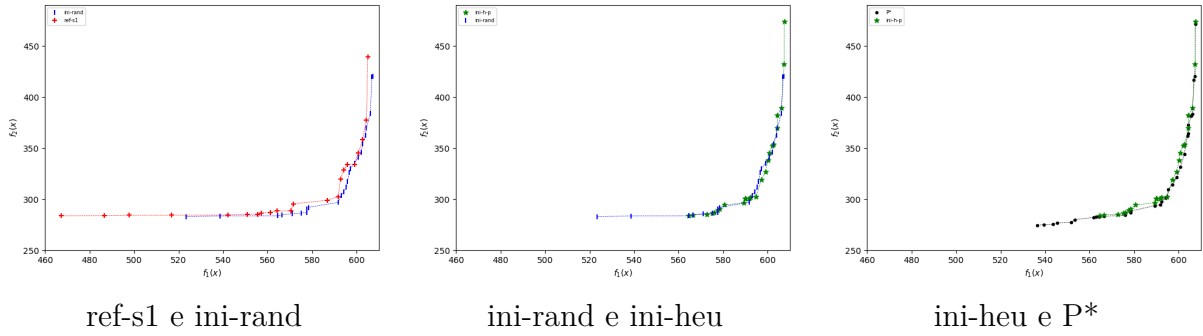


Figura 24 – Melhores fronteiras de não dominados em relação ao E.

Através da análise gráfica, ambos os modelos propostos apresentam consistentemente um desempenho superior a ref para a maioria das soluções. No entanto, o aumento da métrica NS para ini-heu pode estar associado a áreas de alta concentração de soluções e menor exploração dos extremos da fronteira. Técnicas para melhorar a diversidade da população, assim como uma heurística para otimizar $f_2(x)$, poderiam ser investigadas para lidar com essas questões.

5.4 Modelo baseado em Diversidade

Nesta seção, o objetivo é preservar ou melhorar as métricas multiobjetivo do modelo desenvolvido, denominado ini-heu. Apesar do alto número de gerações, emprega-se uma taxa de cruzamento reduzida. Assim, analisou-se a sensibilidade das taxas de cruzamento ao longo das gerações, selecionando uma taxa de cruzamento e um número máximo de gerações adequados que resultem em populações finais sem perdas significativas nas métricas avaliadas. A métrica E para o modelo ini-heu continua alta. Portanto, investigamos o operador de seleção particionado baseado em restrições, com a finalidade de aumentar a diversidade da população e aprimorar a exploração do espaço de busca. Com a intenção de utilizar a informação genética existente e permitir uma exploração mais eficaz do problema, investigou-se a incorporação de uma busca local pelo número de bateladas ao operador de mutação.

5.4.1 Sensibilidade das Taxas de Cruzamento ao longo da Evolução

O número de gerações para o modelo ini-heu é elevado (1000 gerações), portanto, avaliou-se a sensibilidade do aumento da probabilidade de cruzamento. Ao longo das gerações, selecionando para cada taxa de cruzamento um número máximo de gerações adequado. Em contrapartida, avaliou-se uma diminuição do número de gerações, sem um impacto significativo na qualidade das métricas da fronteira de Pareto.

Análise de diferentes taxas de cruzamento

Os resultados da sensibilidade do modelo ini-heu em relação às taxas de cruzamento de 30% (ini-heu), 60%(pC60), 75%(pC75) e 90%(pC90), serão comparados também com o modelo de referência (ref-s1). Os diagramas de caixas ilustram a evolução das métricas Figura 28.

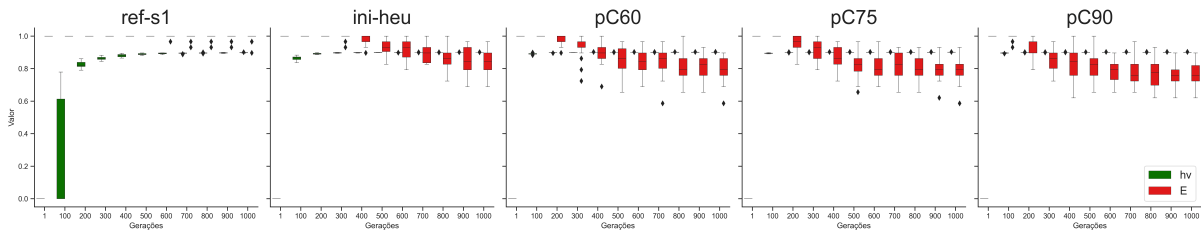


Figura 25 – Métricas hv e E

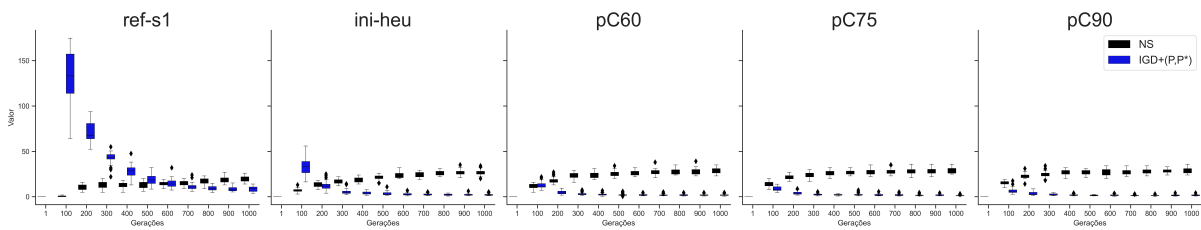


Figura 26 – Métricas IGD+ e NS

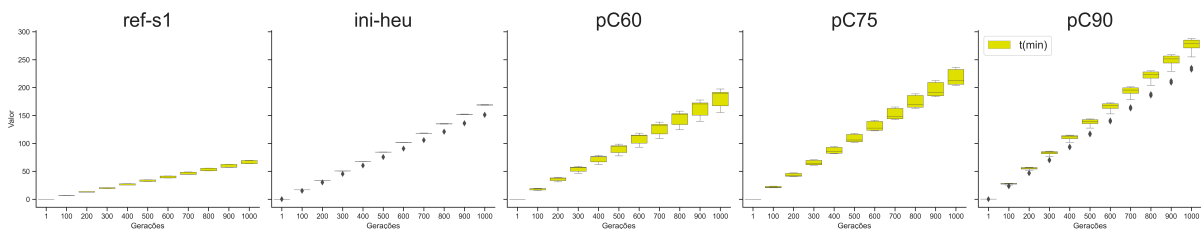


Figura 27 – Métrica Tempo médio (min)

Figura 28 – Evolução das métricas hv, E, IGD+, NS e tempo ao longo das gerações para comparando o teste de referência e o modelo ini-heu com diferentes taxas de cruzamento.

Avaliando a Figura 25, em relação à média do HV, o modelo ini-heu, em comparação à ref-s1, apresenta uma distribuição de valores médios maior mesmo em gerações baixas. Na geração 100, ini-heu (0,865) apresenta uma média 296,1% superior à ref-s1 (0,218). Um ponto interessante é que a mediana de ref-s1 para a geração 100 demonstrou uma grande diferença (valor 0) em relação à média (0,218), evidenciando a necessidade de inclusão da avaliação da mediana. Os resultados numéricos são apresentados na Tabela 38.

Tabela 38 – Métricas avaliação para os testes

metrica	gen testName	100	200	300	400	500	600	700	800	900	1000
E mediana	ini-heu	1,000	1,000	1,000	0,967	0,933	0,933	0,900	0,867	0,850	0,850
	pC60	1,000	1,000	0,933	0,900	0,867	0,850	0,867	0,800	0,833	0,800
	pC80	1,000	0,967	0,900	0,833	0,833	0,833	0,833	0,833	0,800	0,800
	pC90	1,000	0,967	0,867	0,850	0,833	0,800	0,767	0,783	0,767	0,767
	ref-s1	1,000	1,000	1,000	1,000	1,000	1,000	1,000	1,000	1,000	1,000
E medio	ini-heu	1,000	1,000	0,990	0,970	0,934	0,914	0,896	0,872	0,854	0,843
	pC60	1,000	0,979	0,935	0,899	0,854	0,839	0,841	0,817	0,811	0,801
	pC80	0,995	0,947	0,888	0,828	0,820	0,812	0,823	0,814	0,815	0,788
	pC90	0,994	0,939	0,861	0,831	0,823	0,791	0,787	0,782	0,766	0,777
	ref-s1	1,000	1,000	1,000	1,000	1,000	0,999	0,994	0,992	0,992	0,989
IGD+ mediana	ini-heu	33,148	11,702	5,031	4,033	2,725	2,539	2,530	2,160	1,916	1,975
	pC60	11,816	4,170	2,768	2,031	1,768	1,792	1,595	1,696	1,656	1,617
	pC80	8,017	3,192	2,558	1,916	1,651	1,658	1,702	1,738	1,554	1,551
	pC90	6,036	2,649	2,240	1,713	1,551	1,408	1,357	1,157	1,332	1,401
	ref-s1	133,202	67,433	43,399	28,409	18,453	13,871	10,808	8,727	8,374	8,772
IGD+ medio	ini-heu	33,687	12,494	5,602	4,155	3,632	2,848	2,613	2,293	2,261	2,250
	pC60	12,644	4,603	3,156	2,379	2,101	2,093	1,968	1,875	1,850	1,885
	pC80	8,284	3,679	3,146	2,049	1,891	1,772	1,944	1,810	1,737	1,808
	pC90	6,709	3,550	2,517	1,907	1,637	1,575	1,466	1,317	1,532	1,610
	ref-s1	133,034	70,695	42,248	28,008	18,943	14,618	11,313	9,243	8,787	8,470
NS mediana	ini-heu	7,00	13,00	16,50	19,00	21,00	22,50	23,00	25,00	27,00	27,00
	pC60	12,00	17,50	23,50	24,50	26,00	26,00	27,50	27,00	28,00	28,00
	pC80	15,00	21,00	24,00	26,00	27,00	28,00	29,00	29,00	29,00	29,00
	pC90	15,00	22,50	25,00	27,00	26,50	27,50	28,00	28,00	28,00	28,00
	ref-s1	0,00	10,00	13,00	13,50	12,50	15,00	15,00	17,50	19,00	20,00
NS medio	ini-heu	7,23	13,33	16,57	18,80	21,33	23,27	23,57	25,47	26,60	26,80
	pC60	11,85	18,00	23,44	24,12	25,47	26,03	27,59	27,62	28,09	28,62
	pC80	14,68	21,42	24,23	26,16	27,84	28,35	28,77	29,39	29,16	28,77
	pC90	14,87	22,70	25,10	26,67	26,83	27,13	27,53	27,83	28,03	28,93
	ref-s1	0,40	10,17	13,20	12,60	12,90	14,47	14,93	17,43	18,70	19,93
hv mediana	ini-heu	0,8655	0,8923	0,8970	0,8981	0,8996	0,8999	0,9001	0,9002	0,9003	0,9003
	pC60	0,8918	0,8993	0,8998	0,9003	0,9004	0,9004	0,9005	0,9005	0,9005	0,9005
	pC80	0,8953	0,8997	0,9003	0,9004	0,9004	0,9004	0,9005	0,9005	0,9005	0,9005
	pC90	0,8966	0,8999	0,9003	0,9004	0,9004	0,9004	0,9004	0,9004	0,9005	0,9005
	ref-s1	0,0000	0,8216	0,8641	0,8832	0,8893	0,8937	0,8949	0,8956	0,8972	0,8977
hv medio	ini-heu	0,8646	0,8920	0,8971	0,8984	0,8994	0,8997	0,9000	0,9001	0,9003	0,9004
	pC60	0,8910	0,8990	0,8997	0,9003	0,9004	0,9005	0,9006	0,9009	0,9009	0,9009
	pC80	0,8954	0,8995	0,9002	0,9005	0,9007	0,9008	0,9008	0,9009	0,9009	0,9011
	pC90	0,8964	0,8998	0,9003	0,9004	0,9005	0,9005	0,9006	0,9006	0,9007	0,9008
	ref-s1	0,2182	0,8244	0,8637	0,8803	0,8891	0,8933	0,8948	0,8960	0,8970	0,8977

Comparando ini-heu em relação a pC60, pC75 e pC90 na geração 100, observa-se uma leve melhora de 3,0 - 3,6%; contudo, a partir da geração 200, os valores médios tornam-se semelhantes, com incrementos entre 100 gerações inferiores a 1,0%.

Analisando a média da métrica E, comparando ini-heu com ref-s1, observa-se que ini-heu apresenta melhorias a partir da geração 300, com uma diferença de 0,085 na geração 600. Comparando ini-heu na geração 1000 com os outros testes, nota-se que pC60 é capaz de encontrar médias semelhantes a partir da geração 600; pC75 a partir da geração 500 e pC90 a partir da geração 400. Isso indica que, em relação à métrica E, há potencial para reduzir o número de gerações ao elevar a taxa de cruzamento, otimizando assim a eficiência do algoritmo. Avaliando a Figura 26, em relação à média do NS, observa-se que ref-s1, em comparação com ini-heu, apresenta uma redução de 94,4% na geração 100. Analisando o NS médio para ini-heu na geração 1000 (26,80), observa-se que pC60 é capaz de obter resultados superiores a partir da geração 700; pC75 a partir de 600 gerações e

pC90 a partir da geração 500. Isso indica que, em relação ao NS, a taxa de cruzamento mais elevada também permite a redução do número de gerações.

Avaliando a Figura 26, em relação à média do NS, observa-se que ref-s1, em comparação com ini-heu, apresenta uma redução de 94,4% na geração 100. Analisando o NS médio para ini-heu na geração 1000 (26,80), observa-se que pC60 é capaz de obter resultados superiores a partir da geração 700; pC75 a partir de 600 gerações e pC90 a partir da geração 500. Isso indica que, em relação ao NS, a taxa de cruzamento mais elevada também permite a redução do número de gerações, mantendo a qualidade da solução encontrada.

Em relação à média do IGD+, ref-s1, em comparação com ini-heu, apresenta uma piora de 294,49% na geração 100. Ini-heu é capaz de encontrar médias melhores que ref-s1 na geração 1000 a partir da geração 100. Comparando o resultado de ini-heu na geração 1000 (2,250), observa-se que pC60 é capaz de encontrar resultados melhores a partir da geração 500; pC75 e pC90 a partir da geração 400. Isso demonstra que, em relação ao IGD+, também é possível realizar a redução do número de gerações ao aumentar a taxa de cruzamento. Avaliando a média do tempo (Figura 27), observa-se que ref-s1 apresenta um tempo total 59,98% inferior ao ini-heu. Contudo, ini-heu apresenta um tempo semelhante ao tempo total de ref-s1 caso seja executado até a geração 500; para pC60 e pC75 até a geração 400 e para pC90 até a geração 300.

Uma possível explicação para esses resultados é que taxas de cruzamento mais altas promovem uma maior diversidade genética na população ao longo das gerações. Isso pode facilitar a exploração do espaço de soluções, permitindo ao algoritmo escapar de mínimos locais e encontrar soluções mais próximas do ótimo global. A taxa de cruzamento de 90% (pC90) mostrou ser particularmente eficaz, atingindo resultados comparáveis ao modelo de referência em um menor número de gerações, especialmente para as métricas de HV, IGD+ e NS.

Outra hipótese é a relação entre a taxa de cruzamento e a convergência do algoritmo. Taxas de cruzamento mais altas podem acelerar a convergência inicial, resultando em melhorias rápidas nas primeiras gerações, como observado na métrica de IGD+ e no número de soluções não dominadas (NS). No entanto, é importante balancear a taxa de cruzamento para evitar uma convergência prematura, onde o algoritmo pode perder a diversidade necessária para explorar novas soluções.

Adicionalmente, a análise do tempo médio de execução sugere que, embora modelos com taxas de cruzamento mais altas exijam mais tempo computacional por geração, a redução no número total de gerações necessárias para atingir um desempenho satisfatório pode compensar esse aumento de tempo. Isso indica que, para aplicações práticas, uma taxa de cruzamento elevada pode ser benéfica. Os resultados obtidos confirmam que o aumento das taxas de cruzamento tende a melhorar a qualidade das soluções do algoritmo, conforme evidenciado pelas métricas de desempenho analisadas. Destaca-se a taxa de

cruzamento de 90% (pC90) utilizando 600 gerações. Em comparação com ini-heu ao final das 1000 gerações, pC90 apresenta uma redução de 5,8% na mediana E, redução de 28,7% na mediana do IGD+, aumento de 1,8% na mediana do NS e uma redução de 0,8% na mediana do tempo. Esta configuração foi utilizada nos testes realizados na sequência.

Análise de dominância A Tabela 39 apresenta os valores dos modelos na geração 1000. Os resultados demonstram que o melhor desempenho ocorre para o modelo pC90. Observa-se que 60,49% dos elementos de um conjunto de soluções não dominadas de ini-heu são dominados por pC90, enquanto apenas 17,36% dos elementos de um conjunto de soluções não dominadas de pC90 são dominados por ini-heu.

Tabela 39 – Métrica CS das estratégias

A\B	ini-heu	pC60	pC80	pC90	ref-s1
ini-heu	-	0,223	0,212	0,174	0,885
pC60	0,565	-	0,290	0,287	0,907
pC80	0,517	0,383	-	0,306	0,920
pC90	0,605	0,409	0,324	-	0,920
ref-s1	0,054	0,041	0,032	0,029	-

Teste de Hipóteses por taxa de cruzamento em relação à última geração

Foram realizados testes de hipóteses comparando o resultado do modelo ini-heu na geração 1000 com os outros modelos com diferentes probabilidades de cruzamento em várias gerações (100, 200, ..., 900 e 1000). O objetivo foi identificar a partir de qual geração não há diferença estatisticamente significativa na qualidade das soluções para cada probabilidade de cruzamento.

Tabela 40 – Teste de hipóteses para as métricas comparando ini-heu-g1000 em relação à pC90 em diferentes gerações

	pC90g100	pC90g200	pC90g300	pC90g400	pC90g500	pC90g600	pC90g700	pC90g800	pC90g900	pC90g1000
hv	>	>	=	=	<	<	<	<	<	<
NS	>	>	=	=	=	=	=	=	=	<
IGD+	<	<	=	=	>	>	>	>	>	>
E	<	<	=	=	=	>	>	>	>	>

Conforme a Tabela 40, na geração 300 (pC90g300) não há diferença estatisticamente significativa em comparação com o modelo ini-heu na geração 1000 (ini-heu-g1000). Isso indica o potencial de reduzir o número de gerações necessárias durante as execuções. Destaca-se a geração 600 (pC90-g600), que apresenta melhoria significativa em relação a todas as métricas, exceto o NS.

Tabela 41 – Teste de hipóteses para as métricas comparando ini-heu-g1000 em relação à pC80 para diferentes gerações

	pC80g100	pC80g200	pC80g300	pC80g400	pC80g500	pC80g600	pC80g700	pC80g800	pC80g900	pC80g1000
hv	>	>	>	=	<	<	<	<	<	<
NS	>	>	>	=	=	=	<	<	<	<
IGD+	<	<	<	=	=	=	=	=	>	>
E	<	<	<	=	=	=	=	=	=	>

A Tabela 41 apresenta a comparação com o modelo pC80. Observa-se que, a partir da geração 400, obtêm-se resultados semelhantes, mas o modelo apresenta melhoria significativa em relação ao E, IGD+ e HV apenas na geração 1000.

Tabela 42 – Teste de hipóteses para as métricas comparando ini-heu-g1000 em relação ao pC60 em diferentes gerações

	pC60g100	pC60g200	pC60g300	pC60g400	pC60g500	pC60g600	pC60g700	pC60g800	pC60g900	pC60g1000
hv	>	>	>	>	=	<	<	<	<	<
NS	>	>	>	>	=	=	=	=	=	<
IGD+	<	<	<	=	=	=	=	=	=	=
E	<	<	<	<	=	=	=	=	=	>

Avaliando a Tabela 42, observa-se que, a partir da geração 500, os resultados são semelhantes. No entanto, apenas na geração 1000 há resultados estatisticamente significativos em relação ao HV, NS e E.

Esses resultados indicam que taxas de cruzamento mais elevadas promovem uma maior diversidade genética, facilitando a exploração do espaço de soluções e acelerando a convergência inicial. A presença de outliers nas primeiras gerações pode influenciar as médias, mas as melhorias nas medianas ao longo das gerações corroboram a eficácia das taxas de cruzamento elevadas na otimização multiobjetivo. Portanto, selecionou-se a configuração do modelo pC90 para os testes subsequentes, reduzindo o número de gerações para 600. Com um tempo de execução semelhante ao modelo ini-heu na geração 1000, o pC90 apresentou resultados estatisticamente superiores em relação às métricas analisadas, com uma redução de 5,8% na mediana do E, uma redução de 28,7% no IGD+, um aumento de 0,7% no HV e um aumento de 1,8% no NS.

5.4.2 Seleção Particionada Baseada em Restrições

Nesta seção analisou-se a influencia da incorporação da seleção para cruzamento particionado e da reinserção particionada ao modelo desenvolvido em Subseção 5.4.1.

5.4.2.1 Seleção para Cruzamento Particionada

Incorporou-se ao modelo a seleção para o cruzamento avaliou-se as probabilidades de utilização do critério original (pMat), priorizando restrição, e em caso de empate o NSGA-II, de 50% (d-pMat50), 70%(d-pMat70), 80%(d-pMat80), 90%(d-pMat90) e 100%(pC90). A Tabela 43 apresenta os resultados. Em relação à mediana do E, o melhor modelo é

Tabela 43 – Métricas avaliação para os testes

metrica	gen testName	100	200	300	400	500	600
E mediana	d-pMat50	1,000	0,967	0,900	0,883	0,867	0,867
	d-pMat70	1,000	0,967	0,933	0,867	0,867	0,867
	d-pMat80	1,000	0,933	0,900	0,833	0,800	0,817
	d-pMat90	1,000	0,933	0,900	0,867	0,833	0,833
	pC90	1,000	0,967	0,867	0,850	0,833	0,800
E medio	d-pMat50	1,000	0,961	0,907	0,881	0,849	0,853
	d-pMat70	0,999	0,963	0,912	0,879	0,853	0,847
	d-pMat80	0,994	0,929	0,893	0,817	0,801	0,801
	d-pMat90	0,997	0,932	0,887	0,867	0,829	0,829
	pC90	0,994	0,939	0,861	0,831	0,823	0,791
IGD+ mediana	d-pMat50	11,320	3,687	2,544	1,767	1,852	1,799
	d-pMat70	9,137	2,813	2,440	1,777	1,696	1,618
	d-pMat80	7,694	2,646	1,950	1,979	1,646	1,530
	d-pMat90	8,255	2,889	2,051	1,899	1,639	1,683
	pC90	6,036	2,649	2,240	1,713	1,551	1,408
IGD+ medio	d-pMat50	13,268	4,360	2,684	1,959	2,085	1,971
	d-pMat70	9,128	3,645	2,459	1,978	1,773	1,591
	d-pMat80	8,084	3,417	2,267	2,155	1,985	1,646
	d-pMat90	8,744	3,253	2,204	2,063	1,896	1,904
	pC90	6,709	3,550	2,517	1,907	1,637	1,575
NS mediana	d-pMat50	11,50	19,50	24,00	26,00	27,00	28,00
	d-pMat70	14,00	20,00	24,50	26,00	27,00	28,00
	d-pMat80	16,00	21,50	23,00	26,00	26,00	29,00
	d-pMat90	15,00	22,00	26,00	28,00	27,00	28,50
	pC90	15,00	22,50	25,00	27,00	26,50	27,50
NS medio	d-pMat50	11,57	19,87	24,23	24,83	27,30	28,10
	d-pMat70	14,33	19,83	25,27	26,70	27,67	27,90
	d-pMat80	15,47	21,44	23,69	26,16	26,72	28,78
	d-pMat90	15,34	22,06	25,44	26,59	26,94	28,50
	pC90	14,87	22,70	25,10	26,67	26,83	27,13
hv mediana	d-pMat50	0,8925	0,8996	0,9002	0,9003	0,9004	0,9004
	d-pMat70	0,8950	0,8996	0,9001	0,9003	0,9004	0,9005
	d-pMat80	0,8957	0,8998	0,9002	0,9003	0,9004	0,9005
	d-pMat90	0,8967	0,8999	0,9003	0,9004	0,9004	0,9004
	pC90	0,8966	0,8999	0,9003	0,9004	0,9004	0,9004
hv medio	d-pMat50	0,8926	0,8995	0,9002	0,9005	0,9006	0,9008
	d-pMat70	0,8944	0,8995	0,9001	0,9003	0,9004	0,9006
	d-pMat80	0,8955	0,8996	0,9002	0,9004	0,9005	0,9006
	d-pMat90	0,8962	0,8998	0,9004	0,9006	0,9007	0,9007
	pC90	0,8964	0,8998	0,9003	0,9004	0,9005	0,9005
t(min) mediana	d-pMat50	17,852	35,703	53,555	71,407	89,258	107,110
	d-pMat70	17,902	35,803	53,705	71,607	89,508	107,410
	d-pMat80	17,765	35,530	53,295	71,060	88,825	106,590
	d-pMat90	17,731	35,462	53,193	70,923	88,654	106,385
	pC90	27,943	55,887	83,830	111,773	139,717	167,660
t(min) medio	d-pMat50	17,814	35,628	53,442	71,256	89,069	106,883
	d-pMat70	17,890	35,780	53,670	71,560	89,451	107,341
	d-pMat80	17,777	35,555	53,332	71,109	88,887	106,664
	d-pMat90	17,767	35,534	53,300	71,067	88,834	106,601
	pC90	27,215	54,429	81,644	108,858	136,073	163,287

o d-pMat80, contudo ainda apresenta um aumento de 2,0% em comparação com a probabilidade de 100%(pC90). Avaliando a mediana do IGD+, comportamento semelhante ao E ocorre. Tal que o melhor modelo é o d-pMat80, mas apresenta um resultado pior

em relação ao pC90 (8,6% superior). Para a mediana do hv, têm-se resultados semelhantes contudo o melhor permanece sendo o pC90. Para a métrica NS, o teste d-pMat80 apresenta um aumento de 5,4% na mediana em comparação com pC90 na geração 600.

A Tabela 44 sugere que a estratégia de seleção para cruzamento particionada baseada em restrição não apresenta melhoria significativa. Uma vez que pC90 apresenta melhor dominância em relação aos modelos avaliados d-pMat50, d-pMat70, d-pMat80 e d-pMat90. Dentre eles o d-pMat90 apresenta a melhor dominância, com 31,6% dos elementos de pC90 são dominadas por d-pMat90. Contudo 37,4% dos elementos de d-pMat90 são dominadas por pC90.

Tabela 44 – Métrica CS das estratégias

A\B	d-pMat50	d-pMat70	d-pMat80	d-pMat90	pC90
d-pMat50	-	0,347	0,323	0,335	0,299
d-pMat70	0,378	-	0,343	0,354	0,313
d-pMat80	0,379	0,363	-	0,352	0,304
d-pMat90	0,386	0,366	0,350	-	0,316
pC90	0,406	0,385	0,363	0,374	-

Na Tabela 45, observa-se que o teste d-pMat80 apresenta métricas semelhantes em relação ao hv, IGD+ e E. Contudo, o pC90 não apresenta superioridade estatística em relação ao teste d-pMat80 para o NS. No entanto, é importante ressaltar que a métrica NS, embora relevante, pode não refletir diretamente a qualidade das soluções, pois uma boa solução pode dominar muitas outras, reduzindo assim o valor de NS.

Tabela 45 – Teste de hipóteses para as métricas comparando pC90 em relação à outros

	d-pMat50	d-pMat70	d-pMat80	d-pMat90
hv	=	=	=	=
NS	=	=	<	=
IGD+(P,P*)	<	=	=	<
E	<	<	=	<

Avaliando os resultados, observa-se que a incorporação da reinserção particionada na seleção para o cruzamento não apresenta evidências significativas de melhorias, portanto, o pC90 permanece sendo a melhor configuração.

5.4.2.2 Reinserção Particionada

Incorporou-se ao modelo a reinserção particionada, utilizando-se percentuais de população com a ordenação de reinserção original (pRe) de 20% (d-pRe20), 40% (d-pRe40), 60% (d-pRe60), 70% (d-pRe70), 80% (d-pRe80), 90% (d-pRe90) e 100% (pC90-g600).

Para d-pRe60, por exemplo, em cada reinserção, primeiro seleciona-se 60% da população baseando-se no critério de ordenação pela restrição e, em caso de empate, utiliza-se os critérios do NSGA-II. Os 40% restantes são selecionados da população corrente remanescente com base nos critérios do NSGA-II. Analisando os valores da taxa de erro, observa-se

Tabela 46 – Desempenho das abordagens que empregam reinserção particionada.

metrica	gen teste	100	200	300	400	500	600
E mediana	d-pRe20	1,000	1,000	0,967	0,933	0,900	0,900
	d-pRe40	1,000	0,967	0,900	0,867	0,833	0,800
	d-pRe60	1,000	0,933	0,833	0,800	0,767	0,733
	d-pRe70	1,000	0,900	0,867	0,833	0,767	0,767
	d-pRe80	1,000	0,933	0,850	0,800	0,800	0,767
	d-pRe90	1,000	0,933	0,867	0,867	0,833	0,800
E medio	pC90	1,000	0,967	0,867	0,850	0,833	0,800
	d-pRe20	1,000	0,988	0,963	0,935	0,908	0,887
	d-pRe40	0,998	0,952	0,884	0,873	0,826	0,801
	d-pRe60	0,992	0,920	0,848	0,798	0,764	0,726
	d-pRe70	0,996	0,903	0,858	0,822	0,795	0,766
	d-pRe80	0,990	0,924	0,864	0,828	0,799	0,784
IGD+ mediana	d-pRe90	0,995	0,914	0,867	0,850	0,840	0,807
	pC90	0,994	0,939	0,861	0,831	0,823	0,791
	d-pRe20	12,267	4,059	2,401	1,742	1,245	1,161
	d-pRe40	8,048	2,857	1,618	1,633	1,234	1,274
	d-pRe60	6,634	2,252	1,683	1,699	1,244	1,299
	d-pRe70	5,529	1,899	1,777	1,574	1,410	1,291
IGD+ medio	d-pRe80	5,743	2,308	1,842	1,893	1,577	1,311
	d-pRe90	6,080	2,659	1,957	1,770	1,747	1,456
	pC90	6,036	2,649	2,240	1,713	1,551	1,408
	d-pRe20	12,833	4,617	2,767	1,812	1,488	1,476
	d-pRe40	8,113	3,302	2,001	1,949	1,342	1,453
	d-pRe60	6,842	2,631	2,031	1,849	1,409	1,469
NS mediana	d-pRe70	6,046	2,306	1,821	1,672	1,641	1,440
	d-pRe80	6,534	3,022	2,056	1,927	1,754	1,649
	d-pRe90	6,581	2,900	2,243	2,001	1,865	1,660
	pC90	6,709	3,550	2,517	1,907	1,637	1,575
	d-pRe20	10,00	17,00	17,00	20,00	20,00	21,00
	d-pRe40	15,00	21,00	22,00	24,00	25,00	24,00
NS medio	d-pRe60	16,00	23,00	25,00	26,00	26,00	26,00
	d-pRe70	15,50	22,50	25,00	26,00	27,00	28,00
	d-pRe80	16,00	22,00	25,00	26,00	26,50	27,00
	d-pRe90	17,00	23,00	25,00	27,50	28,00	29,00
	pC90	15,00	22,50	25,00	27,00	26,50	27,50
	d-pRe20	10,71	16,42	18,00	19,87	20,13	20,39
hv mediana	d-pRe40	14,81	21,38	22,75	23,97	24,53	24,62
	d-pRe60	16,34	22,88	25,03	25,78	26,12	26,50
	d-pRe70	15,94	22,31	25,00	26,72	27,25	27,56
	d-pRe80	16,44	22,25	25,22	26,44	27,03	27,25
	d-pRe90	16,88	22,72	25,22	27,25	28,19	28,41
	pC90	14,87	22,70	25,10	26,67	26,83	27,13
hv medio	d-pRe20	0,8885	0,8975	0,8983	0,8994	0,8992	0,8996
	d-pRe40	0,8958	0,8995	0,9003	0,9004	0,9004	0,9004
	d-pRe60	0,8970	0,9000	0,9003	0,9004	0,9004	0,9004
	d-pRe70	0,8977	0,9000	0,9003	0,9004	0,9005	0,9005
	d-pRe80	0,8974	0,9000	0,9003	0,9004	0,9004	0,9004
	d-pRe90	0,8972	0,8999	0,9003	0,9004	0,9004	0,9005
	pC90	0,8966	0,8999	0,9003	0,9004	0,9004	0,9004
	d-pRe20	0,8879	0,8969	0,8984	0,8991	0,8993	0,8995
	d-pRe40	0,8960	0,8994	0,9002	0,9004	0,9004	0,9005
	d-pRe60	0,8973	0,8999	0,9003	0,9004	0,9005	0,9006
	d-pRe70	0,8976	0,8999	0,9004	0,9005	0,9006	0,9007
	d-pRe80	0,8970	0,8999	0,9003	0,9004	0,9006	0,9006
	d-pRe90	0,8967	0,8998	0,9003	0,9005	0,9007	0,9008
	pC90	0,8964	0,8998	0,9003	0,9004	0,9005	0,9005

que os modelos que adotam a reinserção particionada com probabilidades de 60, 70 e 80%

alcançaram médias e medianas menores que do pC90. O algoritmo d-pRe60 retorna os melhores valores, promovendo uma redução de 8,4% na mediana e 8,2% na média, em relação ao algoritmo sem reinserção particionada.

Avaliando a mediana do IGD+, observa-se que a redução da probabilidade melhora os resultados do IGD+, com exceção do d-pRe90, todas as demais abordagens com reinserção particionada superam o pC90, com o d-pRe20 apresentando a menor mediana, atingindo um valor 17,5% menor que aquele obtido pelo pC90. Ao avaliar a média do IGD+, os resultados são semelhantes para as probabilidades abaixo de 70%, com d-pRe70 apresentando o menor valor médio, com uma redução de 8,6% em relação ao pC90. Observando os dados da métrica NS, nota-se a redução no número de soluções válidas encontradas a medida que a probabilidade usada na reinserção particionada também diminui. Tal comportamento pode ser justificado pela própria essência da seleção particionada, que desconsidera a restrição do problema durante a escolha de uma parte da população sobrevivente. Apesar disso, as abordagens com probabilidades de 70% e 90% foram capazes de superar pC90 (pRe = 100%) tanto na mediana, quanto no valor médio. O modelo d-pRe90 apresenta os melhores valores para essa métrica, proporcionando um aumento de 5,5% na mediana e 4,7% no valor médio. Em relação ao hipervolume, todas as abordagens avaliadas apresentaram valores médios e medianos similares e, com exceção do d-pRe20, os modelos que usam reinserção particionada retornam valores iguais ou ligeiramente superiores ao pC90.

Os resultados sugerem que a estratégia de reinserção particionada apresenta potencial em especial devido supera a reinserção uniforme usada nas demais abordagens. Diferentemente da seleção para o cruzamento, esses resultados indicam que a maior diversidade gerada pela reinserção particionada aprimora a eficiência da exploração do AG, favorecendo, em especial, o direcionamento para indivíduos pertencentes ou próximos a fronteira de Pareto, mas sem comprometer significativamente a obtenção de soluções que atendam a restrição. Isso é obtido devido a um maior balanceamento na exploração entre as fronteiras de não dominância e os indivíduos válidos ao longo das gerações.

No geral, os algoritmos com às probabilidades 40 e 60% apresentam bons resultados, sugerindo que o balanceamento entre soluções que atendam ou se aproximem da restrição e indivíduos melhor posicionados nas fronteiras de não dominância tem um efeito positivo em relação às métricas observadas. Em outras palavras, há indícios que o AGMO se beneficia com uma maior diversidade na população. A Tabela 47 apresenta os resultados de dominância entre os modelos. Os resultados indicam que o operador apresenta melhor dominância em relação à pC90 para a maior parte dos modelos: d-pRe20, d-pRe40, d-pRe60 e d-pRe80. Destacam-se os modelos d-pRe40 e d-pRe60; d-pRe40 45,5% domina de pC90, contudo apenas 22,2% de d-pRe40 são dominadas por pC90. O modelo d-pRe60 domina 41,3% de pC90, contudo pC90 domina apenas 25,5%. Apesar de pC90 apresentar melhor dominância em relação à d-pRe90, a diferença é muito pequena; o modelo pC90

domina 33,5% de d-pRe90 e d-pRe90 domina um valor levemente superior de 34,0%.

Tabela 47 – Métrica CS das estratégias

A\B	d-pRe20	d-pRe40	d-pRe60	d-pRe70	d-pRe80	d-pRe90	pC90
d-pRe20	-	0,236	0,310	0,344	0,361	0,389	0,385
d-pRe40	0,375	-	0,358	0,398	0,433	0,458	0,455
d-pRe60	0,348	0,275	-	0,352	0,392	0,418	0,413
d-pRe70	0,327	0,257	0,300	-	0,373	0,393	0,389
d-pRe80	0,316	0,242	0,280	0,305	-	0,367	0,360
d-pRe90	0,300	0,227	0,264	0,280	0,325	-	0,335
pC90	0,292	0,222	0,255	0,278	0,322	0,340	-

Análise do teste de hipóteses

Conforme a Tabela 48, as probabilidades de 70,80 e 90%, não apresentaram diferença estatística significativa em comparação com a execução pC90-g600. Contudo, o modelo d-pRe60 apresenta uma diferença estatística em relação à pC90-g600 será selecionada como a melhor configuração.

Tabela 48 – Teste de hipóteses para as métricas comparando pC90-g600 em relação à outros

	d-pRe20	d-pRe40	d-pRe60	d-pRe70	d-pRe80	d-pRe90
hv	>	=	=	=	=	=
NS	>	>	=	=	=	=
IGD+(P,P*)	=	=	=	=	=	=
E	<	=	>	=	=	=

A incorporação da reinserção particionada baseada em restrições demonstrou através do teste de hipóteses que ocorreu uma melhoria significativa para o E. O melhor modelo é o d-pRe60 que apresentou uma redução da mediana do E de 8,3% e redução de 7,7% da mediana do IGD+. Através da análise de dominância, verificou-se que a maior parte dos modelos avaliados apresentou uma melhoria em relação à dominância comparando-se com o modelo pC90. Adicionalmente confirmou-se que o modelo d-pRe60 também apresenta uma melhor dominância em relação à pC90. Portanto a melhor configuração ocorre para o modelo d-pRe60.

5.4.3 Sensibilidade em Relação à Mutação de Gene

Avaliou-se no operador de mutação a sensibilidade em relação à probabilidade de adição de gene de modo a permitir a redução do número de genes. Inicialmente, analisou-se a sensibilidade da probabilidade de mutação de gene (pMG) utilizando uma probabilidade

de adição de gene fixa de 50% (pAG), selecionando pMG com performance semelhante ou superior à d-pRe60. Em seguida, utilizando a pMG selecionada, realizou-se a análise de sensibilidade em relação à pAG. Os resultados detalhados são apresentados em Apêndice A.1. A partir da análise de sensibilidade, selecionou-se o modelo pMG40, que utiliza uma probabilidade de pMG de 40% e pAG de 50%. Adicionalmente, o modelo pMG40pAG0, com pMG de 40% e pAG de 0%, não encontrou soluções válidas, indicando a importância da capacidade de adicionar genes aos indivíduos para o problema. Uma hipótese é que o operador de mutação, especialmente a operação de adição de novo gene, assemelha-se a um operador de busca local.

5.4.4 Operador de Cruzamento Produto-Batelada

Nesta seção, investigou-se a incorporação do cruzamento uniforme com probabilidades de troca de produto (pCP) e de troca de batelada (pCB) (conforme Seção 4.3). Ao invés de realizar a troca conjunta do par produto e número de bateladas, pode-se trocar o produto e o número de bateladas independentemente. Avaliaram-se as combinações de pCP = 0% e pCB = 100% (d-pCP0pCB100); pCP = 100% e pCB = 0% (d-pCP100pCB0); pCP = 100% e pCB = 100% (d-pCP100pCB100); pCP = 100% e pCB = 50% (d-pCP100pCB50); pCP = 50% e pCB = 100% (d-pCP50pCB100); e pCP = 75% e pCB = 100% (d-pCP75pCB100). Os resultados detalhados são apresentados em Apêndice A.2. Os experimentos não apresentaram evidências de que o cruzamento proposto, permitindo a troca de produtos e bateladas de forma independente, melhora a qualidade das soluções em comparação com o cruzamento simultâneo do par produto e batelada. Contudo, há indícios de maior influência do número de bateladas em relação ao produto neste cruzamento, o que será investigado na próxima seção.

5.4.5 Operador de Mutação com Busca Local de Batelada

Nesta seção, avaliou-se a utilização do operador de mutação com a incorporação da busca local de batelada (conforme Subseção 4.4.2), analisando a sensibilidade em relação ao valor do delta de número de bateladas através de diferentes % da média do número de bateladas pDM , utilizando um número máximo de iterações (MI) de 2. Avaliaram-se os valores de pDM de 25% (d-BLB25), 50% (d-BLB50), 65% (d-BLB65), 75% (d-BLB75), 85% (d-BLB85) e 100% (d-BLB100). Compararam-se os resultados com o modelo sem a introdução da busca local (d-pMG40). A Tabela 49 apresenta os resultados. Em relação à mediana do E, observa-se que pDM mais próximos de 100% demonstram melhores resultados. O d-BLB75 apresenta uma redução de 8,7% na mediana e 4,5% no valor médio. Em relação à métrica IGD+, a melhor mediana é obtida com a busca local e uma probabilidade de 85%, reduzindo em 12,2% o valor de d-pMG40. O d-BLB75 retorna o melhor valor médio entre as abordagens investigadas, sendo 10% menor que aquele obtido

pelo algoritmo sem busca local. Para ambas as métricas, com exceção da abordagem que adota probabilidade de 25%, todas as demais apresentam resultados melhores que aqueles alcançados por d-pMG40. Isso indica que a busca local de bateladas auxilia na aproximação da população do AG à fronteira de Pareto.

Tabela 49 – Métricas avaliação para os testes

metrica	gen testName	100	200	300	400	500	600
E mediana	d-BLB100	1,000	0,900	0,783	0,750	0,750	0,733
	d-BLB25	1,000	0,933	0,800	0,767	0,800	0,767
	d-BLB50	1,000	0,867	0,800	0,767	0,767	0,767
	d-BLB65	1,000	0,900	0,767	0,767	0,733	0,733
	d-BLB75	1,000	0,867	0,800	0,800	0,733	0,700
	d-BLB85	1,000	0,867	0,783	0,733	0,733	0,717
	d-pMG40	1,000	0,967	0,900	0,833	0,767	0,767
E medio	d-BLB100	0,994	0,887	0,784	0,753	0,747	0,730
	d-BLB25	0,999	0,932	0,803	0,766	0,772	0,770
	d-BLB50	0,994	0,868	0,782	0,757	0,753	0,754
	d-BLB65	0,996	0,874	0,779	0,763	0,750	0,743
	d-BLB75	0,999	0,865	0,794	0,769	0,723	0,722
	d-BLB85	0,999	0,874	0,782	0,749	0,745	0,726
	d-pMG40	1,000	0,953	0,892	0,832	0,774	0,756
IGD+ mediana	d-BLB100	9,423	2,633	1,939	1,486	1,256	1,127
	d-BLB25	18,202	4,112	1,911	1,586	1,556	1,643
	d-BLB50	9,480	2,274	1,715	1,348	1,330	1,167
	d-BLB65	11,077	2,227	1,692	1,491	1,345	1,106
	d-BLB75	10,799	2,214	1,458	1,486	1,209	1,223
	d-BLB85	9,730	2,250	1,534	1,400	1,284	1,096
	d-pMG40	10,656	3,409	2,277	1,809	1,581	1,249
IGD+ medio	d-BLB100	9,412	2,940	2,238	1,685	1,505	1,360
	d-BLB25	19,691	4,073	2,764	2,155	1,958	2,075
	d-BLB50	9,747	3,101	2,072	1,697	1,593	1,399
	d-BLB65	11,740	2,735	2,089	1,608	1,548	1,446
	d-BLB75	10,380	2,756	1,910	1,845	1,242	1,331
	d-BLB85	9,865	2,524	1,649	1,600	1,544	1,399
	d-pMG40	11,026	4,232	2,710	1,951	1,630	1,479
NS mediana	d-BLB100	14,00	22,00	24,00	23,50	24,00	24,00
	d-BLB25	11,00	18,00	22,00	23,00	24,00	24,00
	d-BLB50	13,00	20,00	22,00	24,00	25,00	25,00
	d-BLB65	14,00	20,00	23,50	24,00	25,00	24,00
	d-BLB75	15,00	20,00	22,00	24,00	24,00	25,00
	d-BLB85	14,00	19,00	23,00	24,00	25,00	24,00
	d-pMG40	14,00	20,50	22,50	25,50	26,00	25,00
NS medio	d-BLB100	14,87	21,47	23,60	24,70	24,70	24,67
	d-BLB25	11,41	19,24	21,52	22,97	23,76	24,17
	d-BLB50	13,52	20,79	23,48	24,21	25,03	25,31
	d-BLB65	14,17	20,23	23,67	24,83	25,57	24,97
	d-BLB75	15,78	19,33	22,30	24,41	24,63	25,19
	d-BLB85	14,68	20,00	23,25	24,43	24,64	25,07
	d-pMG40	14,40	20,57	22,63	25,37	25,23	25,97
hv mediana	d-BLB100	0,8957	0,8993	0,8996	0,8996	0,8996	0,8996
	d-BLB25	0,8859	0,8990	0,8993	0,8995	0,8995	0,8996
	d-BLB50	0,8956	0,8994	0,8995	0,8996	0,8996	0,8996
	d-BLB65	0,8940	0,8993	0,8995	0,8996	0,8996	0,8997
	d-BLB75	0,8956	0,8992	0,8994	0,8995	0,8995	0,8996
	d-BLB85	0,8953	0,8992	0,8994	0,8996	0,8996	0,8997
	d-pMG40	0,8940	0,8997	0,9001	0,9003	0,9004	0,9004
hv medio	d-BLB100	0,8953	0,8994	0,8996	0,8996	0,8997	0,8998
	d-BLB25	0,8860	0,8988	0,8993	0,8995	0,8996	0,8996
	d-BLB50	0,8952	0,8993	0,8995	0,8996	0,8997	0,8997
	d-BLB65	0,8939	0,8994	0,8996	0,8997	0,8997	0,8998
	d-BLB75	0,8949	0,8991	0,8994	0,8996	0,8996	0,8996
	d-BLB85	0,8947	0,8991	0,8995	0,8996	0,8996	0,8997
	d-pMG40	0,8942	0,8993	0,9000	0,9005	0,9007	0,9007

. Analisando a métrica NS, as medianas ficaram entre 24 e 25 soluções válidas encontradas, com os melhores resultados sendo obtidos com a busca local usando as probabilidades de 50% e 75% e pelo d-pMG40. Já o melhor valor médio é alcançado pelo d-pMG40, sendo 2,5% maior que o valor obtido pelo d-BLB50 e 3% superior ao d-BLB75. Em relação ao hv, os resultados são semelhantes para os modelos com a busca local incorporada, embora o melhor resultado também seja alcançado pela abordagem sem busca local (d-pMG40). No geral, a busca local de bateladas mostrou-se eficiente na melhoria das métricas de qualidade das soluções, especialmente em termos de IGD+ e E, com d-BLB75 destacando-se como o melhor modelo. A busca local de bateladas mostrou-se eficiente na qualificação dos indivíduos provenientes do cruzamento (filhos) e, conseqüentemente, na melhoria do conjunto de soluções retornado pelo AG, principalmente em relação às métricas IGD+ e taxa de erro com d-BLB75 destacando-se como o melhor modelo. A maior efetividade da busca local de bateladas pode ser atribuída à capacidade deste operador de explorar a vizinhança das soluções e encontrar o número de bateladas mais adequado para um determinado produto, promovendo uma convergência mais rápida e precisa para a fronteira de Pareto (P^*).

A Tabela 50, apresenta a análise de dominância, calculada através da métrica CS, entre os conjuntos de soluções retornados pelas abordagens que empregam a busca local de bateladas e pelo algoritmo com melhor desempenho até o momento (d-pMG40).

Tabela 50 – Análise de dominância entre as abordagens, considerando o percentual de uso da busca local nas operações de mutação.

	d-BLB100	d-BLB25	d-BLB50	d-BLB65	d-BLB75	d-BLB85	d-pMG40
A\B							
d-BLB100	-	0,347	0,322	0,279	0,302	0,301	0,276
d-BLB25	0,209	-	0,240	0,209	0,230	0,235	0,233
d-BLB50	0,237	0,292	-	0,229	0,251	0,258	0,245
d-BLB65	0,267	0,332	0,304	-	0,287	0,291	0,262
d-BLB75	0,248	0,313	0,288	0,246	-	0,267	0,258
d-BLB85	0,248	0,324	0,292	0,251	0,268	-	0,258
d-pMG40	0,354	0,413	0,404	0,370	0,371	0,361	-

. Como pode ser observado, apesar das melhorias nas métricas apresentadas pelo uso da busca local, o modelo d-pMG40 ainda apresenta a melhor dominância. O modelo d-BLB100 domina 27.6% de d-pMG40, enquanto 35.3% de d-BLB100 são dominados por d-pMG40. Resultado semelhante ocorre ao se comparar d-BLB85, 25.83% dos elementos de um conjunto de não dominadas de d-pMG40 são dominadas por d-BLB85 e 36.14% dos elementos de um conjunto de não dominadas de d-BLB85 são dominadas por d-pMG40.

A Tabela 51 apresenta o teste de hipóteses para os testes avaliados. Percebe-se que, embora o uso da mutação baseada em busca local de bateladas tenha melhorado os valores absolutos das métricas taxa de erro e o IGD+, tais diferenças não são estatisticamente relevantes. Por outro lado, o algoritmo d-pMG40 apresenta um hipervolume significativamente melhor em relação a todas as abordagens que usam o operador de mutação com busca local.

Tabela 51 – Teste de hipóteses comparando d-pMG40 em relação às abordagens que utilizam busca local de bateladas nas operações de mutação.

	d-BLB25	d-BLB50	d-BLB65	d-BLB75	d-BLB85	d-BLB100
hv	>	>	>	>	>	>
NS	>	=	=	=	=	=
IGD+	<	=	=	=	=	=
E	=	=	=	=	=	=

A Tabela 52 apresenta os resultados do teste de hipóteses comparando a abordagem d-BLB75 em relação às outras abordagens utilizando a busca local com percentuais diferentes de pDM . O percentual 25% demonstra resultado significativamente inferior para as métricas IGD+ e o E. Os percentuais 50, 65, 85 e 100% não demonstram evidência de desigualdade. Sugerindo que buscas de número de bateladas com deltas pequenos podem demorar muito para permitir uma exploração eficaz do espaço de busca.

Tabela 52 – Teste de hipóteses para as métricas comparando d-BLB75 em relação com outros percentuais da média do número de bateladas.

	d-BLB50	d-BLB100	d-BLB85	d-BLB65	d-BLB25
hv	=	=	=	=	=
NS	=	=	=	=	=
IGD+	=	=	=	=	<
E	=	=	=	=	<

Com base na análise de dominância e nos resultados dos testes de hipóteses, no geral, o algoritmo d-pMG40 ainda apresenta o melhor desempenho entre as abordagens investigadas. Entretanto, a incorporação do mecanismo de busca local ao operador de mutação também demonstra potencial para a otimização multiobjetivo de problemas restritos, dada a tendência de redução observada nos valores da taxa de erro e do IGD+, indicando sua capacidade de convergir em direção à fronteira de Pareto. Dentre as abordagens que utilizam mutação com busca local de batelada, destaca-se o d-BLB75 que apesar de não apresentar diferenças significativas entre as outras abordagens de busca local, alcançou o

melhor desempenho geral avaliando os valores das métricas, portanto, será o representante dessa estratégia na análise comparativa apresentada na próxima seção.

5.4.6 Análise Comparativa entre os Melhores Modelos

Nesta seção, realiza-se uma análise comparativa entre os melhores modelos investigados (ini-rand, ini-heu, d-BLB75 e d-pMG40) e o modelo de referência (ref). A Tabela 53 apresenta os valores médios e medianas das métricas de desempenho, considerando as 30 execuções de cada abordagem.

Comparando o ini-rand com o modelo de referência, observa-se uma redução de 13,3% na mediana da taxa de erro, uma redução de 78,6% na mediana de IGD+, um aumento de 0,2% na mediana do hipervolume e um aumento de 25,0% em relação ao ao número de soluções válidas presentes no conjunto retornado pelos algoritmos. Esses valores indicam claramente uma superioridade do algoritmo base desenvolvido neste trabalho em relação ao modelo de referência.

Comparando o desempenho das estratégias de inicialização da população (ini-heu vs. ini-rand), a abordagem baseada na inicialização heurística apresenta uma redução de 1,9% na mediana de E, uma redução de 5,4% na mediana de IGD+ e um aumento de 8% na mediana do NS. Analisando o desempenho alcançado com as demais estratégias propostas, em comparação com ini-heu, observa-se que o modelo d-pMG40 proporciona uma redução de 9,8% na mediana de E, uma redução de 35,7% na mediana de IGD+ e uma redução de 7,4% na mediana do NS. Por outro lado, d-BLB75 em comparação com o d-pMG40 proporciona uma diminuição de 8,7% na mediana da taxa de erro e uma redução de 2,1% do IGD+. Ambos os modelos (d-BLB75 e d-pMG40) conseguem aprimorar os valores das métricas de desempenho, em um número menor de gerações do AG (600 contra 1000 das demais abordagens).

Tabela 54 – Análise de dominância entre as abordagens evolutivas propostas e o modelo de referência.

A\B	d-BLB75	d-pMG40	ini-heu	ini-rand	ref
d-BLB75	-	0,258	0,593	0,620	0,921
d-pMG40	0,371	-	0,712	0,736	0,981
ini-heu	0,173	0,119	-	0,435	0,908
ini-rand	0,138	0,102	0,361	-	0,912
ref-s1	0,009	0,006	0,044	0,045	-

Como pode ser observado, o modelo de referência é amplamente dominado por todas as abordagens propostas e domina menos de 5% dos indivíduos de ini-rand e ini-heu e em torno de 1% dos conjuntos de d-BLB75 e d-pMG40. Isso evidencia que as soluções encontradas pelo modelo de referência estão mais distantes da fronteira de Pareto, indicando

Tabela 53 – Comparação entre as principais abordagens investigadas e o modelo de referência.

metrica	gen testName	100	200	300	400	500	600	700	800	900	1000
E mediana	d-BLB75	1,000	0,867	0,800	0,800	0,733	0,700				
	d-pMG40	1,000	0,967	0,900	0,833	0,767	0,767				
	ini-heu	1,000	1,000	1,000	0,967	0,933	0,933	0,900	0,867	0,850	0,850
	ini-rand	1,000	1,000	1,000	0,967	0,967	0,933	0,933	0,933	0,900	0,867
	ref	1,000	1,000	1,000	1,000	1,000	1,000	1,000	1,000	1,000	1,000
E medio	d-BLB75	0,999	0,865	0,794	0,769	0,723	0,722				
	d-pMG40	1,000	0,953	0,892	0,832	0,774	0,756				
	ini-heu	1,000	1,000	0,990	0,970	0,934	0,914	0,896	0,872	0,854	0,843
	ini-rand	1,000	1,000	0,990	0,971	0,954	0,928	0,928	0,904	0,887	0,868
	ref-s1	1,000	1,000	1,000	1,000	1,000	0,999	0,994	0,992	0,992	0,989
IGD+ mediana	d-BLB75	10,799	2,214	1,458	1,486	1,209	1,223				
	d-pMG40	10,656	3,409	2,277	1,809	1,581	1,249				
	ini-heu	33,148	11,702	5,031	4,033	2,725	2,539	2,530	2,160	1,916	1,975
	ini-rand	42,003	11,785	5,361	4,106	3,112	2,772	2,297	1,916	1,889	1,873
	ref-s1	133,202	67,433	43,399	28,409	18,453	13,871	10,808	8,727	8,374	8,772
IGD+ medio	d-BLB75	10,380	2,756	1,910	1,845	1,242	1,331				
	d-pMG40	11,026	4,232	2,710	1,951	1,630	1,479				
	ini-heu	33,687	12,494	5,602	4,155	3,632	2,848	2,613	2,293	2,261	2,250
	ini-rand	40,242	13,196	6,543	4,551	3,317	2,739	2,543	2,285	2,263	2,327
	ref-s1	133,034	70,695	42,248	28,008	18,943	14,618	11,313	9,243	8,787	8,470
NS mediana	d-BLB75	15,00	20,00	22,00	24,00	24,00	25,00				
	d-pMG40	14,00	20,50	22,50	25,50	26,00	25,00				
	ini-heu	7,00	13,00	16,50	19,00	21,00	22,50	23,00	25,00	27,00	27,00
	ini-rand	5,50	12,00	16,50	19,00	20,00	23,00	23,50	23,50	25,00	25,00
	ref-s1	0,00	10,00	13,00	13,50	12,50	15,00	15,00	17,50	19,00	20,00
NS medio	d-BLB75	15,78	19,33	22,30	24,41	24,63	25,19				
	d-pMG40	14,40	20,57	22,63	25,37	25,23	25,97				
	ini-heu	7,23	13,33	16,57	18,80	21,33	23,27	23,57	25,47	26,60	26,80
	ini-rand	5,94	12,28	16,78	19,44	20,34	22,34	24,09	23,84	24,94	24,81
	ref-s1	0,40	10,17	13,20	12,60	12,90	14,47	14,93	17,43	18,70	19,93
hv mediana	d-BLB75	0,8956	0,8992	0,8994	0,8995	0,8995	0,8996				
	d-pMG40	0,8940	0,8997	0,9001	0,9003	0,9004	0,9004				
	ini-heu	0,8655	0,8923	0,8970	0,8981	0,8996	0,8999	0,9001	0,9002	0,9003	0,9003
	ini-rand	0,8459	0,8922	0,8973	0,8988	0,8995	0,8998	0,8999	0,9002	0,9002	0,9003
	ref-s1	0,0000	0,8216	0,8641	0,8832	0,8893	0,8937	0,8949	0,8956	0,8972	0,8977
hv medio	d-BLB75	0,8949	0,8991	0,8994	0,8996	0,8996	0,8996				
	d-pMG40	0,8942	0,8993	0,9000	0,9005	0,9007	0,9007				
	ini-heu	0,8646	0,8920	0,8971	0,8984	0,8994	0,8997	0,9000	0,9001	0,9003	0,9004
	ini-rand	0,8467	0,8913	0,8970	0,8987	0,8993	0,8998	0,8999	0,9002	0,9003	0,9003
	ref-s1	0,2182	0,8244	0,8637	0,8803	0,8891	0,8933	0,8948	0,8960	0,8970	0,8977
t(min) mediana	d-BLB75	7,257	14,513	21,770	29,027	36,283	43,540				
	d-pMG40	3,846	7,692	11,537	15,383	19,229	23,075				
	ini-heu	16,905	33,810	50,715	67,620	84,525	101,430	118,335	135,240	152,145	169,050
	ini-rand	10,222	20,443	30,665	40,887	51,108	61,330	71,552	81,773	91,995	102,217
	ref-s1	6,395	12,790	19,185	25,580	31,975	38,370	44,765	51,160	57,555	63,950
t(min) medio	d-BLB75	7,233	14,467	21,700	28,933	36,167	43,400				
	d-pMG40	3,836	7,672	11,509	15,345	19,181	23,017				
	ini-heu	16,609	33,218	49,827	66,436	83,045	99,654	116,263	132,872	149,480	166,089
	ini-rand	9,788	19,576	29,364	39,152	48,941	58,729	68,517	78,305	88,093	97,881
	ref-s1	6,646	13,293	19,939	26,586	33,232	39,879	46,525	53,172	59,818	66,464

que o modelo apresenta um baixo desempenho em relação à otimização multiobjetivo. Dentre as abordagens propostas neste trabalho, ini-rand e ini-heu têm os piores desempenhos em termos de dominância, com mais de 50% de seus indivíduos sendo dominados pelas soluções dos modelos que empregam os operadores genéticos aprimorados (d-pMG40 e d-BLB85), enquanto dominam em torno de 14 e 17% dos indivíduos de d-BLB85 e 11% o conjunto de d-pMG40. O melhor modelo em relação à dominância é o d-pMG40 que domina 37% das soluções da abordagem com busca local e tem cerca de 26% de seus indivíduos dominados por d-BLB85. Estes resultados evidenciam a melhoria na capacidade dos modelos propostos em encontrar soluções com maior qualidade em relação aos objetivos. Testes de hipóteses foram realizados a fim de avaliar a relevância estatísticas das diferenças observadas nas métricas de desempenho obtidas pelas melhores abordagens evolutivas propostas. A Tabela 55 apresenta os resultados desses testes, comparando o desempenho do modelo d-pMG40 com as demais abordagens investigadas e com o modelo de referência.

Como pode ser observado, o melhor modelo d-pMG40 é superior ao modelo de referência em todas as métricas de desempenho avaliadas. Ele também é melhor que os modelos ini-rand e ini-heu nas métricas multiobjetivo, tendo um desempenho similar em relação ao número de soluções válidas encontradas (métrica NS). Em comparação com a versão que adota a mutação com busca local de bateladas (d-BLB75), ambas as abordagens apresentam desempenhos estatisticamente similares, exceto pelo hipervolume, onde o d-pMG40 demonstra uma ligeira superioridade em relação ao d-BLB75.

Tabela 55 – Teste de hipóteses comparando o algoritmo d-pMG40 com as melhores abordagens investigadas e o modelo de referência.

	d-BLB75	ref-s1	ini-rand	ini-heu
hv	>	>	>	>
NS	=	>	=	=
IGD+	=	<	<	<
E	=	<	<	<

Portanto, com base nas análises realizadas, conclui-se que o melhor modelo investigado é o d-pMG40, o qual incorpora a inicialização heurística e a reinserção particionada, juntamente com o aumento da taxa de cruzamento e a redução no número de gerações totais. Essa combinação propiciou uma exploração mais eficiente, gerando um conjunto de soluções não dominadas mais qualificada, o que pode ser evidenciado pela superioridade observada nas métricas de desempenho.

Também é importante destacar o desempenho alcançado pela abordagem que incorpora a busca local de bateladas ao operador de mutação. Ela apresenta um desempenho estatisticamente similar ao d-pMG40 e muito superior aos demais modelos evolutivos. Embora sem relevância estatística, essa abordagem atingiu valores médios e medianos

superiores para a taxa de erro e IGD+, o que é um indício que a busca local auxilia na convergência das soluções em direção a fronteira de Pareto. Assim, acredita-se que novas pesquisas devem ser realizadas a fim de investigar novas formas de utilização dessa estratégia para melhorar, ainda mais, a eficiência e eficácia do algoritmo evolutivo quando aplicado a problemas de otimização multiobjetivo com restrições.

5.5 Avaliação da Eficiência de Exploração do AGMO

Esse experimento tem por objetivo avaliar a capacidade do AGMO em direcionar a busca para regiões promissoras e encontrar soluções válidas. Nesse sentido, o desempenho do algoritmo é comparado com abordagens que usam estratégias de amostragem (aleatória e heurística) na construção de um superconjunto de soluções, o qual é classificado de acordo com a restrição e os critérios multiobjetivo e os melhores indivíduos são retornados como resposta do algoritmo. A ideia é verificar se a construção de uma quantidade de indivíduos similar àquela avaliada durante todas as execuções do AG é suficiente para gerar um conjunto de soluções com qualidade similar ao obtido pela abordagem evolutiva. Portanto, considerando as 30 execuções da abordagem evolutiva, onde em cada execução o AGMO é evoluído 50 vezes (rodadas) para construir o conjunto de soluções final; as 1000 gerações por rodada; e os 100 indivíduos da população por geração, o modelo de amostragem aleatória usado neste experimento cria e avalia $1,5 \times 10^8$ indivíduos por execução.

Nos modelos de amostragem implementados, o indivíduo é caracterizado pela sequência de número de produtos dentre os 4 fármacos possíveis e número de bateladas que pode variar conforme valores mínimos e máximos por produto, conforme apresentado na

$$Seq_{prod} = \sum_{i=1}^{17} 4^i = 2,29 \times 10^{10} \quad (12)$$

Considerando que, para o produto D, a produção ocorre apenas em múltiplos de 3, existem, em média, 43 opções de quantidade de bateladas por produto. Assim, o total de variações das possíveis sequências de bateladas (Seq_{bat}) é dado por:

$$Seq_{bat} = \sum_{i=1}^{17} 43^i \approx 6,01 \times 10^{27} \quad (13)$$

Portanto, a composição das possíveis sequências de produtos e bateladas resulta em $6,01 \times 10^{37}$ indivíduos diferentes. Foram avaliados três variações do modelo de amostragem de indivíduos:

- **fb-aleat**: nesse modelo, o tamanho do cromossomo varia entre 1 e 17, e os produtos e bateladas de cada gene são escolhidos de forma totalmente aleatória, sendo que o número de bateladas deve satisfazer os valores de mínimo e máximo definidos na Tabela 1.

- **fb-heu**: esse modelo também gera cromossomos que variam de 1 a 17 genes, mas emprega a estratégia de inicialização baseada em heurística para selecionar o produto de cada gene e define o número de bateladas entre 10 e 30.
- **fb-heu***: é uma variação do modelo anterior, que impede a geração de indivíduos pequenos, selecionando a quantidade de genes por cromossomo entre 10 e 17.

Esses modelos foram executados 30 vezes e seus respectivos desempenhos foram comparados com nossa abordagem evolutiva (d-pMG40), em relação ao número de soluções válidas encontradas (NS). Esses resultados são apresentados na Tabela 56.

Tabela 56 – Desempenho dos modelos de amostragem em relação à abordagem evolutiva.

	Abordagens	fb-aleat	fb-heu	fb-heu*	d-pMG40
NS	max	1	0	0	30,00
	mediana	0,000	0,000	0,000	25,00
	média	0,036±0,189	0,000±0,0	0,000±0,0	25,97±2,41
	min	0	0	0	22,00

A análise dos valores apresentados na tabela evidencia a superioridade do modelo evolutivo em relação aos modelos baseados em amostragem, os quais não conseguiram ou encontraram apenas uma solução válida, com um número de avaliações semelhante ao d-pMG40. Esses resultados evidenciam a dificuldade em encontrar indivíduos que atendam à restrição do problema, bem como que nosso AGMO é capaz de direcionar a busca para regiões promissoras, sendo eficiente na exploração de conjunto de soluções não dominadas para o problema de escalonamento de bateladas, em comparação às abordagens baseadas em amostragem.

5.6 Sensibilidade das Abordagens a Variações de Demandas

No problema investigado, o déficit total de estoque ($f_2(\mathbf{x}, SD^s)$) e o backlog total ($\mathbf{e}(\mathbf{x}, SD^s)$) são dependentes das demandas do mercado. Portanto, a aptidão dos indivíduos é diretamente influenciada pelo conjunto de simulações de demanda (SD) usado durante as avaliações. Nos experimentos anteriores, o algoritmo é capaz de encontrar soluções adequadas para aquele cenário bem comportado (demandas similares), onde um único conjunto de simulações durante a evolução da população. Entretanto, essa estratégia pode comprometer a eficiência do algoritmo em ambientes mais voláteis, com maior variabilidade nas vendas. Nesse sentido, este experimento visa analisar a robustez das abordagens evolutivas em relação às variações nos pedidos de venda. Para simular

Tabela 57 – Desempenho das abordagens evolutivas considerando diferentes simulações de demanda.

		din-ref	din-ini-rand	din-ini-heu	din-pMG40	din-BLB
metrica	tipo					
E	dp	0,0000	0,0473	0,0475	0,0610	0,0747
	max	1,0000	1,0000	1,0000	0,7742	0,7742
	mediana	1,0000	0,9677	0,9355	0,6129	0,5484
	medio	1,0000	0,9410	0,9321	0,6285	0,5731
	min	1,0000	0,8387	0,8387	0,4839	0,4516
IGD+	dp	2,7558	1,6421	1,3073	0,5066	0,8049
	max	22,1543	9,9891	7,9069	3,3963	4,4232
	mediana	11,5654	4,3603	5,0174	1,8799	1,8787
	medio	12,1006	4,8225	4,9053	1,9769	2,0929
	min	6,4327	2,5333	2,8785	1,1374	1,0343
NS	dp	3,2439	3,6294	3,7643	2,6607	2,3596
	max	23,0000	30,0000	33,0000	34,0000	31,0000
	mediana	15,0000	20,0000	24,0000	28,0000	25,0000
	medio	15,2885	20,6207	22,7931	27,2727	25,1333
	min	10,0000	13,0000	16,0000	21,0000	20,0000
hv	dp	0,0029	0,0011	0,0008	0,0007	0,0002
	max	0,8993	0,9006	0,9006	0,9031	0,9011
	mediana	0,8940	0,8994	0,8999	0,9012	0,9001
	medio	0,8934	0,8991	0,8995	0,9011	0,9001
	min	0,8874	0,8956	0,8982	0,8991	0,8996

tal volatilidade, o conjunto de simulações de Monte Carlo é modificado a cada 10 gerações (PCDGA, conforme Subseção 5.1.2). Os principais algoritmos investigados foram reavaliados nesse novo ambiente dinâmico, a saber: o modelo de referência (din-ref), o modelo baseado nos métodos clássicos (din-ini-rand), o modelo com inicialização heurística (din-ini-heu), e o modelo focado na diversidade (din-pMG40) e o modelo que emprega a mutação com busca local (din-BLB). Dessa forma, espera-se avaliar a capacidade de adaptação das abordagens evolutivas às mudanças de demanda, sem comprometer consideravelmente a qualidade do conjunto de soluções encontrados. A Tabela 57, mostra os valores das métricas de desempenho de cada algoritmo nesse ambiente dinâmico.

Como pode ser observado, os algoritmos possuem hipervolumes semelhantes, sendo o melhor resultado obtido pelo modelo din-pMG40, que apresenta um aumento no valor médio 0,8% em relação à din-ref. Avaliando o número de soluções não dominadas válidas (NS), o din-pMG40 apresenta uma melhoria da média de 78,3% e da mediana de 86,6% em relação à din-ref. Em relação ao IGD+, o din-pMG40 apresenta uma redução da média

de 83,6% e da mediana de 83,7% em relação à din-ref. Avaliando a taxa de erro (E), o modelo din-BLB apresenta uma redução da média de 37,1% e da mediana de 45,1% em relação à din-ref, além de uma redução de 10,5% em relação ao din-pMG40. Os resultados indicam que os modelos din-pMG40 e din-BLB mantêm um desempenho robusto, com melhorias significativas na taxa de erro, número de soluções válidas e IGD+ em relação ao modelo de referência (din-ref), com uma ligeira vantagem para o din-pMG40, mostrando a capacidade de adaptação em ambientes com maior volatilidade.

A Tabela 58 apresenta os resultados de dominância entre os algoritmos investigados. O modelo din-pMG40 apresenta uma leve superioridade em relação à dominância, com 41,0% de suas soluções dominadas por din-BLB, enquanto domina 43,5% das soluções desse algoritmo. Comparado aos demais algoritmos, a superioridade do din-pMG40 é mais significativa, dominando mais de 90% das soluções. A maior diferença é em relação ao din-ref, onde 99,7% das soluções são dominadas por din-pMG40 e domina apenas 0,1% das soluções do algoritmo.

Tabela 58 – Análise de dominância entre as abordagens evolutivas em ambientes dinâmicos.

A \ B	din-BLB	din-ini-heu	din-ini-rand	din-pMG40	din-ref
din-BLB	-	0,829	0,846	0,410	0,965
din-ini-heu	0,068	-	0,451	0,044	0,940
din-ini-rand	0,057	0,464	-	0,038	0,933
din-pMG40	0,435	0,937	0,935	-	0,997
din-ref	0,003	0,028	0,029	0,001	-

O teste de hipóteses apresentado na Tabela 59, avalia a relevância estatística das diferenças entre os valores das métricas de desempenho dos AGMOs investigados. Como esperado, o din-pMG40 alcançou valores significativamente superiores aos demais algoritmos em praticamente todas as métricas consideradas. A única exceção é em relação ao d-BLB75, que obteve uma melhor taxa de erro e um IGD+ similar, mas sendo superado quanto ao hipervolume e o número de soluções válidas encontradas. Isso indica, sugerindo que o aumento da diversidade proporcionado pelos operadores propostos em din-pMG40 é benéfico para a exploração e adaptação do algoritmo às mudanças no conjunto de demandas.

Tabela 59 – Teste de hipóteses comparando o din-pMG40 em relação às demais abordagens em um ambiente com variação na demanda.

	din-ref	din-ini-rand	din-ini-heu	d-BLB
hv	>	>	>	>
NS	>	>	>	>
IGD+	<	<	<	=
E	<	<	<	>

CONCLUSÃO

Neste trabalho, investigamos modelos de AGMO aplicados aos POMR no escalonamento de bateladas de produtos farmacêuticos com demandas estocásticas. Nossa investigação partiu do AGMO publicado em (JANKAUSKAS; FARID, 2019), que foi proposto para uma instância do escalonamento de bateladas aplicados em dados reais de uma produção farmacêutica, sujeito a restrições de demanda. A partir da reprodução desse algoritmo, que denominamos modelo de referência ref-s1, conduzimos uma série de investigações para aumentar o desempenho dos AGMO nesse problema, que se mostrou complexo e com baixa convergência para soluções viáveis. Propusemos quatro modelos principais: **ini-rand**, que incorpora inicialização aleatória, mutação e crossover ao modelo de referência (ref-s1); **ini-heu**, que incorpora à ini-rand uma heurística de inicialização; **d-pMG40**, que utiliza uma estratégia de reinserção particionada. No modelo ini-heu. Na próxima seção apresentamos as principais conclusões dessa investigação. Esta pesquisa teve como objetivo principal investigar e aprimorar os operadores genéticos dos AGMO para o problema de POMR no contexto de SB de produtos farmacêuticos, utilizando SMC para selecionar as soluções capazes de lidar com variações esperadas de demanda. O foco foi encontrar soluções abordando os desafios de convergência de soluções viáveis de maior qualidade em relação às principais métricas avaliadas nessa investigação o hipervolume (hv), distância generacional invertida mais (IGD+), taxa de erro (E) e número de soluções válidas encontradas na população final(NS). O primeiro modelo proposto ini-rand em comparação com ref-s1, apresentou uma redução de 71,9% para o IGD+ e 11,9% para o E, além de um aumento de 24,9% para o NSV. Isso sugere que os operadores propostos de inicialização, mutação e crossover apresentaram significativa melhoria na busca por soluções mais próximas de P^* e melhoraram a convergência das soluções viáveis. O segundo modelo ini-heu em comparação com a referência ref-s1 aprimorou os resultados obtidos por ini-rand, apresentando uma redução de 73,4% no IGD+ e 13,4% no E, acompanhada por um aumento de 34,9% no NS. Além disso, demonstrou superioridade na métrica CS em comparação com o modelo de referência, com 88,5% das soluções do modelo original sendo dominadas por soluções de ini-rand e apenas 5,4% das soluções de ini-heu sendo do-

minadas por soluções de referência. Dessa forma, os modelos propostos ini-rand e ini-heu se mostraram eficazes para melhorar a exploração do espaço de soluções restritas durante a busca do AGMO.

A análise gráfica revelou que o modelo ini-heu foi capaz de encontrar fronteiras próximas de P^* . No entanto, também apresentou áreas com alta concentração de soluções na métrica NS. Apesar das melhorias, a métrica E permaneceu alta em 85,2%, denotando a complexidade do problema restrito. Os dois últimos modelos, d-pMG40 e d-BLB, buscaram aprimorar o modelo ini-heu para melhorar a convergência ao Pareto. Analisando o modelo d-pMG40 em comparação com ini-heu demonstrou uma redução de 9,8% no E; redução de 36,7% no IGD+; redução de 0,007% no hv e redução de 7,4% no NS. Em relação à dominância, d-pMG40 domina 71,2% das soluções de ini-heu; e ini-heu domina apenas 11,9% das soluções de d-pMG40. Por fim, d-pMG40 apresentou melhoria estatisticamente significativa em relação ao E, IGD+ e hv. Isto sugere que o operador de seleção particionado baseado em restrições beneficia a exploração de soluções mais próximas de P^* através de uma maior diversidade de soluções, devido ao conjunto de soluções não restritas incorporadas. Sobre o operador de mutação incorporando a busca local, é importante destacar que o modelo d-BLB, apesar de não demonstrar diferença significativa em relação ao d-pMG40, apresentou potencial em relação ao d-pMG40 com uma redução de 8,7% no E; redução de 2,1% na mediana do IGD+ e redução de 0,09% na mediana do hv.

Portanto, apresenta potencial para ser investigado utilizando outras estratégias de busca local. Por fim, os quatro modelos propostos foram avaliados utilizando um ambiente de testes com simulações de Monte Carlo, variando a cada 10 gerações. Os resultados obtidos com o novo ambiente de teste confirmam a eficácia dos modelos desenvolvidos, utilizando um ambiente mais dinâmico. Ao se avaliar os resultados com o novo ambiente de teste, comparando com o modelo de referência (din-ref), os melhores modelos são os baseados nos modelos d-pMG40 e d-BLB, avaliados para o novo cenário de teste e denominados, respectivamente, din-pMG40 e din-BLB. O modelo din-pMG40, em comparação com din-ref, apresenta melhorias estatisticamente significativas de 38,7% para o E, redução de 83,7% para o IGD+, redução de 86,6% para o NS e uma redução de 0,8% para o hv. O din-pMG40 domina 99,7% de din-ref e apenas 0,1% de din-pMG40 são dominados por din-ref. O modelo din-BLB, em relação ao din-pMG40, demonstra uma redução estatisticamente significativa para o E de 10,5%. Adicionalmente, apresenta uma redução de 0,1% para o IGD+, uma pequena redução de 0,1% para o hv e uma redução de 10,7% para o NS. Analisando a dominância, 43,47% de din-BLB são dominados por din-pMG40, enquanto 40,99% de din-pMG40 são dominados por din-BLB. Isto indica que a incorporação da busca local no número de bateladas auxilia na redução de E para o caso de ambiente mais dinâmico.

Em relação às hipóteses realizados, a adoção de estratégias para diversificação da popu-

lação inicial. A partir dos experimentos realizados permitiu que o AGMO encontre um conjunto de soluções mais qualificado, conforme demonstrado pelos resultados do modelo ini-heu. Adicionalmente, a melhoria da diversificação permitiu a simplificar-se os operadores genéticos sem comprometer a qualidade.

Em relação à utilização de heurísticas, a incorporação de na inicialização da população aprimorou a convergência do algoritmo. Adicionalmente, mesmo em ambientes mais diversos o AGMO foi capaz de proporcionar melhores resultados em comparação com o modelo de referência.

A utilização de técnicas de busca local podem melhorar a convergência e a qualidade do conjunto de soluções encontrado conforme demonstrado com a estratégia d-BLB. Contudo a melhor estratégia permanece sendo o modelo d-pMG40, que não utiliza a busca local. O custo computacional para avaliação de aptidão é significativo, logo como trabalhos futuros estratégias baseadas em modelos *surrogate* podem ser investigados.

Ao comparar-se o melhor modelo com buscas aleatórias, estas não são capazes de realizar uma busca efetiva no problema restrito, uma vez que poucas soluções válidas foram encontradas. Desta forma confirmando a dificuldade do problema restrito.

Por fim a combinação das diferentes estratégias para lidar com o problema restrito demonstraram melhores resultados em relação ao modelo tradicional. Adicionalmente, os modelos desenvolvidos apresentam resultados mesmo em ambientes dinâmicos de demanda.

6.1 Trabalhos Futuros

A utilização de estratégias de busca local demonstrou-se relevante para a redução da taxa de erro médio nos ambientes de teste com maior variabilidade durante o processo evolutivo. Contudo, o desafio permanece em reduzir o custo computacional de avaliação dos objetivos e restrições sem uma perda substancial na qualidade das soluções. Uma alternativa é realizar uma investigação da sensibilidade em relação ao número de simulações de Monte Carlo mais adequado para permitir a significância estatística desejada e proporcionar uma avaliação do indivíduo com um custo computacional aceitável. Outra potencial estratégia a ser investigada é a utilização de modelos preditivos (*surrogate model*), em substituição às simulações de Monte Carlo, visando agilizar o processo de avaliação dos indivíduos e permitir uma busca local mais extensiva. Além disso, investigar outras heurísticas e técnicas de diversidade para melhorar a variabilidade e a exploração das soluções, especialmente nas faixas inferiores do objetivo de déficit de estoque $f_2(x)$, pode contribuir ainda mais para o aprimoramento do processo evolutivo, reduzindo seu tempo de execução, melhorando a eficiência do algoritmo e qualificando o conjunto de soluções encontrado, tanto em relação aos critérios multiobjetivos, quanto na quantidade de soluções válidas. Quebrar linha, criando um novo parágrafo que deve começar com: Outras técnicas inteligentes, como colônia de formigas e *simulated annealing*, também

podem ser adaptadas para direcionar a construção incremental de soluções potenciais, em um contexto multiobjetivo. Adicionalmente, investigações em relação ao critério de parada como, por exemplo, o tempo de processamento, podem ser realizadas a fim de mapear seus respectivos impactos na qualidade do conjunto de soluções viáveis não dominadas resultante. Em relação à comparação com outras abordagens baseadas em amostragem aleatória, uma questão ainda em aberto é se o aumento significativo na quantidade de soluções geradas por essas abordagens (ex: 10 vezes a quantidade de soluções avaliadas pelo AGMO) possibilitaria a obtenção de um conjunto de soluções válidas similar ou próximo àquele alcançado pelos algoritmos evolutivos propostos ou é necessária a adoção de uma heurística mais refinada para lidar com problemas de otimização com restrição. Também é importante comparar o desempenho das nossas abordagens com outros métodos de otimização bem estabelecidos, como a programação linear inteira. O estudo de algoritmos híbridos que integrem diferentes métodos inteligentes a fim de aprimorar a busca por soluções válidas é outra frente de pesquisa que pode ser explorada.

Outra evolução natural da pesquisa é ampliar a complexidade do problema investigado, adicionando novas linhas de produção e produtos ou considerando um número maior de objetivos, tornando-o um problema com muitos objetivos (*many-objective*). Assim, será possível investigar outros tipos de algoritmos para o escalonamento com restrições. Por fim, recomenda-se explorar outras métricas de desempenho multiobjetivo na análise das abordagens, em especial aquelas que permitem uma comparação da distribuição espacial entre os conjuntos de não dominância encontrados (ex: *attainment function*).

6.2 Contribuições Em Produção Bibliográfica

Os resultados alcançados ao longo desta pesquisa possibilitaram a elaboração de artigos científicos:

- ❑ (KOHARA; OLIVEIRA; MARTINS, 2023b): publicado no Encontro Nacional de Inteligência Artificial e Computacional (ENIAC), esse artigo apresenta uma análise do operador de mutação otimizada, cruzamento e inicialização aleatória da população, ou seja, se refere ao desenvolvimento do modelo ini-rand.
- ❑ (KOHARA; OLIVEIRA; MARTINS, 2023a): publicado no Congresso Internacional de Ferramentas baseadas em Inteligência Artificial (*International Conference on Tools with Artificial Intelligence - ICTAI*), esse artigo apresenta uma análise do impacto em adotar o operador de inicialização heurística da população inicial no desempenho do AGMO. Ou seja, se refere ao desenvolvimento do modelo ini-heu.
- ❑ (KOHARA; OLIVEIRA; MARTINS, 2024): aceito na Conferência Brasileira de Sistemas Inteligentes (*Brazilian Conference On Intelligent Systems - BRACIS*), esse artigo apresenta uma análise da influência dos operadores de reinserção particionada

e de mutação incorporada à busca local de bateladas no desempenho da otimização multiobjetivo. Ou seja, se refere ao desenvolvimento dos modelos d-pMG40 e d-BLB.

Além dos artigos supracitados, pretende-se preparar e submeter uma nova publicação, contemplando os resultados e análises referentes à abordagem d-pMG40, comparando seu desempenho com as demais abordagens investigadas, considerando ambientes estáticos e dinâmicos.

Referências

- AHMADI, E. et al. A multi objective optimization approach for flexible job shop scheduling problem under random machine breakdown by evolutionary algorithms. **Comput. and Operations Research**, v. 73, p. 56–66, 2016. Disponível em: <<https://doi.org/10.1016/j.cor.2016.03.009>>.
- AZADEH, A.; GOLDANSAZ, S. M.; ZAHEDI-ANARAKI, A. H. Solving and optimizing a bi-objective open shop scheduling problem by a modified genetic algorithm. **International Journal of Advanced Manufacturing Technology**, Springer-Verlag London Ltd, v. 85, p. 1603–1613, 7 2016. Disponível em: <<https://doi.org/10.1007/s00170-015-8069-z>>.
- BAECK, T. Selective pressure in evolutionary algorithms: A characterization of selection mechanisms. **IEEE Conference on Evol. Comput.**, v. 1, p. 57–62, 1994. Disponível em: <<https://doi.org/10.1109/ICEC.1994.350042>>.
- BAGHERI, A. et al. An artificial immune algorithm for the flexible job-shop scheduling problem. **Future Generation Computer Systems**, v. 26, n. 4, p. 533–541, 2010. Disponível em: <<https://doi.org/10.1016/j.future.2009.10.004>>.
- BATTELLE. **Biopharmaceutical Industry-Sponsored Clinical Trials: Impact on State Economies**. 2015. 1–22 p.
- BEZDAN, T. et al. Multi-objective task scheduling in cloud computing environment by hybridized bat algorithm. **J. of Intelligent and Fuzzy Systems**, v. 42, p. 411–423, 1 2022. Disponível em: <<https://doi.org/10.3233/JIFS-219200>>.
- CAI, X.; GAO, L.; LI, X. Efficient Generalized Surrogate-Assisted Evolutionary Algorithm for High-Dimensional Expensive Problems. **IEEE Trans. on Evol. Comput.**, v. 24, p. 365–379, 2020. Disponível em: <<https://doi.org/10.1109/TEVC.2019.2919762>>.
- CASTILLO, E. D. An application of network scheduling optimization in a pharmaceutical firm. **Comput. in Industry**, v. 18, p. 279–287, 1992. Disponível em: <[https://doi.org/10.1016/0166-3615\(92\)90031-H](https://doi.org/10.1016/0166-3615(92)90031-H)>.
- CASTRO, P. M.; GROSSMANN, I. E.; ROUSSEAU, L. M. Decomposition Techniques for Hybrid MilP/cP models applied to Scheduling and Routing Problems. In: _____. **Hybrid Optimization: The Ten Years of CPAIOR**. New York, NY: Springer New York, 2011. p. 135–167. Disponível em: <https://doi.org/10.1007/978-1-4419-1644-0_4>.

CHEN, Y. et al. Digital twins in pharmaceutical and biopharmaceutical manufacturing: a literature review. **Processes**, MDPI, v. 8, n. 9, p. 1088, 2020. Disponível em: <<https://doi.org/10.3390/pr8091088>>.

COTT, B. J.; MACCHIETTO, S. Minimizing the effects of batch process variability using online schedule modification. **Comput. and Chem. Eng.**, v. 13, p. 105–113, 1989. Disponível em: <[https://doi.org/10.1016/0098-1354\(89\)89011-8](https://doi.org/10.1016/0098-1354(89)89011-8)>.

DARWIN, C. **On the origin of species by means of natural selection**. [S.l.]: Good Press, 2023.

DEB, K. Multi-objective Genetic Algorithms: Problem Difficulties and Construction of Test Problems. 1999. Disponível em: <<https://doi.org/10.3390/mca26010005>>.

DEB, K. et al. Adenosine in difficult aneurysm surgeries: report of two cases. **Journal of Neuroanaesthesiology and Critical Care**, Thieme Medical and Scientific Publishers Private Ltd., v. 1, n. 01, p. 066–068, 2014. Disponível em: <<https://doi.org/10.3390/mca26010005>>.

_____. A fast and elitist multiobjective genetic algorithm: NsgA-II. **IEEE Trans. Evol. Comput.**, v. 6, p. 182–197, 2002. Disponível em: <<https://doi.org/10.3390/mca26010005>>.

DEB, K.; ROY, P. C.; HUSSEIN, R. Surrogate Modeling Approaches for Multiobjective Optimization: Methods, Taxonomy, and Results. **Mathematical and Computational Applications**, v. 26, p. 5, 2020. Disponível em: <<https://doi.org/10.3390/mca26010005>>.

DIAZ-GOMEZ, P.; HOUGEN, D. Initial Population for Genetic Algorithms: A metric Approach. In: . [S.l.: s.n.], 2007. p. 43–49.

DREXL, A.; KIMMS, A. Lot sizing and scheduling —Survey and extensions. **European J. of Operational Research**, v. 99, p. 221–235, 1997. Disponível em: <[https://doi.org/10.1016/S0377-2217\(97\)00030-1](https://doi.org/10.1016/S0377-2217(97)00030-1)>.

FONSECA, C. M.; FLEMING, P. J. Genetic algorithms for multiobjective optimization: formulation, discussion and generalization. In: . [s.n.], 1993. v. 93, n. July, p. 416–423. Disponível em: <<https://doi.org/10.1109/3468.650319>>.

_____. Multiobjective optimization and multiple constraint handling with evolutionary algorithms - Part I: A unified formulation. **IEEE Trans. Systems, Man, and Cybernetics Part A: Systems and Humans**, v. 28, p. 26–37, 1998. Disponível em: <<https://doi.org/10.1109/3468.650319>>.

FORTIN, F. A. et al. DEap: Evol. algorithms made easy. **J. of Machine Learning Research**, v. 13, p. 2171–2175, 2012.

FRAGA, L. M. et al. Estrategias evolutivas para o ajuste de parametros de um modelo epidemiologico baseado em automatos celulares probabilisticos. Universidade Federal de Uberlandia, 2021. Disponível em: <<http://doi.org/10.14393/ufu.di.2022.5042>>.

FRANÇA, T. P. et al. A comparative analysis of moeas considering two discrete optimization problems. In: IEEE. **2017 Brazilian Conference on Intelligent Systems (BRACIS)**. IEEE, 2017. p. 402–407. Disponível em: <<https://doi.org/10.1016/j.asoc.2017.09.017>>.

_____. Estratégias bio-inspiradas aplicadas em problemas discretos com muitos objetivos. Universidade Federal de Uberlândia, 2018. Disponível em: <<http://dx.doi.org/10.14393/ufu.di.2019.368>>.

FRANCCA, T. P.; MARTINS, L. G.; OLIVEIRA, G. M. Maconds: Many-objective ant colony optimization based on non-dominated sets. In: IEEE. **2018 IEEE Congress on Evol. Comput.** 2018. p. 1–8. Disponível em: <<https://doi.org/10.1109/CEC.2018.8477958>>.

FU, Y. et al. Scheduling dual-objective stochastic hybrid flow shop with deteriorating jobs via bi-population evolutionary algorithm. **IEEE Trans. on Systems, Man, and Cybernetics: Systems**, IEEE, v. 50, n. 12, p. 5037–5048, 2019. Disponível em: <<https://doi.org/10.1109/TSMC.2019.2907575>>.

_____. Stochastic multi-objective integrated disassembly-reprocessing-reassembly scheduling via fruit fly optimization algorithm. **J. of Cleaner Production**, v. 278, 1 2021. Disponível em: <<https://doi.org/10.1016/j.jclepro.2020.123364>>.

GAO, K. et al. A review on swarm intelligence and evolutionary algorithms for solving flexible job shop scheduling problems. **IEEE/CAA J. of Automatica Sinica**, v. 6, p. 904–916, 2019. Disponível em: <<https://doi.org/10.1016/10.1109/JAS.2019.1911540>>.

GAO, K. Z. et al. An improved artificial bee colony algorithm for flexible job-shop scheduling problem with fuzzy processing time. **Expert Syst. Appl.**, Pergamon Press, Inc., USA, v. 65, n. C, p. 52–67, dec 2016. Disponível em: <<https://doi.org/10.1016/j.eswa.2016.07.046>>.

GEORGIADIS, G. P. Optimal Production Planning and Scheduling of Mixed Batch and Continuous Ind. Processes. 2021.

GOLBERG, D. E. Genetic algorithms in search, optimization, and machine learning. **Addion wesley**, 1989.

GOLDBERG, D. E.; DEB, K. A comparative Analysis of Selection Schemes Used in Genetic Algorithms. p. 69–93, 1991. Disponível em: <<https://doi.org/10.1023/A:1022602019183>>.

GOLDBERG, D. E. et al. Rapid, accurate optimization of difficult problems using messy genetic algorithms. In: **Proceedings of the Fifth International Conference on Genetic Algorithms**. [s.n.], 1993. p. 59–64. Disponível em: <<https://doi.org/10.1023/A:1022602019183>>.

GOLDBERG, D. E.; HOLLAND, J. H. Genetic Algorithms and Machine Learning. **Machine Learning**, v. 3, p. 95–99, 1988. Disponível em: <<https://doi.org/10.1023/A:1022602019183>>.

- HABIB, A. et al. A multiple Surrogate Assisted Decomposition Based Evolutionary Algorithm for Expensive Multi/Many-Objective Optimization. **IEEE Trans. on Evol. Comput.**, v. 23, p. 1000–1014, 02 2019. Disponível em: <<https://doi.org/10.1109/TEVC.2019.2899030>>.
- HAN, Y. et al. Discrete evolutionary multi-objective optimization for energy-efficient blocking flow shop scheduling with setup time. **Applied Soft Computing Journal**, Elsevier Ltd, v. 93, 8 2020. Disponível em: <<https://doi.org/10.1016/j.asoc.2020.106343>>.
- HARRER, S. et al. Artificial Intelligence for Clinical Trial Design. **Trends in Pharmacological Sciences**, v. 40, p. 577–591, 2019. Disponível em: <<https://doi.org/10.1016/j.tips.2019.05.005>>.
- HOLLAND, J. H. **Adaptation in natural and artificial systems: an introductory analysis with applications to biology, control, and artificial intelligence**. MIT press, 1992. Disponível em: <<https://doi.org/10.7551/mitpress/1090.001.0001>>.
- HONG, M. S. et al. Challenges and opportunities in biopharmaceutical manufacturing control. **Comput. and Chem. Eng.**, v. 110, p. 106–114, 2018. Disponível em: <<https://doi.org/10.1016/j.compchemeng.2017.12.007>>.
- HONKOMP, S. et al. The curse of reality —why process scheduling optimization problems are difficult in practice. **Comput. and Chem. Eng.**, v. 24, n. 2, p. 323–328, 2000. Disponível em: <[https://doi.org/10.1016/S0098-1354\(98\)00296-8](https://doi.org/10.1016/S0098-1354(98)00296-8)>.
- HONKOMP, S. J.; MOCKUS, L.; REKLAITIS, G. V. Robust scheduling with processing time uncertainty. **Comput. and Chem. Eng.**, v. 21, p. 1055–1060, 1997. Disponível em: <[https://doi.org/10.1016/S0098-1354\(98\)00296-8](https://doi.org/10.1016/S0098-1354(98)00296-8)>.
- _____. A framework for schedule evaluation with processing uncertainty. **Comput. and Chem. Eng.**, v. 23, p. 595–609, 1999. Disponível em: <[https://doi.org/10.1016/S0098-1354\(98\)00296-8](https://doi.org/10.1016/S0098-1354(98)00296-8)>.
- ISHIBUCHI, H. et al. Modified distance calculation in generational distance and inverted generational distance. **Evol. Multi-criterion Optimization**, p. 110–125, 2015. Disponível em: <https://doi.org/10.1007/978-3-319-15892-1_8>.
- JÁHN-ERDÖS, S.; KÖVÁRI, B. Exploring the Potential of a Genetic Algorithm on a Real-World Complex Scheduling Problem. In: **2022 9th International Conference on Soft Computing and Machine Intelligence (ISCMI)**. [s.n.], 2022. p. 113–117. Disponível em: <<https://doi.org/10.1109/ISCMI56532.2022.10068465>>.
- JANKAUSKAS, K.; FARID, S. S. Multi-objective biopharma capacity planning under uncertainty using a flexible genetic algorithm approach. **Comput. and Chem. Eng.**, v. 128, p. 35–52, 9 2019. Disponível em: <<https://doi.org/10.1016/j.compchemeng.2019.05.023>>.
- JANKAUSKAS, K.; PAPAGEORGIOU, L. G.; FARID, S. S. Fast genetic algorithm approaches to solving discrete-time mixed integer linear programming problems of capacity planning and scheduling of biopharmaceutical manufacture. **Comput. and Chem. Eng.**, v. 121, p. 212–223, 2 2019. Disponível em: <<https://doi.org/10.1016/j.compchemeng.2018.09.019>>.

- JIA, Z.-H.; GAO, L.-Y.; ZHANG, X.-Y. A new history-guided multi-objective evolutionary algorithm based on decomposition for batching scheduling. **Expert Systems With Applications**, v. 141, 2020. Disponível em: <<https://doi.org/10.1016/j.eswa.2019.112920>>.
- JOHNSON, P. O. **Statistical methods in research**. Prentice-Hall, 1949. Disponível em: <<https://doi.org/10.1016/j.tibtech.2013.05.011>>.
- JUNGBAUER, A. Continuous downstream processing of biopharmaceuticals. **Trends in Biotechnology**, v. 31, p. 479–492, 2013. Disponível em: <<https://doi.org/10.1016/j.tibtech.2013.05.011>>.
- KANAKAMEDALA, K. B.; REKLAITIS, G. V.; VENKATASUBRAMANIAN, V. Reactive schedule modification in multipurpose batch chemical plants. **Ind. and Eng. chemistry research**, v. 33, p. 77–90, 1994. Disponível em: <<https://doi.org/10.1021/ie00025a011>>.
- KHOUKHI, F. E.; BOUKACHOUR, J.; ALAOUI, A. E. H. The “Dual-Ants Colony”: A novel hybrid approach for the flexible job shop scheduling problem with preventive maintenance. **Computers and Industrial Engineering**, v. 106, p. 236–255, 2017. Disponível em: <<https://doi.org/10.1016/j.cie.2016.10.019>>.
- KOHARA, D. T.; OLIVEIRA, G. M. B. D.; MARTINS, L. G. A. Improving a multiobjective evolutionary algorithm applied to batch scheduling in pharmaceutical manufacturing. In: IEEE. **2023 IEEE 35th International Conference on Tools with Artificial Intelligence (ICTAI)**. 2023. p. 399–403. Disponível em: <<https://doi.org/10.5753/eniac.2023.234618>>.
- KOHARA, D. T.; OLIVEIRA, G. M. B. de; MARTINS, L. G. A. Aprimoramento de um modelo evolutivo multiobjetivo aplicado ao problema de sequenciamento de bateladas na manufatura de produtos farmacêuticos. In: SBC. **Anais do XX Encontro Nacional de Inteligência Artificial e Computacional**. 2023. p. 1099–1113. Disponível em: <<https://doi.org/10.5753/eniac.2023.234618>>.
- _____. Enhancing multiobjective genetic algorithms for pharmaceutical batch scheduling: A study on partitioned selection with constraints and mutation with greedy local search strategy. In: SBC. 2024. Disponível em: <<https://doi.org/10.5753/eniac.2023.234618>>.
- KOPANOS, G. M.; MÉNDEZ, C. A.; PUIGJANER, L. MIP-based decomposition strategies for large-scale scheduling problems in multiproduct multistage batch plants: A benchmark scheduling problem of the pharmaceutical industry. **European journal of operational research**, v. 207, p. 644–655, 2010. Disponível em: <<https://doi.org/10.1016/j.ejor.2010.06.002>>.
- KÜSTER, T. et al. Multi-objective optimization of energy-efficient production schedules using genetic algorithms. **Optimization and Eng.**, 2021. Disponível em: <<https://doi.org/10.1007/s11081-021-09691-3>>.
- LAFETA, T. et al. MEanDs: A many-objective Evolutionary Algorithm based on Non-dominated Decomposed Sets applied to multicast routing. **Applied Soft Computing**, Elsevier, v. 62, p. 851–866, 2018. Disponível em: <<https://doi.org/10.1016/j.asoc.2017.09.017>>.

- LAÍNEZ, J. M.; SCHAEFER, E.; REKLAITIS, G. V. Challenges and opportunities in enterprise-wide optimization in the pharmaceutical industry. **Comput. and Chem. Eng.**, v. 47, p. 19–28, 2012. Disponível em: <<https://doi.org/10.1016/j.compchemeng.2012.07.002>>.
- LAKHDAR, K. et al. Medium term planning of biopharmaceutical manufacture using mathematical programming. **Biotechnology progress**, v. 21, p. 1478–1489, 2005. Disponível em: <<https://doi.org/10.1021/bp0501571>>.
- LAUMANNS, M. et al. Combining convergence and diversity in evolutionary multiobjective optimization. **Evol. computation**, v. 10, p. 263–282, 2002. Disponível em: <<https://doi.org/10.1162/106365602760234108>>.
- LI, Y. et al. An optimization method for energy-conscious production in flexible machining job shops with dynamic job arrivals and machine breakdowns. **Journal of Cleaner Production**, Elsevier Ltd, v. 254, 5 2020. Disponível em: <<https://doi.org/10.1016/j.knosys.2019.02.027>>.
- LI, Z. et al. An elitist nondominated sorting hybrid algorithm for multi-objective flexible job-shop scheduling problem with sequence-dependent setups. **Knowledge-Based Systems**, v. 173, p. 83–112, 6 2019. Disponível em: <<https://doi.org/10.1016/j.knosys.2019.02.027>>.
- LIU, Y. et al. Optimal design of an integrated discontinuous water using network coordinating with a central continuous regeneration unit. **Ind. and Eng. chemistry research**, v. 48, p. 10924–10940, 2009. Disponível em: <<https://doi.org/10.1021/ie9000053>>.
- LOLIC, M. **NDA at the FDA Professional Affairs and Stakeholder Engagement**. [S.l.], 2016.
- MAJOZI, T.; SEID, E. R.; LEE, J.-Y. **Synthesis, design, and resource optimization in batch chemical plants**. [S.l.: s.n.], 2015.
- MANN, H. B.; WHITNEY, D. R. On a test of whether one of two random variables is stochastically larger than the other. **The annals of mathematical statistics**, p. 50–60, 1947. Disponível em: <<https://doi.org/10.1214/aoms/1177730491>>.
- MARCHETTI, P. A.; MENDEZ, C. A.; CERDA, J. Simultaneous lot sizing and scheduling of multistage batch processes handling multiple orders per product. **Ind. and Eng. Chemistry Research**, v. 51, p. 5762–5780, 2012. Disponível em: <<https://doi.org/10.1021/ie202275y>>.
- MICHALEWICZ, Z. **GAs: Why Do They Work?** Springer, 1996. Disponível em: <<https://doi.org/10.1016/j.jacbts.2016.03.002>>.
- MITCHELL, M. **An introduction to genetic algorithms**. [S.l.]: MIT press, 1998.
- NORMAN, G. A. V. Drugs, Devices, and the Fda: Part 1: An Overview of Approval Processes for Drugs. **JACC: Basic to Translational Science**, v. 1, p. 170–179, 2016.
- OYEBOLU, F. B. Meta-Heuristic and Hyper-Heuristic Scheduling Tools for Biopharmaceutical Production. 2019. Disponível em: <https://doi.org/10.1007/978-1-4419-8191-2_17>.

- PACCIARELLI, D.; MELONI, C.; PRANZO, M. Models and methods for production scheduling in the pharmaceutical industry. **Planning Production and Inventories in the Extended Enterprise: A State-of-the-Art Handbook, Volume 2**, p. 429–459, 2011.
- PEZZELLA, F.; MORGANTI, G.; CIASCHETTI, G. A genetic algorithm for the Flexible Job-shop Scheduling Problem. **Computers and Operations Research**, v. 35, n. 10, p. 3202–3212, 2008. Disponível em: <<https://doi.org/10.1016/j.cor.2007.02.014>>.
- QUADT, D.; KUHN, H. A taxonomy of flexible flow line scheduling procedures. **European J. of Operational Research**, v. 178, p. 686–698, 2007. Disponível em: <<https://doi.org/10.1016/j.ejor.2006.01.042>>.
- RIDLEY, M. Evolutionary vistas. **Nature**, v. 381, p. 287–288, 1996. Disponível em: <<https://doi.org/10.1038/381287a0>>.
- ROTHLAUF, F. Design of Modern Heuristics: Principles and Application. **Springer Berlin Heidelberg**, 2011. Disponível em: <<https://doi.org/10.1007/978-3-540-72962-4>>.
- SAFE, M. et al. On stopping criteria for genetic algorithms. In: SPRINGER. **Advances in Artificial Intelligence—SBIA 2004: 17th Brazilian Symposium on Artificial Intelligence, Sao Luis, Maranhao, Brazil, September 29-October 1, 2004. Proceedings 17**. 2004. p. 405–413. Disponível em: <https://doi.org/10.1007/978-3-540-28645-5_41>.
- SAMAK-KULKARNI, S. M.; RAJHANS, N. R. Determination of optimum inventory model for minimizing total inventory cost. **Procedia Eng.**, v. 51, p. 803–809, 2013. Disponível em: <<https://doi.org/10.1016/j.proeng.2013.01.115>>.
- SAND, G. et al. Engineered versus standard evolutionary algorithms: A case study in batch scheduling with recourse. **Comput. and Chem. Eng.**, v. 32, p. 2706–2722, 2008. Disponível em: <<https://doi.org/10.1016/j.compchemeng.2007.09.006>>.
- SCHAFFER, J. D. Some experiments in machine learning using vector evaluated genetic algorithms. Vanderbilt Univ., Nashville, TN (USA), 1985.
- SCHWEFEL, H. **Evolution and Optimum Seeking**. [S.l.]: Wiley, 1995. (Sixth Generation Computer Technologies).
- SHAH, N. **Planning and scheduling-single-and multisite planning and scheduling: Current status and future challenges**. [S.l.: s.n.], 1998. v. 94. 75-90 p.
- SILVER, E. A. A heuristic for selecting lot size quantities for the case of a deterministic time-varying demand rate and discrete opportunities for replenishment. **Prod. Inventory Manage.**, v. 2, p. 64–74, 1973.
- SOUIER, M.; DAHANE, M.; MALIKI, F. An NsgA-II-based multiobjective approach for real time routing selection in a flexible manufacturing system under uncertainty and reliability constraints. **The International Journal of Advanced Manufacturing Technology**, v. 100, p. 2813–2829, 2019. Disponível em: <<https://doi.org/10.1007/s00170-018-2897-6>>.

SRINIVAS, N.; DEB, K. Multiobjective optimization using nondominated sorting in genetic algorithms. **Evol. Comput.**, MIT Press, v. 2, p. 221–248, 1994. Disponível em: <<https://doi.org/10.1162/evco.1994.2.3.221>>.

STÜTZLE, T. Local search algorithms for combinatorial problems. **Darmstadt University of Technology PhD Thesis**, v. 20, 1998.

TIAN, Y. et al. A coevolutionary Framework for Constrained Multiobjective Optimization Problems. **IEEE Trans. Evol. Comput.**, v. 25, p. 102–116, 2021. Disponível em: <<https://doi.org/10.1016/j.jclepro.2015.09.097>>.

TILL, J. et al. A hybrid evolutionary algorithm for solving two-stage stochastic integer programs in chemical batch scheduling. **Comput. and Chem. Eng.**, v. 31, p. 630–647, 2007. Disponível em: <<https://doi.org/10.1016/j.compchemeng.2006.09.003>>.

ULLMAN, J. NP-complete scheduling problems. **Journal of Computer and System Sciences**, v. 10, n. 3, p. 384–393, 1975. Disponível em: <[https://doi.org/10.1016/S0022-0000\(75\)80008-0](https://doi.org/10.1016/S0022-0000(75)80008-0)>.

VELEZ, S.; MARAVELIAS, C. T. Multiple and nonuniform time grids in discrete-time Mip models for chemical production scheduling. **Comput. and Chem. Eng.**, v. 53, p. 70–85, 2013. Disponível em: <<https://doi.org/10.1016/j.compchemeng.2013.01.014>>.

WANG, Y. et al. Global and Local Surrogate-Assisted Differential Evolution for Expensive Constrained Optimization Problems With Inequality Constraints. **IEEE Trans. on Cybernetics**, v. 49, n. 5, p. 1642–1656, 2019. Disponível em: <<https://doi.org/10.1109/TCYB.2018.2809430>>.

ZHANG, B. et al. An automatic multi-objective evolutionary algorithm for the hybrid flowshop scheduling problem with consistent sublots. **Knowledge-Based Systems**, v. 238, p. 107819, 2022. Disponível em: <<https://doi.org/10.1016/j.jclepro.2015.09.097>>.

ZHANG, J. et al. Review of job shop scheduling research and its new perspectives under Industry 4.0. **J. of Intelligent Manufacturing**, v. 30, p. 1809–1830, 2019. Disponível em: <<https://doi.org/10.1016/j.jclepro.2015.09.097>>.

ZHANG, R.; CHIONG, R. Solving the energy-efficient job shop scheduling problem: A multi-objective genetic algorithm with enhanced local search for minimizing the total weighted tardiness and total energy consumption. **J. of Cleaner Production**, v. 112, p. 3361–3375, 2016. Disponível em: <<https://doi.org/10.1016/j.jclepro.2015.09.097>>.

ZHOU, S. et al. A multi-objective differential evolution algorithm for parallel batch processing machine scheduling considering electricity consumption cost. **Comput. and Operations Research**, v. 96, p. 55–68, 2018. Disponível em: <<https://doi.org/10.1016/j.cor.2018.04.009>>.

ZITZLER, E. et al. SPea2: Improving the strength pareto evolutionary algorithm. 2001. Disponível em: <<https://doi.org/10.1109/4235.797969>>.

ZITZLER, E.; THIELE, L. An evolutionary algorithm for multiobjective optimization: The strength pareto approach. **TIK-report**, v. 43, 1998. Disponível em: <<https://doi.org/10.1109/4235.797969>>.

_____. Multiobjective Evol. Algorithms: A comparative Case Study and the Strength Pareto Approach. **IEEE Trans. Evol. Comput.**, v. 3, 1999. Disponível em: <<https://doi.org/10.1109/4235.797969>>.

Apêndices

Investigações adicionais

Nesta seção será apresentada investigações adicionais realizadas durante o mestrado.

A.1 Sensibilidade em Relação à Mutação de Gene

Nesta seção, avaliou-se a alteração da probabilidade de adição de gene para permitir a redução do número de genes. Inicialmente, analisou-se a sensibilidade da probabilidade de mutação de gene (pMG) utilizando uma probabilidade de adição de gene fixa de 50% (pAG), selecionando pMG com performance semelhante ou superior à d-pRe60. Em seguida, fixou-se a pMG selecionadae, analisando a sensibilidade do algoritmo em relação à pAG.

A.1.1 Análise da Probabilidade de Mutação de Gene

Avaliou-se a sensibilidade do modelo d-pRe60 (conforme Seção 5.4) em relação à probabilidade de mutação de gene (pMG), considerando, no caso em que ocorre a mutação de gene, uma probabilidade de adição de gene (pAG) de 50%. Avaliaram-se inicialmente as probabilidades de pMG de 7% (pMG7), 20% (pMG20), 35% (pMG35), 40% (pMG40), 45% (pMG45), 60% (pMG60) e 80% (pMG80). Compararam-se os resultados com o modelo d-pRe60, que utiliza um pMG de 7% e pAG de 100%. .

Os resultados são apresentados na Tabela 60. Em relação ao hv, têm-se medianas semelhantes para todas as probabilidades. Em relação à mediana do NS, os resultados são similares, variando entre 25 e 26. Avaliando a mediana do E, ao se comparar d-pRe60 com pMG7, a redução de pAG de 100% para 50% resultou em um aumento de cerca de 9,1%. Contudo, os resultados são semelhantes para as probabilidades entre 35% e 60%, com um aumento de 4,5% em relação ao d-pRe60. Em relação à mediana do IGD+, comparando d-pRe60 com pMG7, assim como para o E, a redução de pAG de 100% para 50% apresentou um aumento de 83,0%. Contudo, ao se comparar o modelo pMG40 com o d-pRe60, este apresentou uma redução de 3,8%. Isto indica que pode-se encontrar re-

Tabela 60 – Desempenho do algoritmo em função da probabilidade de mutação de gene.

metrica	gen testName	100	200	300	400	500	600
E mediana	d-pMG20	1,000	0,933	0,900	0,883	0,833	0,800
	d-pMG35	1,000	0,950	0,900	0,850	0,800	0,767
	d-pMG40	1,000	0,967	0,900	0,833	0,767	0,767
	d-pMG45	1,000	0,950	0,900	0,833	0,800	0,767
	d-pMG60	1,000	0,967	0,867	0,850	0,800	0,767
	d-pMG7	1,000	0,983	0,933	0,900	0,850	0,800
	d-pMG80	1,000	0,933	0,900	0,833	0,817	0,800
	d-pRe60	1,000	0,933	0,833	0,800	0,767	0,733
E medio	d-pMG20	0,998	0,958	0,900	0,863	0,828	0,803
	d-pMG35	0,999	0,943	0,877	0,847	0,799	0,784
	d-pMG40	1,000	0,953	0,892	0,832	0,774	0,756
	d-pMG45	0,999	0,943	0,888	0,839	0,788	0,753
	d-pMG60	0,998	0,948	0,863	0,840	0,800	0,763
	d-pMG7	1,000	0,977	0,935	0,883	0,860	0,812
	d-pMG80	0,996	0,941	0,893	0,849	0,807	0,784
	d-pRe60	0,992	0,920	0,848	0,798	0,764	0,726
IGD+ mediana	d-pMG20	11,618	3,786	2,142	2,405	2,225	1,890
	d-pMG35	10,505	3,468	2,208	1,606	1,466	1,386
	d-pMG40	10,656	3,409	2,277	1,809	1,581	1,249
	d-pMG45	9,957	3,869	2,073	1,845	1,529	1,391
	d-pMG60	10,184	3,001	1,932	2,016	1,745	1,784
	d-pMG7	17,995	5,078	3,614	2,619	2,334	2,378
	d-pMG80	8,485	2,630	1,755	1,821	1,438	1,551
	d-pRe60	6,634	2,252	1,683	1,699	1,244	1,299
IGD+ medio	d-pMG20	11,998	4,479	2,836	2,997	2,820	1,868
	d-pMG35	10,578	3,942	2,557	2,143	1,792	1,490
	d-pMG40	11,026	4,232	2,710	1,951	1,630	1,479
	d-pMG45	10,504	4,144	2,592	2,221	1,952	1,874
	d-pMG60	10,533	3,526	2,313	2,161	1,960	1,870
	d-pMG7	18,938	5,646	4,015	2,970	2,717	2,470
	d-pMG80	8,760	3,293	1,873	1,934	1,586	1,726
	d-pRe60	6,842	2,631	2,031	1,849	1,409	1,469
NS mediana	d-pMG20	13,00	19,50	23,00	24,00	25,00	25,50
	d-pMG35	14,00	20,00	23,00	24,50	26,00	26,00
	d-pMG40	14,00	20,50	22,50	25,50	26,00	25,00
	d-pMG45	13,00	20,50	22,00	24,00	25,00	25,00
	d-pMG60	15,00	19,50	23,00	25,00	25,50	26,50
	d-pMG7	11,00	19,50	22,00	24,00	26,00	26,00
	d-pMG80	13,00	20,00	22,00	24,00	25,00	26,00
	d-pRe60	16,00	23,00	25,00	26,00	26,00	26,00
NS medio	d-pMG20	13,93	20,10	23,13	24,50	25,63	26,00
	d-pMG35	14,13	20,17	23,03	24,60	26,13	26,37
	d-pMG40	14,40	20,57	22,63	25,37	25,23	25,97
	d-pMG45	13,80	20,50	22,33	23,67	24,90	25,13
	d-pMG60	13,95	20,31	23,12	25,07	25,64	26,02
	d-pMG7	11,75	19,50	22,75	23,79	26,00	26,43
	d-pMG80	13,50	19,63	22,20	24,07	25,30	26,23
	d-pRe60	16,34	22,88	25,03	25,78	26,12	26,50
hv mediana	d-pMG20	0,8927	0,8986	0,9001	0,9003	0,9003	0,9004
	d-pMG35	0,8944	0,8996	0,9001	0,9003	0,9004	0,9004
	d-pMG40	0,8940	0,8997	0,9001	0,9003	0,9004	0,9004
	d-pMG45	0,8948	0,8997	0,9002	0,9003	0,9004	0,9004
	d-pMG60	0,8945	0,8997	0,9002	0,9003	0,9003	0,9004
	d-pMG7	0,8887	0,8980	0,8998	0,9002	0,9003	0,9003
	d-pMG80	0,8943	0,8997	0,9001	0,9003	0,9004	0,9004
	d-pRe60	0,8970	0,9000	0,9003	0,9004	0,9004	0,9004
hv medio	d-pMG20	0,8927	0,8987	0,8999	0,9002	0,9005	0,9006
	d-pMG35	0,8940	0,8992	0,9000	0,9003	0,9005	0,9005
	d-pMG40	0,8942	0,8993	0,9000	0,9005	0,9007	0,9007
	d-pMG45	0,8947	0,8994	0,9001	0,9004	0,9005	0,9006
	d-pMG60	0,8945	0,8995	0,9002	0,9004	0,9005	0,9006
	d-pMG7	0,8874	0,8980	0,8996	0,9002	0,9004	0,9005
	d-pMG80	0,8950	0,8994	0,9001	0,9004	0,9004	0,9005
	d-pRe60	0,8973	0,8999	0,9003	0,9004	0,9005	0,9006

sultados semelhantes ao d-pRe60, permitindo também a redução do número de genes.

A Tabela 61 apresenta os resultados em relação à dominância. Para os modelos com pAG de 50%, a melhor cobertura ocorre para d-pMG45, que domina 31,3% de d-pRe60; contudo, d-pRe60 domina 35,0% de d-pMG45. Portanto, d-pRe60 apresenta uma dominância um pouco superior. O segundo melhor modelo é o d-pMG40, que domina 30,3% de d-pRe60; contudo, d-pRe60 domina 36,3% de d-pMG40.

Tabela 61 – Métrica CS considerando as probabilidades de mutação de gene.

A\B	d-pMG20	d-pMG35	d-pMG40	d-pMG45	d-pMG60	d-pMG7	d-pMG80	d-pRe60
d-pMG20	-	0,332	0,297	0,289	0,307	0,426	0,337	0,272
d-pMG35	0,380	-	0,317	0,307	0,328	0,447	0,353	0,286
d-pMG40	0,404	0,369	-	0,325	0,343	0,467	0,375	0,303
d-pMG45	0,414	0,377	0,348	-	0,356	0,474	0,387	0,314
d-pMG60	0,402	0,369	0,335	0,320	-	0,461	0,374	0,298
d-pMG7	0,297	0,257	0,232	0,224	0,241	-	0,277	0,217
d-pMG80	0,381	0,351	0,317	0,308	0,323	0,445	-	0,283
d-pRe60	0,427	0,397	0,363	0,350	0,367	0,492	0,401	-

Avaliando o teste de hipóteses (conforme Tabela 62), observa-se que os modelos pMG40 e pMG45 não apresentam diferenças significativas em relação ao d-pRe60. Contudo, pMG40 apresentou uma mediana de IGD+ 3,8% inferior e, ao realizar uma mutação de gene, permite também a probabilidade de redução do número de genes (com pAG igual a 50%), portanto foi selecionado como modelo base para as análises seguintes.

Avaliando o teste de hipóteses (Conforme Tabela 62), têm se que o modelo pMG40 e pMG45 não apresentam diferenças significativas em relação à d-pRe60. Contudo, pMG40AGen50 apresentou uma mediana em relação ao IGD+ 3,8% inferior e ao realizar uma mutação de gene permite também a probabilidade de redução de número de genes (com pAG igual à 50%), portanto foi selecionado como modelo base para as análises seguintes.

Tabela 62 – Teste de hipóteses comparando d-pRe60 com as demais probabilidades de mutação de gene.

	d-pMG7	d-pMG10	d-pMG20	d-pMG30	d-pMG35	d-pMG40	d-pMG45	d-pMG50	d-pMG60	d-pMG80
hv	>	>	>	>	>	=	=	=	=	=
NS	=	=	=	=	=	=	=	=	=	=
IGD+	<	<	=	=	=	=	=	=	<	=
E	<	<	<	<	<	=	=	<	<	<

Devido aos melhores resultados em relação às métricas multiobjetivo, especialmente a menor mediana de IGD+, somado a resultados semelhantes de dominância e à semelhança estatística, selecionou-se o modelo pMG40 para os experimentos seguintes.

A.1.2 Análise da Probabilidade de Adição de Gene

Nesta seção, avaliou-se a sensibilidade em relação à probabilidade de adição de gene (pAG). A partir do melhor modelo selecionado em Apêndice A.1.1, que adota um pMG de 40% (pMG40) e pAG de 50%, verificou-se o efeito de diferentes valores para a probabilidade de adição de gene. A Tabela 63, mostra os valores médios e medianas para as métricas de desempenho, alcançados a partir de 30 execuções do AGMO ao empregar pAG de 0% (pMG40pAG0), 25% (pMG40pAG25), 75% (pMG40pAG75) e 100% (pMG40pAG100)., respectivamente. Analisando a tabela, em relação à mediana do NS, os testes apresentaram resultados semelhantes, com o melhor resultado ocorrendo para pMG40pAG100, apresentando um aumento de 8,0% em relação a pMG40. A exceção foi pAG igual a 0% (pMG40pAG0), que não conseguiu encontrar soluções em nenhuma das execuções. Isso indica a importância da adição de novos genes para encontrar indivíduos válidos e demonstra uma elevada dependência do modelo em relação à mutação para encontrar soluções válidas, diferente do objetivo inicial da mutação, que é proporcionar pequenas alterações nos indivíduos. Uma hipótese é que a adição de genes assemelha-se a uma busca local. Em relação ao hv, resultados semelhantes foram encontrados.

Em relação à mediana do E, a abordagem pMG40pAG100 apresentou o melhor resultado, com uma redução de 4,3% em relação a pMG40. Avaliando a mediana do IGD+, a abordagem pMG40pAG100 apresentou uma redução de 4,0% em relação a pMG40. A análise das métricas sugere que o modelo pMG40pAG100 é o mais eficiente, apresentando as melhores reduções em E e IGD+, além de um aumento significativo em NS. Assim, o modelo pMG40pAG100 foi selecionado para os experimentos subsequentes.

Conforme pode ser observado na Tabela 64, que apresenta a relação de dominância entre as variações de probabilidade investigadas, a estratégia d-pMG40pAG100 apresenta um conjunto de soluções não dominadas ligeiramente superior, dominando 36,04% dos elementos de d-pMG40, enquanto esse último domina 30,33% dos elementos de d-pMG40pAG100. Avaliando o teste de hipóteses na Tabela 65, observa-se que d-pMG40pAG75 e d-pMG40pAG100 não apresentam diferença significativa em comparação com d-pMG40. Portanto, devido à capacidade de permitir também a redução do número de genes, o modelo pMG40 será selecionado. Adicionalmente, destaca-se que pMG40pAG0 não foi capaz de encontrar soluções válidas, o que indica a importância da capacidade de adicionar genes aos indivíduos para o problema. Uma hipótese é que o operador de mutação, especialmente a operação de adição de novo gene, assemelha-se a um operador de busca local.

A.2 Cruzamento Produto-Batelada

O operador de mutação demonstra significativa sensibilidade em relação à probabilidade de adição de gene, sugerindo que o cruzamento não está proporcionando diversidade

Tabela 63 – Métricas avaliação para os testes

metrica	gen testName	100	200	300	400	500	600
E mediana	d-pMG40	1,000	0,967	0,900	0,833	0,767	0,767
	d-pMG40pAG0	1,000	1,000	1,000	1,000	1,000	1,000
	d-pMG40pAG100	1,000	0,900	0,833	0,800	0,750	0,733
	d-pMG40pAG25	1,000	1,000	0,917	0,867	0,833	0,800
	d-pMG40pAG75	1,000	0,933	0,833	0,800	0,800	0,767
E medio	d-pMG40	1,000	0,953	0,892	0,832	0,774	0,756
	d-pMG40pAG0	1,000	1,000	1,000	1,000	1,000	1,000
	d-pMG40pAG100	0,994	0,910	0,849	0,786	0,752	0,724
	d-pMG40pAG25	1,000	0,979	0,920	0,869	0,826	0,796
	d-pMG40pAG75	0,997	0,921	0,832	0,808	0,776	0,767
IGD+ mediana	d-pMG40	10,656	3,409	2,277	1,809	1,581	1,249
	d-pMG40pAG100	7,444	2,262	1,744	1,684	1,383	1,199
	d-pMG40pAG25	16,290	5,373	3,058	2,036	1,702	1,692
	d-pMG40pAG75	6,598	2,741	1,751	1,590	1,339	1,250
IGD+ medio	d-pMG40	11,026	4,232	2,710	1,951	1,630	1,479
	d-pMG40pAG100	7,370	2,621	2,082	1,974	1,483	1,344
	d-pMG40pAG25	16,236	5,438	3,401	2,239	2,031	2,000
	d-pMG40pAG75	7,055	2,813	1,857	1,862	1,592	1,444
NS mediana	d-pMG40	14,00	20,50	22,50	25,50	26,00	25,00
	d-pMG40pAG0	0,00	0,00	0,00	0,00	0,00	0,00
	d-pMG40pAG100	15,50	21,00	24,50	26,00	26,00	27,00
	d-pMG40pAG25	11,50	19,00	23,00	24,00	24,00	25,50
	d-pMG40pAG75	15,50	21,00	24,50	25,50	27,00	26,50
NS medio	d-pMG40	14,40	20,57	22,63	25,37	25,23	25,97
	d-pMG40pAG0	0,00	0,00	0,00	0,00	0,00	0,00
	d-pMG40pAG100	15,93	21,70	24,50	26,73	26,73	26,77
	d-pMG40pAG25	11,77	19,57	22,33	24,00	24,30	25,87
	d-pMG40pAG75	15,23	20,67	24,47	25,53	27,03	26,90
hv mediana	d-pMG40	0,8940	0,8997	0,9001	0,9003	0,9004	0,9004
	d-pMG40pAG0	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	d-pMG40pAG100	0,8975	0,9000	0,9003	0,9004	0,9004	0,9005
	d-pMG40pAG25	0,8898	0,8981	0,8999	0,9002	0,9003	0,9003
	d-pMG40pAG75	0,8960	0,8999	0,9003	0,9004	0,9004	0,9004
hv medio	d-pMG40	0,8942	0,8993	0,9000	0,9005	0,9007	0,9007
	d-pMG40pAG0	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	d-pMG40pAG100	0,8971	0,9000	0,9003	0,9005	0,9007	0,9007
	d-pMG40pAG25	0,8892	0,8981	0,8995	0,9001	0,9004	0,9004
	d-pMG40pAG75	0,8956	0,8997	0,9004	0,9004	0,9005	0,9006

Tabela 64 – Métrica CS considerando as probabilidades de adição de gene.

A\B	d-pMG40	d-pMG40pAG0	d-pMG40pAG100	d-pMG40pAG25	d-pMG40pAG75
d-pMG40	-	1,000	0,303	0,458	0,332
d-pMG40pAG0	0,000	-	0,000	0,000	0,000
d-pMG40pAG100	0,360	1,000	-	0,476	0,359
d-pMG40pAG25	0,249	1,000	0,226	-	0,251
d-pMG40pAG75	0,340	1,000	0,303	0,460	-

Tabela 65 – Teste de hipóteses comparando d-pMG40 com as demais probabilidades de adição de gene.

	d-pMG40pAG0	d-pMG40pAG25	d-pMG40pAG75	d-pMG40pAG100
hv	>	>	=	=
NS	>	=	=	=
IGD+		<	=	=
E	<	<	=	=

suficiente entre os pares de produto e número de bateladas para explorar soluções de forma eficaz. Portanto, investigou-se a incorporação do cruzamento uniforme (Seção 4.3) com probabilidade de troca de produto (pCP) e probabilidade de troca de batelada (pCB) independentes. Isto é, em vez de realizar a troca conjunta do par produto e número de bateladas, os elementos de um gene podem ser trocados independentemente, de acordo com pCP e pCB. Durante os experimentos, avaliou-se as combinações de pCP=0% e pCB=100% (d-pCP0pCB100), pCP=100% e pCB=0% (d-pCP100pCB0), pCP=100% e pCB=100% (d-pCP100pCB100), pCP=100% e pCB=50% (d-pCP100pCB50), pCP=50% e pCB=100% (d-pCP50pCB100), e pCP=75% e pCB=100% (d-pCP75pCB100). Os resultados foram comparados com o modelo d-pMG40 selecionado em Apêndice A.1.2, que utiliza o cruzamento de um ponto. A Tabela 66 apresenta as medianas e médias das métricas obtidas a partir de 30 execuções do AGMO, considerando as combinações investigadas. Em relação à mediana do hv, os resultados foram semelhantes entre as configurações avaliadas, contudo d-pCP100pCB0, apresentou um valor 5,3% inferior em comparação com d-pMG40. Avaliando a mediana do NS, o d-pCP100pCB100 apresentou um resultado 4% melhor em relação ao d-pMG40, sendo que o pior resultado ocorre para d-pCP100pCB0 com uma redução de 8%. Em relação à mediana do E, os modelos d-pMG40 e d-pCP100pCB100 apresentam o melhor resultado, enquanto o pior resultado

Tabela 66 – Desempenho do algoritmo em função da probabilidade de adição de gene adotada.

metrica	gen testName	100	200	300	400	500	600
E mediana	d-pCP0pCB100	1,000	0,967	0,933	0,900	0,867	0,833
	d-pCP100pCB0	1,000	1,000	1,000	0,967	0,933	0,900
	d-pCP100pCB100	1,000	0,967	0,900	0,833	0,800	0,767
	d-pCP100pCB50	1,000	1,000	0,933	0,900	0,867	0,833
	d-pCP50pCB100	1,000	0,967	0,933	0,867	0,833	0,800
	d-pCP75pCB100	1,000	1,000	0,933	0,867	0,833	0,817
	d-pMG40	1,000	0,967	0,900	0,833	0,767	0,767
E medio	d-pCP0pCB100	1,000	0,973	0,929	0,891	0,852	0,817
	d-pCP100pCB0	1,000	0,995	0,986	0,957	0,936	0,915
	d-pCP100pCB100	1,000	0,969	0,886	0,842	0,796	0,761
	d-pCP100pCB50	1,000	0,986	0,928	0,883	0,856	0,831
	d-pCP50pCB100	0,999	0,976	0,914	0,866	0,829	0,798
	d-pCP75pCB100	0,999	0,975	0,907	0,865	0,835	0,797
	d-pMG40	1,000	0,953	0,892	0,832	0,774	0,756
IGD+ mediana	d-pCP0pCB100	11,737	4,163	3,582	2,257	1,724	1,324
	d-pCP100pCB0	25,319	7,836	4,819	3,511	2,629	2,420
	d-pCP100pCB100	13,439	3,909	2,594	1,553	1,578	1,553
	d-pCP100pCB50	15,297	4,301	3,130	2,132	1,891	1,518
	d-pCP50pCB100	11,667	4,799	2,545	2,035	1,536	1,555
	d-pCP75pCB100	13,043	4,612	2,869	1,787	1,965	1,561
	d-pMG40	10,656	3,409	2,277	1,809	1,581	1,249
IGD+ medio	d-pCP0pCB100	11,791	5,012	4,357	2,749	2,398	1,693
	d-pCP100pCB0	26,162	8,176	5,160	3,839	3,118	2,594
	d-pCP100pCB100	13,624	4,596	3,203	2,268	2,033	1,892
	d-pCP100pCB50	15,474	5,087	3,418	2,591	2,415	1,796
	d-pCP50pCB100	12,718	4,714	2,790	2,508	2,174	1,954
	d-pCP75pCB100	13,514	4,969	3,473	2,420	2,169	1,816
	d-pMG40	11,026	4,232	2,710	1,951	1,630	1,479
NS mediana	d-pCP0pCB100	14,00	19,00	22,00	23,00	23,50	24,00
	d-pCP100pCB0	10,50	16,00	19,00	20,00	22,00	23,00
	d-pCP100pCB100	12,00	20,00	22,00	24,00	24,50	26,00
	d-pCP100pCB50	11,00	19,00	22,00	24,00	25,50	25,00
	d-pCP50pCB100	12,50	20,00	21,00	24,00	24,00	24,50
	d-pCP75pCB100	12,00	20,00	21,50	23,00	24,00	25,50
	d-pMG40	14,00	20,50	22,50	25,50	26,00	25,00
NS medio	d-pCP0pCB100	13,63	20,07	22,00	22,90	23,83	24,03
	d-pCP100pCB0	10,76	16,06	18,88	20,65	22,38	22,65
	d-pCP100pCB100	12,07	20,40	21,93	23,93	24,80	25,60
	d-pCP100pCB50	11,03	19,60	21,97	23,73	25,37	25,57
	d-pCP50pCB100	12,27	20,80	21,70	23,47	24,10	24,50
	d-pCP75pCB100	11,83	19,38	21,62	23,21	24,33	25,38
	d-pMG40	14,40	20,57	22,63	25,37	25,23	25,97
hv mediana	d-pCP0pCB100	0,8932	0,8992	0,8999	0,9001	0,9003	0,9003
	d-pCP100pCB0	0,8789	0,8963	0,8977	0,8987	0,8993	0,8999
	d-pCP100pCB100	0,8909	0,8994	0,9000	0,9002	0,9003	0,9003
	d-pCP100pCB50	0,8900	0,8985	0,8998	0,9000	0,9002	0,9003
	d-pCP50pCB100	0,8907	0,8993	0,8999	0,9001	0,9002	0,9003
	d-pCP75pCB100	0,8913	0,8992	0,8998	0,9000	0,9002	0,9003
	d-pMG40	0,8940	0,8997	0,9001	0,9003	0,9004	0,9004
hv medio	d-pCP0pCB100	0,8930	0,8988	0,8995	0,9000	0,9001	0,9003
	d-pCP100pCB0	0,8778	0,8960	0,8978	0,8988	0,8993	0,8996
	d-pCP100pCB100	0,8905	0,8989	0,8997	0,9001	0,9003	0,9003
	d-pCP100pCB50	0,8891	0,8984	0,8997	0,9000	0,9002	0,9003
	d-pCP50pCB100	0,8907	0,8988	0,8997	0,9001	0,9002	0,9003
	d-pCP75pCB100	0,8911	0,8988	0,8998	0,9000	0,9002	0,9003
	d-pMG40	0,8942	0,8993	0,9000	0,9005	0,9007	0,9007

ocorre para d-pCP100pCB0. Isso indica que o par de informações de produto e batelada pode transmitir informações relevantes para o aprimoramento das soluções. Comparando o modelo d-pCP0pCB100 com o d-pCP100pCB100, observa-se um aumento de 8,6% no E; e para d-pCP100pCB0 em relação ao d-pCP100pCB100, um aumento de 17,4%. Isso sugere que a pCB possui maior importância em comparação com a pCP no E. Analisando a mediana do IGD+, o melhor resultado ocorre para o d-pMG40, e entre os modelos com o novo cruzamento, o melhor resultado é para d-pCP0pCB100 com um aumento de 5,9%. Os resultados indicam que incorporação da troca independente de produtos e bateladas não apresenta melhorias em relação à d-pMG40. A maior importância da pCB em comparação à pCP sugere que a troca de bateladas desempenha um papel mais crucial na melhoria das soluções. Além disso, a sensibilidade do operador de mutação à adição de genes destaca a importância da diversificação das soluções, indicando que a adição de novos genes pode funcionar como um operador de busca local, auxiliando na exploração de soluções válidas.

Tabela 67 – Métrica CS considerando as probabilidades de adição de gene.

A\B	d-pCP0pCB100	d-pCP100pCB0	d-pCP100pCB100	d-pCP100pCB50	d-pCP50pCB100	d-pCP75pCB100	d-pMG40
d-pCP0pCB100	-	0,575	0,310	0,373	0,341	0,356	0,275
d-pCP100pCB0	0,207	-	0,186	0,219	0,213	0,214	0,164
d-pCP100pCB100	0,370	0,604	-	0,407	0,367	0,388	0,294
d-pCP100pCB50	0,330	0,564	0,292	-	0,331	0,341	0,258
d-pCP50pCB100	0,348	0,576	0,310	0,384	-	0,364	0,280
d-pCP75pCB100	0,346	0,572	0,311	0,380	0,344	-	0,281
d-pMG40	0,430	0,642	0,391	0,467	0,432	0,444	-

. De acordo com a Tabela 67, o d-pMG40 permanece sendo o melhor modelo em termos de dominância. O d-pMG40 domina 39,1% de d-pCP100pCB100, enquanto d-pCP100pCB100 domina 29,4% de d-pMG40. Em comparação com os modelos d-pCP0pCB100 e d-pCP100pCB0, o d-pCP0pCB100 domina 57,5% de d-pCP100pCB0, enquanto d-pCP100pCB0 domina apenas 20,7% de d-pCP0pCB100. Isso pode indicar uma maior importância do número de bateladas em relação ao produto, conforme verificado também na comparação entre d-pCP100pCB100 e d-pCP100pCB0, onde d-pCP100pCB100 domina 60,4% de d-pCP100pCB0, e d-pCP100pCB0 domina apenas 18,6% de d-pCP100pCB100. Adicionalmente, d-pCP100pCB100 domina 37,0% de d-pCP0pCB100, enquanto d-pCP0pCB100 domina 31,0% de d-pCP100pCB100.

Tabela 68 – Teste de hipóteses considerando as probabilidades de adição de gene.

	d-pCP100pCB0	d-pCP100pCB50	d-pCP0pCB100	d-pCP50pCB100	d-pCP75pCB100	d-pCP100pCB100
hv	>	>	>	>	>	>
NS	>	=	>	=	=	=
IGD+	<	=	=	=	=	=
E	<	<	<	<	<	=

O teste de hipóteses, apresentado na Tabela 68, confirma que o melhor modelo ainda é o d-pMG40. Entre os modelos utilizando o cruzamento, o d-pCP100pCB100 apresenta semelhança estatística em comparação ao d-pMG40 para todas as métricas, exceto o hv. Isso sugere que o cruzamento original simultâneo do par produto e batelada é relevante para o problema. Adicionalmente, ao comparar o d-pCP100pCB0, observa-se uma piora significativa em todas as métricas, enquanto o d-pCP0pCB100 apresenta semelhança em relação ao IGD+. Isto sugere uma maior relevância do número de bateladas em comparação ao produto.

Não há indícios de que o cruzamento proposto, permitindo a troca independente de produtos e bateladas, melhora a qualidade das soluções em comparação com o cruzamento simultâneo do par produto e batelada. Assim, o melhor modelo continua sendo o d-pMG40. No entanto, há indícios de que o número de bateladas tem maior significância em comparação com o produto neste cruzamento em relação ao IGD+. Isto pode indicar potencial para investigar métodos de busca local em relação ao número de bateladas.