



CEFET-MG

**CENTRO FEDERAL DE EDUCAÇÃO
TECNOLÓGICA DE MINAS GERAIS**

Diretoria de Pesquisa e Pós-Graduação

Programa de Mestrado em Modelagem
Matemática e Computacional

HEURÍSTICAS PARA O PROBLEMA DE SEQUENCIAMENTO EM UMA MÁQUINA COM PENALIDADES POR ANTECIPAÇÃO E ATRASO DA PRODUÇÃO

Dissertação de Mestrado apresentada ao Programa de Pós-Graduação em Modelagem Matemática e Computacional do CEFET-MG, como parte dos requisitos necessários para a obtenção do título de Mestre em Modelagem Matemática e Computacional.

Aluno: Bruno Ferreira Rosa

Orientador: Prof. Dr. Marcone Jamilson Freitas Souza (UFOP)

Co-Orientador: Prof. Dr. Sérgio Ricardo de Souza (CEFET-MG)

Belo Horizonte, Dezembro de 2009.

Dedico este trabalho aos meus pais José Lázaro e Isilene, e à minha namorada Fernanda, que estão ao meu lado em todos os momentos e foram os principais motivadores deste trabalho. Amo muito vocês!

Agradecimentos

Agradeço a Deus pela vida, por me guardar nas muitas viagens necessárias para a execução deste trabalho.

Aos meus pais pelo exemplo de amor, união, honestidade e determinação.

Aos meus irmãos pela amizade e por estarem sempre ao meu lado, principalmente nos momentos de descontração.

Ao grande amor da minha vida Fernanda por ser minha confidente, minha companheira, pela compreensão e pelo seu amor incondicional.

A todos os meus familiares que sempre me apoiaram em minhas escolhas.

Aos professores Marcone e Sérgio por me aceitarem, por acreditarem na minha capacidade, pelas essenciais orientações, pela paciência e disposição em me ajudar.

Aos professores e funcionários do Programa de Pós-Graduação do CEFET-MG pela oportunidade de buscar e aprimorar meus conhecimentos.

Aos amigos do MMC que estiveram presentes na execução deste trabalho. A amizade e colaboração de vocês foram fundamentais.

Aos colegas que fiz na “estrada” pelas caronas, pela confiança depositada em mim e por compartilharem suas histórias de vida.

Ao CEFET-MG pelo apoio e incentivo.

Enfim, os meus sinceros agradecimentos a todos aqueles que colaboraram, mesmo que indiretamente, para a realização deste trabalho.

Resumo

Este trabalho trata do problema de sequenciamento de tarefas em uma máquina. No problema abordado, os tempos de preparação da máquina são dependentes da sequência de produção e cada tarefa está associada a um tempo de processamento e uma janela de tempo, dentro da qual ela deve ser preferencialmente concluída. O objetivo é minimizar a soma ponderada dos atrasos e das antecipações na execução de tais tarefas. Em termos práticos, as penalidades por antecipação são devidas aos custos gerados pela necessidade de estocagem, enquanto as penalidades por atraso são consequências de multas contratuais. São propostas duas novas formulações de programação matemática para representar o problema, sendo a primeira delas um aperfeiçoamento de uma formulação da literatura, e, a outra, uma formulação indexada no tempo. Propôs-se, também, um algoritmo heurístico para determinar a melhor sequência de produção. Tal algoritmo é composto de duas fases. Na primeira, gera-se uma solução com base na metaheurística GRASP, no Princípio da Otimalidade Próxima e na Descida em Vizinhança Variável. Na segunda, faz-se o pós-refinamento da solução proveniente da fase anterior por meio de outra Descida em Vizinhança Variável. Com o objetivo de reduzir o custo computacional, adicionalmente é proposta uma estratégia de redução do espaço de busca. Devido a particularidade dos problemas-teste da única base de dados existente na literatura referente ao problema tratado, foram geradas outras duas novas bases de dados, uma das quais mais genérica. Experimentos computacionais mostram que a formulação indexada no tempo possibilita resolver problemas maiores e com maior eficiência, quando comparada com as formulações de programação matemática da literatura. O algoritmo heurístico proposto mostrou-se competitivo com os demais algoritmos existentes na literatura, tendo o baixo custo computacional como o seu grande aliado. Desta forma, foi possível resolver, em tempo computacional viável, problemas de dimensões maiores que os resolvidos até então.

Palavras-Chave: Sequenciamento em uma máquina, Formulação indexada no tempo, GRASP, Princípio da Otimalidade Próxima, Descida em Vizinhança Variável.

Abstract

This work deals with the problem of job scheduling in a single-machine. In the considered problem, the machine's setup time is dependent on production sequence and each job has a processing time and a time window within which it should preferably be completed. The objective is to minimize the weighted sum of the tardiness and earliness in the execution of such jobs. In practical terms, the advance penalties are due to the costs incurred by the need for storage, whereas late-payment penalties are the contractual penalties result. It was proposed two new mathematical programming formulations to problem representation, the first being an improved literature formulation, and the other a time-indexed formulation. It was also proposed an heuristic algorithm to determine the best production sequence. This algorithm is composed of two phases. At first, it generates a solution based on GRASP, Proximate Optimality Principle and Variable Neighborhood Descent. In the second, it is the solution post-refinement from the previous phase through another Variable Neighborhood Descent. In order to reduce the computational cost is further proposed a strategy to reduce the search space. Due to the instances particularity exist only one database in the literature for the problem approached, we generated two other new databases, one of them more generic. Computational experiments show that time-indexed formulation allows solving larger instances and more efficiently when compared with others mathematical programming formulations. The heuristic algorithm was shown to be competitive with other literature algorithms, and having low computational cost as its ally. Thus, it was possible to resolve in computing time feasible, greater dimensions instances than those solved until this date.

Keyword: Single-machine scheduling, Time-indexed formulation, GRASP, Proximate Optimality Principle, Variable Neighborhood Descent.

Sumário

Lista de Abreviaturas e Siglas	viii
Lista de Figuras	x
Lista de Algoritmos	xi
Lista de Tabelas	xii
1 Introdução	1
1.1 Objetivos	3
1.1.1 Objetivo Geral	3
1.1.2 Objetivos Específicos	3
1.2 Motivação	3
1.3 Estrutura do Trabalho	5
2 Problema Abordado e Métodos de Resolução	6
2.1 Características do PSUMAA-JE-TPD	6
2.2 Métodos de Resolução	8
2.2.1 Métodos Exatos	8
2.2.2 Métodos Heurísticos	10
2.2.3 Metaheurísticas	13
2.2.4 Estratégias de Intensificação	23
2.3 Trabalhos relacionados	24
2.3.1 PSUMAA com Datas de Entrega Comuns	24
2.3.2 PSUMAA com Datas de Entrega Distintas	26
2.3.3 PSUMAA com Janelas de Entrega Distintas	27
3 Formulações de Programação Matemática	29
3.1 Modelo de Gomes Júnior <i>et al.</i> (2007) - MPLIM-G	29
3.2 Primeiro Modelo Proposto - MPLIM-BG	30
3.3 Segundo Modelo Proposto - MPLIM-IT	31
4 Algoritmo Proposto	33
4.1 GPV	33
4.2 Representação de uma Solução	33
4.3 Vizinhança de uma Solução	33
4.4 Função de Avaliação	35
4.4.1 Algoritmo de Determinação das Datas Ótimas de Início de Proces- samento das Tarefas	35
4.5 Construção de uma Solução	38

4.6	VND ₁	40
4.7	VND ₂	41
4.8	Redução do Espaço de Busca - GPV-REB	41
5	Resultados Computacionais	44
5.1	Problemas-teste	44
5.2	Resultados Obtidos com as Formulações de Programação Matemática . . .	46
5.2.1	Primeira Bateria de Testes - BDS	47
5.2.2	Segunda Bateria de Testes - BDNS	47
5.2.3	Terceira Bateria de Testes - BDDT	48
5.3	Resultados Obtidos com os Algoritmos GPV e GPV-REB	49
5.3.1	Definição dos Parâmetros	50
5.3.2	Primeira Bateria de Testes - BDS	54
5.3.3	Segunda Bateria de Testes - BDNS	56
5.3.4	Terceira Bateria de Testes - BDDT	58
5.4	Teste de Probabilidade Empírica	60
6	Conclusão	62
	Publicações	64
	Referências Bibliográficas	65

Lista de Abreviaturas e Siglas

ADDOIPT Algoritmo de Determinação das Datas Ótimas de Início de Processamento das Tarefas

AE Algoritmos Evolutivos

AG Algoritmos Genéticos

BB *Branch-and-Bound*

BL Busca Local

EDD *Earliest Due Date* - data de entrega mais cedo

FIFO *First-In-First-Out* - o primeiro que entrar é o primeiro a sair

GRASP *Greedy Randomized Adaptive Search Procedure* - Procedimento de busca adaptativa gulosa e randomizada

HC Heurísticas Construtivas

ILS *Iterated Local Search*

JIT *Just in Time*

LRC Lista Restrita de Candidatos

PD Programação Dinâmica

PLIM Programação Linear Inteira Mista

POP Princípio da Otimalidade Próxima

PR *Path Relinking* - Reconexão por Caminhos

PSUMAA Problema de sequenciamento em uma máquina com penalidades por antecipação e atraso na produção

PSUMAA-JE-TPD PSUMAA com janelas de entrega distintas, tempos de *setup* dependentes da sequência de produção e tempos ociosos entre as execuções de tarefas consecutivas permitidos

SA *Simulated Annealing*

SS *Scatter Search*

TS *Tabu Search* - Busca Tabu

VND *Variable Neighborhood Descent* - Busca em Vizinhaça Variável

Lista de Figuras

1.1	Aplicação em uma indústria siderúrgica.	4
2.1	Exemplo de dois posicionamentos distintos para as tarefas de uma mesma sequência de produção.	7
2.2	Exemplo de vizinhança de uma solução.	11
2.3	Comportamento das heurísticas de refinamento.	14
2.4	Estrutura básica de um Algoritmo Evolutivo.	22
4.1	Exemplo de situação em que a estratégia de redução do espaço de busca é utilizada.	42
5.1	Resultados dos testes para definição de γ	51
5.2	Resultados dos testes para definição de MRD_{max}	51
5.3	Resultados dos testes para definição de VND_{max}	52
5.4	Resultados dos testes para definição de $GRASP_{max}$	53
5.5	Teste de probabilidade empírica.	61

Lista de Algoritmos

1	Método da descida.	12
2	Método primeira melhora.	13
3	Método descida randômica.	14
4	Busca em Vizinhança Variável.	16
5	GRASP.	16
6	Fase de construção de um algoritmo GRASP.	17
7	Busca Tabu.	19
8	<i>Simulated Annealing</i>	20
9	<i>Iterated Local Search</i>	21
10	Método GPV	34
11	Algoritmo para determinar as datas ótimas de início de processamento das tarefas.	37
12	Construção de uma solução com POP.	38

Lista de Tabelas

4.1	Resultados apresentados pelas funções adaptativas gulosas testadas.	39
5.1	Comparação MPLIM-G $\times$ MPLIM-BG aplicados à BDS.	47
5.2	Comparação MPLIM-G $\times$ MPLIM-BG aplicados à BDNS.	48
5.3	Comparação MPLIM-G $\times$ MPLIM-BG $\times$ MPLIM-IT aplicados à BDDT. . .	49
5.4	Parâmetros utilizados.	53
5.5	Comparação GPV $\times$ Gomes Júnior (2007) $\times$ Penna (2009) $\times$ Ribeiro (2009) aplicados à BDS.	55
5.6	Comparação GPV $\times$ GPV-REB aplicados à BDS.	56
5.7	Comparação GPV $\times$ GPV-REB $\times$ Ribeiro (2009) aplicados à BDNS. . . .	57
5.8	Comparação GPV $\times$ GPV-REB $\times$ Ribeiro (2009) aplicados à BDDT. . . .	59
5.9	Resultados do MPLIM-IT associado ao GPV.	60

Capítulo 1

Introdução

O surgimento do sistema de administração *Just in Time* (JIT), que ocorreu nos meados da década de 70, evidenciou a importância de um planejamento criterioso das atividades produtivas. Visando a redução dos custos provenientes de processos produtivos, a filosofia JIT desencoraja, além dos atrasos, também as antecipações das tarefas. Há antecipação de uma tarefa quando ela é concluída antes da data desejada para sua entrega e há atraso quando ela é concluída após tal data. Deste modo, ambas as situações acarretam penalidades.

De acordo com Liaw (1999), concluir uma tarefa com atraso pode resultar em multas contratuais, perda de credibilidade da empresa e redução de vendas. Do mesmo modo, concluir uma tarefa antecipadamente pode resultar em custos financeiros extras pela necessidade de disponibilização antecipada de capital, necessidade de espaço para armazenamento ou necessidade de outros recursos para manter e gerenciar o estoque (França Filho, 2007).

O problema de sequenciamento em uma máquina com penalidades por antecipação e atraso da produção, de agora em diante denotado por PSUMAA, consiste em sequenciar e determinar o momento em que as tarefas devem ser executadas em uma máquina, com o objetivo de minimizar a soma ponderada das antecipações e dos atrasos na produção de tais tarefas. Tal problema, segundo Baker & Scudder (1990), reflete melhor ambientes de produção administrados conforme a filosofia JIT.

Outros objetivos relacionados a essa filosofia de produção também são encontrados na literatura, entre eles: minimizar a soma das penalidades por antecipações e atrasos juntamente com os tempos ociosos gerados durante a produção das tarefas (Sourd, 2005), minimizar a soma da maior antecipação com o maior atraso ocorridos na sequência de execução das tarefas (Tavakkoli-Moghaddam *et al.*, 2006) e minimizar a soma dos quadrados das antecipações e dos atrasos ocorridos no decorrer do processo de produção (Weng & Ventura, 1996; Mondal, 2002).

Com relação às datas de entrega das tarefas, o PSUMAA pode ser dividido em três variantes:

- (i) datas de entrega comuns (*common due date*): todas as tarefas devem ser, preferencialmente, concluídas em uma única data pré-determinada;
- (ii) datas de entrega distintas (*distinct due dates*): existe uma data de entrega específica associada a cada tarefa, na qual a tarefa deve ser preferencialmente concluída;

- (iii) janelas de entrega distintas (*distinct due windows*): há um determinado período de tempo associado a cada tarefa, dentro do qual a tarefa deve ser preferencialmente concluída.

O caso de janelas de entrega distintas ocorre quando existem tolerâncias em torno das datas desejadas para a entrega de cada tarefa (Koulamas, 1996). Estas tolerâncias estão relacionadas às características individuais das tarefas e influenciam nos tamanhos das janelas de entrega. As tarefas concluídas dentro de suas respectivas janelas de entrega não incorrem em nenhuma penalidade. Já aquelas concluídas fora de suas janelas de entrega, são penalizadas. A produção de bens perecíveis é um exemplo relacionado. Assuma que um fabricante de produtos químicos combina uma certa substância A , que deteriora rapidamente, com uma segunda substância B para produzir um produto C . Se A for produzida suficientemente antes de B , ela deteriorará. Por outro lado, se A for produzida muito depois de B , o custo da produção de C será maior.

Nas indústrias em que são produzidos diferentes tipos de produtos e existe uma troca frequente do tipo de tarefa executada em uma máquina, após a conclusão de uma tarefa, geralmente é necessário preparar a máquina antes do início da execução da tarefa seguinte (Allahverdi *et al.*, 1999). Este tempo de preparação, chamado tempo de *setup*, inclui os tempos gastos para trocar as ferramentas, preparar o material, limpar a máquina etc. A maioria dos trabalhos em problemas de sequenciamento assume que os tempos de *setup* são independentes da sequência de produção, isto é, que eles são desprezíveis ou podem ser acrescentados aos tempos de processamento das tarefas (Gupta & Smith, 2006). No entanto, de acordo com Panwalkar *et al.* (1973), *apud* Gupta & Smith (2006), em grande parte das situações práticas, tais tempos são dependentes da sequência de produção. Christoforetti (2002) cita o exemplo de uma fábrica de papel que produz diversos tipos de folhas com diferentes cores, espessuras e texturas e que realiza frequentemente a preparação de certas máquinas para obter os diferentes tipos de produtos. Em problemas como este, os tempos de *setup* variam de acordo com sequência de produção e representam uma parcela de tempo considerável em relação ao tempo total de processamento. Portanto, eles não podem ser desconsiderados.

Considerações sobre a continuidade do funcionamento da máquina também podem ser impostas ao PSUMAA. Parte dos trabalhos da literatura não permite a inserção de tempos ociosos na sequência de produção (Chang, 1999). Conforme Li (1997), existem casos em que o custo por manter a máquina inativa é maior que o preço pago pela antecipação de uma tarefa e, neste caso, vale a pena antecipar a produção. Outra situação em que isto ocorre é quando a capacidade da máquina é inferior à demanda. No entanto, há casos em que vale a pena manter a máquina parada, mesmo que exista uma tarefa disponível para ser processada (Souza *et al.*, 2008).

Este trabalho trata o PSUMAA com janelas de entrega distintas, tempos de preparação dependentes da sequência de produção e são permitidos tempos ociosos entre as execuções de tarefas consecutivas. Para clareza de entendimento e evitar dúvidas, o PSUMAA com tais características será doravante denotado por PSUMAA-JE-TPD. Apesar de este problema ser uma generalização do PSUMAA, ele não tem recebido a atenção merecida na literatura.

1.1 Objetivos

Nesta Seção são apresentados os objetivos deste trabalho.

1.1.1 Objetivo Geral

Este trabalho trata o problema de sequenciamento de tarefas em uma máquina com penalidades por antecipação e atraso da produção na sua forma mais genérica. O objetivo principal é desenvolver um algoritmo heurístico robusto e capaz de encontrar soluções de boa qualidade para tal problema em tempo de execução viável.

1.1.2 Objetivos Específicos

Com a finalidade de alcançar o objetivo geral, os seguintes objetivos específicos foram traçados:

- Fazer um levantamento dos trabalhos da literatura que tratam o problema abordado e problemas relacionados, bem como os métodos utilizados para resolvê-los;
- Gerar novos conjuntos de problemas-teste do problema em estudo;
- Desenvolver um algoritmo heurístico para resolver o problema;
- Desenvolver novos modelos de programação matemática para representar o problema abordado;
- Resolver os problemas-teste de dimensões menores por meio dos modelos de programação matemática desenvolvidos e disponíveis na literatura;
- Resolver os problemas-teste disponíveis por meio do algoritmo heurístico implementado e, em seguida, analisar e discutir os resultados encontrados.
- Divulgar os resultados do trabalho em eventos científicos, bem como em periódicos especializados.

1.2 Motivação

O PSUMAA com janelas de entregas distintas e tempos de *setup* dependentes da sequência de produção possui muitas aplicações práticas em indústrias metalúrgicas, têxteis, de tintas, entre outras. O caso real de uma indústria siderúrgica é tratado em Bustamante (2007). Em tal trabalho, uma máquina é considerada como sendo uma sequência de laminadores e cada tarefa representa a produção de um determinado produto (barra chata, cantoneira, vergalhão etc.). Antes da fabricação de cada produto, é necessário realizar um conjunto de ajustes de mesma natureza na sequência de laminadores. Estes ajustes dependem do produto a ser fabricado e do produto fabricado anteriormente. Como cada ajuste exige um tempo de execução, a soma destes tempos configura o tempo de *setup*. Esta situação é ilustrada na Figura 1.1.

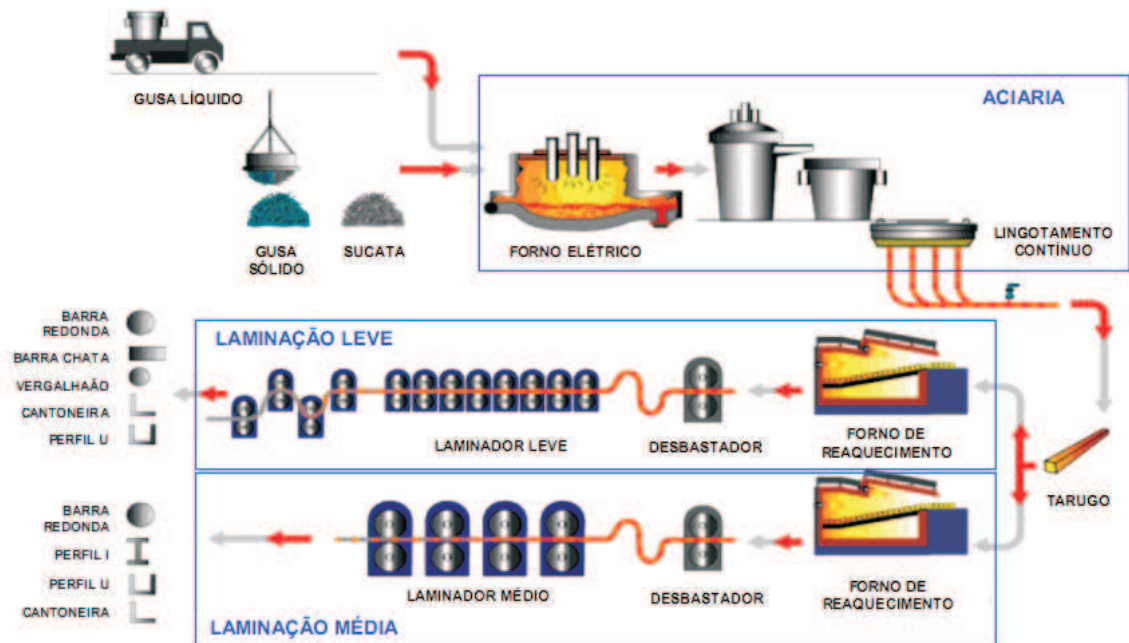


Figura 1.1: Aplicação em uma indústria siderúrgica.

Além do grande número de aplicações, associado ao PSUMAA está a dificuldade de resolvê-lo de forma ótima. Tal dificuldade cresce fatorialmente à medida em que se aumenta a quantidade de tarefas a serem sequenciadas. Esta união entre aplicabilidade e dificuldade de resolução motiva o desenvolvimento de algoritmos eficientes para resolvê-lo.

A utilização de métodos exatos pode demandar um tempo proibitivo para encontrar a solução ótima do PSUMAA, mesmo para problemas de médio porte. Isso justifica a utilização de algoritmos heurísticos, que são procedimentos capazes de encontrar uma “boa solução”, não necessariamente a solução ótima, em tempo computacional viável.

Neste trabalho são utilizados o GRASP (Feo & Resende, 1995) e o Princípio da Otimidade Próxima (Glover & Laguna, 1997) como base para o algoritmo proposto devido aos seguintes motivos: o primeiro por ser uma metaheurística que possui poucos parâmetros para serem ajustados e que se mostrou eficiente na resolução de diversos problemas combinatórios e o segundo, para refinar as soluções parciais e, assim, reduzir o custo computacional da fase de busca local do GRASP.

O desenvolvimento de novas formulações de programação matemática para o problema é motivada pelo fato de existir na literatura apenas um modelo que representa o PSUMAA-JE-TPD. A utilização de um modelo indexado no tempo é motivada pelo trabalho de de Paula (2008), o qual utilizou tal modelagem num problema semelhante e obteve bons resultados. Uma formulação de programação matemática similar à de Gomes Júnior (2007), porém com menor número de variáveis e de restrições, também é proposta na esperança de que ela possibilite resolver problemas de maiores dimensões, quando comparada à formulação de tais autores.

1.3 Estrutura do Trabalho

O restante deste trabalho está organizado como segue.

No Capítulo 2 é feita uma descrição detalhada do PSUMAA-JE-TPD. Um resumo das principais técnicas pertinentes à resolução desse problema também é apresentado. O capítulo é encerrado com a descrição dos trabalhos mais relevantes da literatura em relação ao PSUMAA-JE-TPD.

A formulação de programação matemática proposta por Gomes Júnior (2007) para representar o problema, uma nova formulação baseada nesta última e um modelo de programação matemático indexado no tempo são apresentados no Capítulo 3.

No Capítulo 4 é apresentado o algoritmo heurístico proposto para resolver o PSUMAA-JE-TPD.

Os resultados computacionais obtidos com as formulações de programação matemática propostas e com o algoritmo heurístico desenvolvido são apresentados e discutidos no Capítulo 5.

O Capítulo 6 conclui o trabalho.

Capítulo 2

Problema Abordado e Métodos de Resolução

Neste capítulo é feita uma apresentação das características do problema abordado, seguida pela apresentação dos métodos existentes para resolver problemas similares. Na Seção 2.3 é feita uma revisão da literatura relacionada ao PSUMAA-JE-TPD e problemas correlatos.

2.1 Características do PSUMAA-JE-TPD

O problema de sequenciamento de tarefas em uma máquina tratado neste trabalho (PSUMAA-JE-TPD) possui as seguintes características:

- Uma máquina deve processar um conjunto I de n tarefas;
- Associado a cada tarefa $i \in I$ está:
 - um tempo de processamento P_i ;
 - uma janela de entrega $[E_i, T_i]$ dentro da qual a tarefa i deve ser preferencialmente concluída;
 - um custo α_i por unidade de tempo de antecipação;
 - um custo β_i por unidade de tempo de atraso;
- Há antecipação de uma tarefa $i \in I$ quando seu processamento é concluído antes de E_i ;
- Há atraso de uma tarefa $i \in I$ quando seu processamento é concluído depois de T_i ;
- As tarefas que forem concluídas dentro de suas respectivas janelas de entrega não geram penalidades;
- A máquina executa no máximo uma tarefa por vez e uma vez iniciado o processamento de uma tarefa, não é permitida a sua interrupção;
- Todas as tarefas estão disponíveis para processamento na data 0;

- Entre duas tarefas i e $j \in I$ consecutivas é necessário um tempo S_{ij} de preparação da máquina, chamado tempo de *setup*. Assume-se que o tempo de preparação da máquina para o processamento da primeira tarefa na sequência é igual a 0;
- É permitido tempo ocioso entre a execução de duas tarefas consecutivas.

O objetivo é determinar uma sequência de produção e as datas de início de produção das tarefas de forma a minimizar a soma ponderada das antecipações e atrasos.

É interessante observar que no objetivo do PSUMAA-JE-TPD está embutido dois subproblemas que devem ser resolvidos paralelamente. Um deles é determinar a sequência de produção, isto é, determinar a ordem em que as tarefas serão realizadas. O outro subproblema consiste em programar a data de início de produção de cada tarefa de forma que a soma ponderada das antecipações e atrasos seja a menor possível.

A Figura 2.1 ilustra dois posicionamentos possíveis para as tarefas de uma mesma sequência de produção envolvendo três tarefas. No primeiro posicionamento, as datas de início de processamento das tarefas foram programadas de modo que a tarefa 3 é concluída antecipadamente, acarretando em penalidade por antecipação, e a tarefa 2 dentro de sua janela de entrega, não gerando penalidades. Já no segundo posicionamento, as datas de início de processamento das tarefas foram programadas de forma que a tarefa 2 é concluída com atraso, gerando penalidade, e a tarefa 3 é concluída dentro de sua janela de entrega, não agregando penalidade. Nesta figura, também são ilustrados os tempos de *setup* e de ociosidade de máquina. De acordo com os valores de penalidade incorridos, um destes posicionamentos será o mais adequado para as tarefas desta sequência.

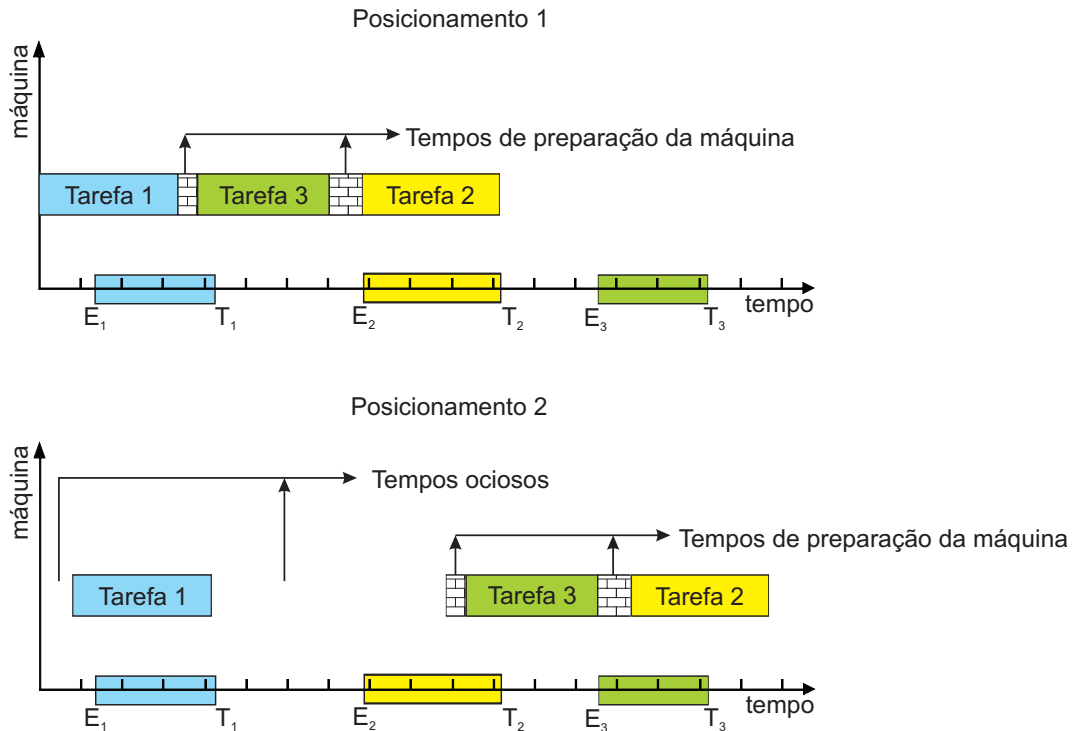


Figura 2.1: Exemplo de dois posicionamentos distintos para as tarefas de uma mesma sequência de produção.

2.2 Métodos de Resolução

O problema de sequenciamento de tarefas é, em seu caso geral, de natureza combinatória. Apesar de possuir domínio finito (para um problema com n tarefas existem $n!$ sequências possíveis) e de simples enumeração, a ideia de testar todas as possíveis sequências é computacionalmente inviável, mesmo para uma quantidade moderada de tarefas. Por exemplo, para o problema com 16 tarefas, um computador que avalia uma sequência em aproximadamente 10^{-6} segundos, gastaria mais de 8 meses para encontrar a melhor sequência.

O PSUMAA com datas de entrega comum é NP-difícil (Baker & Scudder, 1990), ou seja, ainda não é conhecido nenhum algoritmo de complexidade polinomial que o resolva. Como o caso de datas de entrega comum pode ser visto como um caso particular do problema abordado neste projeto, tem-se que o PSUMAA-JE-TPD também é NP-difícil.

Os métodos utilizados para resolver problemas combinatórios, e consequentemente o PSUMAA, podem ser divididos em dois grupos distintos: métodos exatos e métodos heurísticos. Estes métodos são descritos nas seções seguintes.

2.2.1 Métodos Exatos

Os métodos exatos são algoritmos determinísticos, fundamentados em programação matemática e capazes de encontrar a solução ótima de um problema. Uma solução é dita globalmente ótima, ou simplesmente ótima, se, sob um certo critério de otimização, não existir outra solução para o problema com o valor melhor que o dela. Os algoritmos exatos varrem, implicitamente ou não, o conjunto das soluções possíveis para um problema e armazenam a melhor solução encontrada, por isso tem-se a garantia de que a solução retornada é, ou não é, uma solução ótima do problema tratado. Além disso, é possível mensurar o quanto a melhor solução encontrada até um dado instante é inferior à possível solução ótima.

Apesar de serem capazes de encontrar a solução ótima de um problema, devido ao elevado custo computacional, esses métodos são limitados à resolução de problemas de pequeno porte. Como exemplificado anteriormente, para problemas combinatórios, o número de soluções possíveis cresce abruptamente à medida em que se aumenta a dimensão do problema. Deste modo, a simples enumeração das soluções, ainda que implicitamente, se torna inviável para problemas de dimensão mediana.

Os métodos exatos ainda possuem como desvantagens: modelagem mais complexa e nem sempre conseguem encontrar uma solução viável rapidamente. Apesar de suas limitações, os algoritmos exatos são de grande importância na validação de algoritmos não exatos.

A seguir é feita uma breve apresentação de alguns dos métodos exatos já utilizados para resolver o PSUMAA.

Enumeração Total

O método de enumeração total é o mais intuitivo e simples de se implementar entre os algoritmos exatos. Ele consiste em enumerar todas as soluções possíveis e compará-las

entre si. A melhor solução é retornada. Este método é fortemente limitado aos problemas com poucas tarefas.

Branch-and-Bound

O método *Branch-and-Bound* (BB) é um algoritmo de enumeração implícita, ou seja, um método dotado de uma certa “sabedoria”, que retorna a solução ótima de um problema de otimização mesmo enumerando apenas parte do conjunto de soluções possíveis. O BB é baseado no paradigma de divisão e conquista. Tal paradigma consiste em dividir o problema a ser resolvido em partes menores, encontrar soluções para as partes, e então combinar as soluções obtidas em uma solução global (Ziviane, 2004).

De acordo com Gray *et al.* (1997), o método inicia-se considerando o problema original com todas as possíveis soluções, dito problema raiz. Em seguida, são aplicados procedimentos que calculam limitantes superior e inferior para o problema raiz. Se estes limites forem iguais, então uma solução ótima foi encontrada e o procedimento termina. Caso contrário, o problema é dividido em subproblemas menores. Estes subproblemas se tornam filhos do problema raiz. O algoritmo é aplicado recursivamente, gerando assim uma árvore de subproblemas. Quando a solução ótima de um subproblema é encontrada, os demais subproblemas deixam de ser considerados, ou seja, eles são eliminados da árvore de subproblemas. A busca continua até que todos os subproblemas sejam resolvidos ou eliminados, ou até atingir um dado limiar entre a melhor solução encontrada e os limitantes inferiores de todos os subproblemas não resolvidos.

Um limitante superior para um problema de minimização é um valor superior ou igual ao valor ótimo deste problema. O valor da função objetivo de qualquer solução factível, por exemplo, é um limitante superior para o problema. Analogamente, um limitante inferior para este problema é um valor para o qual tem-se a garantia de que a solução ótima é superior ou igual a ele.

Para o cálculo de limites superiores podem ser utilizadas heurísticas, enquanto os limites inferiores são geralmente obtidos por relaxação do problema original. Relaxar um problema consiste em considerar apenas parte de suas restrições. Deste modo, tem-se que a solução do problema original também é uma solução para o problema relaxado. Esta técnica parte do princípio de que o problema relaxado é mais fácil de ser resolvido quando comparado com o problema original. Dentre as técnicas de relaxação existentes, a mais utilizada na resolução do PSUMAA é a relaxação lagrangeana. A técnica de relaxação lagrangeana consiste em adicionar um conjunto de restrições à função objetivo do problema original e resolver o problema relaxado iterativamente. Cada restrição é adicionada com um determinado peso, denominados *multiplicadores de Lagrange*.

Programação Dinâmica

A técnica de programação dinâmica (PD) consiste em dividir o problema em subproblemas. O cálculo das soluções dos subproblemas é feito a partir dos subproblemas menores para os maiores e os resultados são armazenados em uma tabela. Deste modo, uma vez que um subproblema é resolvido, a resposta é armazenada em uma tabela e nunca mais é recalculada. De acordo com Ziviane (2004), a PD é indicada quando a divisão de um problema de tamanho n resulta em n subproblemas de tamanho $n - 1$ facilmente resolvíveis.

2.2.2 Métodos Heurísticos

Como dito anteriormente, o PSUMAA pertence ao conjunto dos problemas de otimização combinatória, para os quais encontrar uma solução ótima de forma exata é uma tarefa difícil. Mesmo com a utilização de estratégias que reduzem o espaço de busca (ver seção 2.2.1), haverá sempre uma quantidade de tarefas a partir da qual a utilização de métodos exatos necessita de um tempo de processamento inviável. Além disso, na maioria das situações práticas, é mais interessante encontrar rapidamente uma “boa solução”, ao invés de perder tempo à procura da solução ótima. Por isso, os métodos heurísticos têm sido muito utilizados na resolução de problemas de natureza combinatória.

Souza (2008) define heurística como sendo uma técnica inspirada em processos intuitivos que procura uma boa solução a um custo computacional aceitável, sem, no entanto, estar capacitada a garantir sua otimalidade, bem como garantir quão próximo está da solução ótima.

Segundo Weise (2008), uma heurística é uma parte de um algoritmo de otimização que usa informações conhecidas para ajudar a decidir qual solução candidata será a próxima a ser testada ou qual a próxima solução a ser produzida. Normalmente, as heurísticas são dependentes das classes de problemas.

Ainda de acordo com Weise (2008), normalmente empregam-se heurísticas em algoritmos determinísticos para definir a ordem de processo das soluções candidatas. Já algoritmos probabilísticos podem considerar somente os elementos do espaço de busca que foram selecionados por uma heurística. Isto se deve ao fato de que, apesar de não garantirem a otimalidade da solução obtida, os métodos heurísticos são de fácil implementação e podem gerar boas soluções rapidamente.

Deste modo, uma heurística pode ser vista como uma técnica que se utiliza de informações particulares dos problemas para encontrar boas soluções a baixo custo computacional.

As heurísticas podem ser divididas em duas classes diferentes: heurísticas construtivas e heurísticas de refinamento. As seções seguintes são destinadas a descrever tais classes.

Heurísticas Construtivas

Heurísticas construtivas (HC) são algoritmos que têm por objetivo a construção de uma solução de forma incremental. A cada passo um elemento é inserido na solução parcial de acordo com uma função de avaliação, que por sua vez depende do problema em questão, até gerar uma solução completa. Nas heurísticas clássicas, os elementos candidatos são geralmente ordenados segundo uma função gulosa, que estima o benefício da inserção de cada elemento, e somente o “melhor” elemento é inserido a cada passo (Souza, 2008). As HC gulosas (heurísticas que utilizam uma função gulosa) são determinísticas e muito utilizadas para gerar soluções iniciais que, possivelmente, serão melhoradas por outros procedimentos heurísticos, ou seja, as soluções gulosas podem ser utilizadas como um limiar superior da solução.

Uma heurística construtiva frequentemente encontrada na literatura para o PSUMAA com datas de entrega distintas é a EDD (*Earliest Due Date* - data de entrega mais cedo). A EDD avalia uma tarefa de acordo com a data de início da sua respectiva janela de entrega. A cada passo, a tarefa com menor data de entrega e ainda não sequenciada é inserida no final da subsequência corrente. O algoritmo é interrompido quando não houver mais tarefa a ser sequenciada.

Outra maneira de se construir uma solução é de forma aleatória, ou seja, em cada passo, o elemento que será inserido na solução parcial é escolhido aleatoriamente dentre os elementos que ainda estão fora dela. Este método é bastante utilizado devido sua facilidade de implementação e por gerar soluções diferentes, isto é, não é determinístico. Porém, testes empíricos mostram que as soluções obtidas geralmente são de baixa qualidade e, por isso, necessitam de um esforço maior na fase de refinamento (Souza, 2008).

Heurísticas de Refinamento

As heurísticas de refinamento são algoritmos que têm por objetivo melhorar uma dada solução inicial por meio de modificações em suas componentes. Uma solução inicial pode ser obtida, por exemplo, através de uma HC. As heurísticas pertencentes a esta classe também são conhecidas como técnicas de busca local (BL) e são baseadas na noção de vizinhança.

Seja S o conjunto de todas as soluções possíveis para um determinado problema de otimização e $f : S \rightarrow \mathbb{R}$ a função objetivo a ser minimizada¹. Uma vizinhança é qualquer função $V : S \rightarrow P(S)$, onde $P(S)$ representa o conjunto das partes de S . Isto é, vizinhança é uma função que leva cada solução $s \in S$ em um subconjunto $V(s) \subseteq S$. Um elemento $s' \in V(s)$ é dito um vizinho de s . Todo vizinho $s' \in V(s)$ é obtido de s por meio de uma operação m , denominada movimento. Esta operação é representada por $s' \leftarrow s \oplus m$.

A Figura 2.2 ilustra a noção de vizinhança. Nela, $V(s) = \{s^1, s^2, s^3, s^4, s^5, s^6\}$ representa a vizinhança da solução $s \in S$. Os movimentos que levam s a cada um dos seus vizinhos são representados pelas setas.

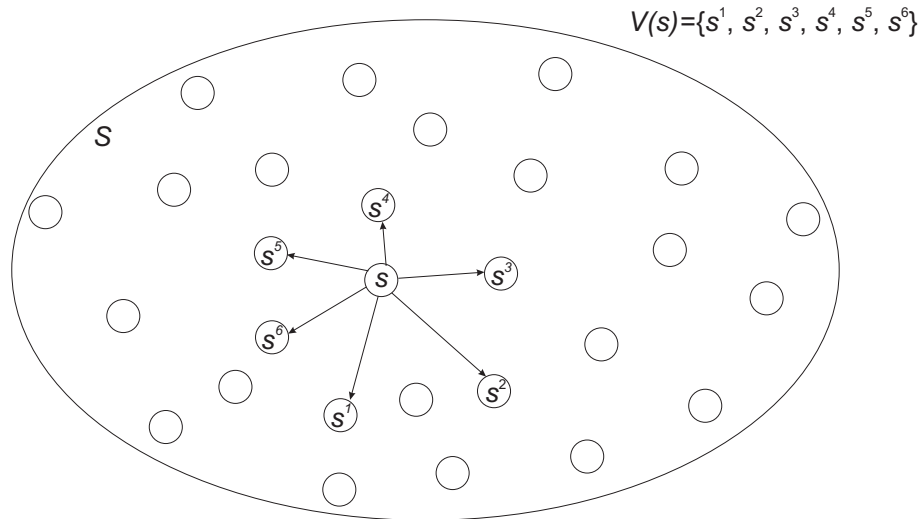


Figura 2.2: Exemplo de vizinhança de uma solução.

O movimento de troca da ordem de processamento de duas tarefas é, por exemplo, um movimento que define uma possível vizinhança para o PSUMAA. Segundo este movimento, a vizinhança de uma solução $s \in S$ é o conjunto das soluções $s' \in S$ obtidas de s por meio da troca da ordem de processamento entre duas tarefas.

¹Caso o problema seja de maximização, ele pode ser facilmente transformado em um problema equivalente de minimização.

Dada uma vizinhança, as técnicas de BL clássicas consistem em caminhar iterativamente de vizinho para vizinho, tentando melhorar a solução construída. Claramente, uma BL é sensível a vizinhança adotada. Deste modo, a escolha de uma vizinhança que permite chegar em qualquer solução do espaço de busca, a partir de qualquer outra solução deste espaço e em um número finito de iterações, é de fundamental importância para o sucesso da técnica.

Outros conceitos importantes surgem com a noção de vizinhança.

Sejam f , V e S como definidos anteriormente. Dizemos que uma solução $s \in S$ é um mínimo local com relação à vizinhança V se $f(s) \leq f(s')$ para todo $s' \in V(s)$. Se uma solução $s^* \in S$ é tal que $f(s^*) \leq f(s')$, para todo $s' \in S$, s^* é dito um mínimo global.

As heurísticas de refinamento clássicas são: método da descida, método de primeira melhora e método de descida randômica.

Dadas uma vizinhança e uma solução inicial, iterativamente, o método da descida analisa todos os vizinhos da solução corrente e o melhor vizinho é selecionado. Se o melhor vizinho representar uma solução melhor que a corrente, ele passa a ser a solução corrente. Caso contrário, o método pára. Portanto a solução retornada é um ótimo local em relação a vizinhança adotada.

O método da descida é descrito no Algoritmo 1, no qual f é a função a ser minimizada, V é uma vizinhança e s uma solução inicial.

procedimento *Descida*($f(\cdot)$, $V(\cdot)$, s)

- 1 $V \leftarrow \{s' \in V(s) \mid f(s') < f(s)\};$
- 2 **enquanto** ($|V| > 0$) **faça**
- 3 Selecione $s' \in V$, onde $s' = \arg \min\{f(s) \mid s \in V\};$
- 4 $s \leftarrow s';$
- 5 $V \leftarrow \{s' \in V(s) \mid f(s') < f(s)\};$
- 6 **fim-enquanto**;
- 7 Retorne s ;

fim *Descida*;

Algoritmo 1: Método da descida.

Como o método da descida analisa toda a vizinhança em cada iteração, ele demanda muito tempo computacional para resolver problemas de maior dimensão. Isto torna necessário técnicas alternativas para estes casos.

O método de primeira melhora (*first improvement method*) é uma variante do método da descida e evita a busca exaustiva pelo melhor vizinho. Neste método, a exploração da vizinhança é interrompida quando um vizinho melhor é encontrado. Quando isto acontece, a solução corrente é atualizada e a exploração da nova vizinhança é inicializada. Deste modo, toda a vizinhança é explorada apenas no pior caso, ou seja, quando somente o último vizinho é de melhora ou todos os vizinhos não são de melhora. O método é interrompido quando um mínimo local é encontrado. É desejável que a ordem de exploração das soluções vizinhas seja alterada a cada passo. Do contrário, privilegia-se apenas um caminho determinístico no espaço de soluções (Souza, 2008).

O Algoritmo 2 representa o método primeiro de melhora. A função a ser minimizada f , a vizinhança V e a solução inicial s são os dados de entrada do método.

```

procedimento Primeira – Melhora( $f(\cdot)$ ,  $V(\cdot)$ ,  $s$ )
1 continua  $\leftarrow$  nao; /* Verifica se a busca na vizinhança corrente pode ser
   interrompida */
2 enquanto (continua = nao) faça
3   continua  $\leftarrow$  sim;
4   para ( $s' \in V(s)$ ) faça
5     se ( $f(s') < f(s)$ ) então
6        $s \leftarrow s'$ ;
7       continua  $\leftarrow$  nao;
8     Sai do para; /* A exploração da vizinhança é interrompida para
       explorar a vizinhança da nova solução corrente */
9   fim-se;
10  fim-para;
11 fim-enquanto;
12 Retorne  $s$ ;
fim Primeira – Melhora;

```

Algoritmo 2: Método primeira melhora.

Uma outra alternativa para se evitar a busca exaustiva de toda a vizinhança é o método de descida randômica. Este método também é uma variante do método da descida. Iterativamente, um vizinho é escolhido de forma aleatória e analisado. Se este vizinho for de melhora, ele passa a ser a solução corrente. Caso contrário, a solução corrente permanece a mesma e passa-se para a próxima iteração. O algoritmo é interrompido após um número pré-determinado de iterações sem melhora na solução corrente.

A solução retornada pelo método randômico de descida não é necessariamente um ótimo local, pois nesta técnica nem todos os vizinhos da solução corrente são explorados.

O refinamento de uma solução s por meio do método randômico de descida, num problema de minimização de uma função f , é representado pelo Algoritmo 3. A vizinhança V e o número máximo de iterações sem melhora *IterMax* também são dados de entrada do método.

2.2.3 Metaheurísticas

Como visto na Seção 2.2.2, as heurísticas de refinamento clássicas são fortemente dependentes da solução inicial e sempre convergem para o primeiro ótimo local em relação à vizinhança adotada. Isto é um grande problema, pois um ótimo local em relação a uma determinada vizinhança pode não ser o ótimo global do problema.

A Figura 2.3 ilustra o comportamento das heurísticas de refinamento para o problema de minimização de uma função f . Observe que os ótimos locais encontrados a partir de s_1 e s_3 não correspondem ao ótimo global, ao contrário do ótimo local relativo à s_2 .

Como o objetivo é encontrar soluções o mais próximo da ótima possível e em tempo reduzido, torna-se necessário a utilização de algoritmos dotados de certa “inteligência”, que sejam capazes de encontrar ótimos locais, mas escapar destes em busca de um ótimo global. Tais algoritmos são comumente chamados de metaheurísticas.

```

procedimento Descida – Randomica( $f(\cdot)$ ,  $V(\cdot)$ ,  $IterMax$ ,  $s$ )
1   $Iter \leftarrow 0$  ;                               /* Contador de iterações sem melhora */
2  enquanto ( $Iter < IterMax$ ) faça
3       $Iter \leftarrow Iter + 1$ ;
4      Selecione aleatoriamente  $s' \in V(s)$ ;
5      se ( $f(s') < f(s)$ ) então
6           $Iter \leftarrow 0$  ;
7           $s \leftarrow s'$  ;
8      fim-se;
9  fim-enquanto;
10 Retorne  $s$ ;
fim Descida – Randomica;

```

Algoritmo 3: Método descida randômica.

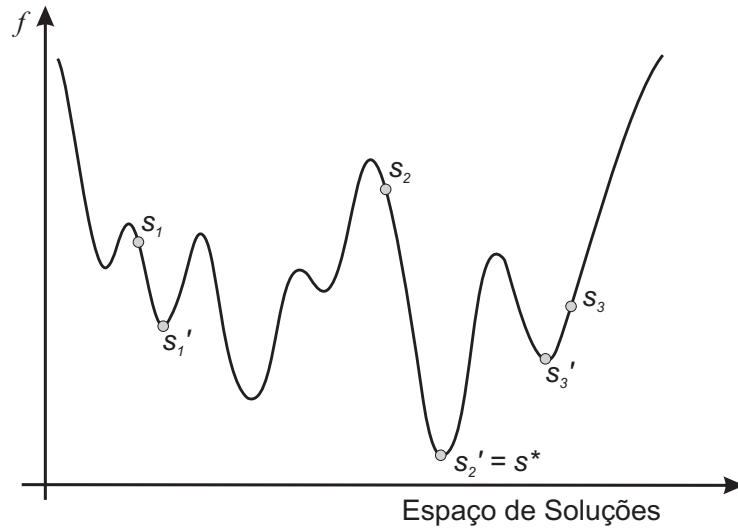


Figura 2.3: Comportamento das heurísticas de refinamento.

De acordo com Weise (2008), uma metaheurística é um método heurístico para resolver uma classe mais geral de problemas. Ela combina funções objetivo ou heurísticas de um modo abstrato e esperançosamente eficiente, normalmente sem utilizar conhecimentos mais profundos do problema, isto é, tratando-o como uma “caixa-preta”.

Para Ribeiro (1996), as metaheurísticas são procedimentos destinados a encontrar uma boa solução, eventualmente a ótima, consistindo na aplicação, em cada passo, de uma heurística subordinada, a qual tem que ser modelada para cada problema específico.

Portanto, neste trabalho considera-se que uma metaheurística é um algoritmo de caráter geral, ao contrário das heurísticas convencionais, com capacidade de escapar de mínimos locais e encontrar o mínimo global de um problema.

Uma metaheurística se diferencia de outra basicamente pela técnica utilizada para fugir de um ótimo local. A forma com que exploram o espaço de busca divide as metaheurísticas em duas categorias distintas: uma baseada em busca local e a outra baseada em busca populacional.

As metaheurísticas fundamentadas em busca local exploram o espaço de busca ele-

mento a elemento, utilizando a noção de vizinhança, a qual varia de acordo com problema abordado (ver Seção 2.2.2). Já os métodos baseados em busca populacional partem de um conjunto de soluções e aplicam, nestas soluções, operadores que visam a melhoria de todo o conjunto.

A seguir serão apresentadas algumas metaheurísticas clássicas.

Busca em Vizinhança Variável

Proposto por Mladenović & Hansen (1997), a Busca em Vizinhança Variável (*Variable Neighborhood Descent*, VND) é um método de busca local que consiste em explorar o espaço de busca por meio de trocas sistemáticas de vizinhanças.

Conforme esses autores, o VND é baseado nos três princípios básicos apresentados a seguir, sendo que o último deles é de natureza empírica e sugere que um ótimo local fornece informações importantes sobre o ótimo global.

- Um ótimo local em relação à uma dada vizinhança pode não ser um ótimo local relativo à outra vizinhança;
- Um ótimo global é um ótimo local em relação a todas as possíveis vizinhanças;
- Para muitos problemas, ótimos locais com relação a uma ou mais vizinhanças são relativamente próximos.

O método VND clássico recebe como entrada uma solução inicial e um conjunto finito e ordenado de vizinhanças. A cada iteração, uma vizinhança é explorada até que se encontre um ótimo local em relação à ela. Se este ótimo local representar uma melhora na solução corrente, ele é aceito e retorna-se à primeira vizinhança. Caso contrário, passa-se a explorar a vizinhança seguinte. O método pára quando um ótimo local em relação a todas as vizinhanças é encontrado.

O Algoritmo 4 representa o método VND básico. Nele é considerado o problema de minimização de uma função f , dados um conjunto $V = \{V^{(1)}, V^{(2)}, \dots, V^{(r)}\}$, composto de r vizinhanças diferentes, e uma solução inicial s .

Segundo Souza (2008), dependendo do problema abordado, a busca pelo ótimo local (linha 3 do Algoritmo 4) pode ser custosa computacionalmente. Nesta situação, é comum fazer a busca pela primeira solução de melhora (ver Seção 2.2.2). Outra alternativa é considerar a exploração apenas em um certo percentual da vizinhança, sendo este percentual um parâmetro do método. Uma terceira alternativa, bastante utilizada nestas situações, é aplicar o método randômico de descida (também descrito na Seção 2.2.2) em cada vizinhança explorada.

Mais informações sobre o VND podem ser encontradas em Mladenović & Hansen (1997) e Hansen & Mladenović (2003), dentre outros trabalhos.

GRASP

O GRASP (*Greedy Randomized Adaptive Search Procedure* - Procedimento de busca adaptativa gulosa e aleatória) é um método iterativo, proposto por Feo & Resende (1995), constituído basicamente de duas fases: uma fase de construção e outra de busca local. Na fase de construção, uma solução é gerada, elemento a elemento, de forma parcialmente

```

procedimento VND( $f(\cdot)$ ,  $V(\cdot)$ ,  $r$ ,  $s$ )
1  $k \leftarrow 1$ ;                                /* Vizinhaça corrente */
2 enquanto ( $k \leq r$ ) faça
3   Encontre o melhor vizinho  $s' \in V^{(k)}(s)$ ;
4   se ( $f(s') < f(s)$ ) então
5      $s \leftarrow s'$ ;
6      $k \leftarrow 1$ 
7   senão
8      $k \leftarrow k + 1$ 
9   fim-se
10 fim-enquanto;
11 Retorne  $s$ ;
fim VND;

```

Algoritmo 4: Busca em Vizinhaça Variável.

gulosa. Já a fase de busca local consiste em refinar a solução construída na primeira fase. Este procedimento é repetido $GRASP_{max}$ vezes e a melhor solução encontrada é retornada como resultado.

O Algoritmo 5 descreve o procedimento básico do GRASP aplicado ao problema de minimização de uma função f . A função adaptativa gulosa g , a vizinhaça V , $GRASP_{max}$ e γ são os parâmetros do método.

```

procedimento GRASP( $f(\cdot)$ ,  $g(\cdot)$ ,  $V(\cdot)$ ,  $GRASP_{max}$ ,  $\gamma$ )
1  $f^* \leftarrow \infty$ ;                          /* Menor valor de  $f$  obtido */
2 para ( $Iter = 1, 2, \dots, GRASP_{max}$ ) faça
3    $s \leftarrow Construc\tilde{a}o - GRASP(g(\cdot), \gamma)$ ;
4    $BuscaLocal(f(\cdot), V(\cdot), s)$ ;
5   se ( $f(s) < f^*$ ) então
6      $s^* \leftarrow s$ ;
7      $f^* \leftarrow f(s)$ ;
8   fim-se
9 fim-para;
10 Retorne  $s^*$ ;
fim GRASP;

```

Algoritmo 5: GRASP.

Como mencionado na Seção 2.2.2, os procedimentos gulosos são determinísticos e, assim, retornam sempre a mesma solução para um determinado problema. Além disso, a solução obtida geralmente converge rapidamente para um ótimo local. Por outro lado, as construções aleatórias geram soluções diversas e, em geral, de baixa qualidade; ou seja, uma construção aleatória gera uma solução diferente em cada execução e uma busca local nestas soluções normalmente demoram para encontrar um ótimo local.

A fase de construção do GRASP visa conciliar as qualidades de uma construção gulosa com as qualidades de uma construção aleatória, isto é, ser capaz de gerar soluções

diferentes e que convergem rapidamente para um ótimo local. A cada iteração desta fase, os elementos que ainda estão fora da solução são ordenados através de uma função que estima o benefício associado a inclusão de cada elemento na solução parcial. Os melhores elementos, segundo esta função, são inseridos em uma lista restrita de candidatos (LRC). Desta lista, aleatoriamente escolhe-se o elemento que será incrementado à solução parcial naquela iteração.

O tamanho da LRC e o nível de gulosidade e de aleatoriedade da fase de construção do GRASP são determinados por um parâmetro $\gamma \in [0, 1]$. É interessante observar que $\gamma = 0$ implica em uma solução completamente gulosa e $\gamma = 1$ implica em uma solução totalmente aleatória.

O Algoritmo 6 descreve a fase de construção de um algoritmo GRASP. A função adaptativa gulosa g é quem estima o benefício da seleção de cada um dos elementos e determina a LRC em cada iteração.

procedimento *Construcao – GRASP*($g(\cdot)$, γ)

- 1 $s \leftarrow \emptyset$;
- 2 Inicialize o conjunto C de candidatos;
- 3 enquanto ($C \neq \emptyset$) faça
- 4 $g(t_{min}) \leftarrow \min\{g(t) \mid t \in C\}$;
- 5 $g(t_{max}) \leftarrow \max\{g(t) \mid t \in C\}$;
- 6 $LRC \leftarrow \{t \in C \mid g(t) \leq g(t_{min}) + \gamma(g(t_{max}) - g(t_{min}))\}$;
- 7 Aleatoriamente, selecione um elemento $t \in LRC$;
- 8 $s \leftarrow s \cup \{t\}$;
- 9 Atualize o conjunto C de candidatos;
- 10 fim-enquanto;
- 11 Retorne s ;

fim *Construcao – GRASP*;

Algoritmo 6: Fase de construção de um algoritmo GRASP.

Na fase de busca local do GRASP pode-se, por exemplo, utilizar uma das técnicas apresentadas na Seção 2.2.2.

Pode-se notar que um ótimo local é encontrado em cada iteração do GRASP. Como os ótimos locais são encontrados a partir de soluções iniciais distintas, eles podem ser diferentes. O melhor ótimo local encontrado é retornado pelo método como a solução para o problema.

Feo & Resende (1995), Resende & Ribeiro (2003) e Souza (2008) apresentam mais informações sobre o GRASP.

Busca Tabu

Proposta por Glover (1986), a Busca Tabu (*Tabu Search*, TS) é uma metaheurística baseada em busca local.

Dada uma vizinhança e uma solução inicial, a TS consiste basicamente em, iterativamente, explorar a vizinhança atual e mudar da solução corrente para o melhor vizinho encontrado, mesmo que o melhor vizinho não represente uma solução de melhora. Isto é feito até que um critério de parada seja satisfeito.

A estratégia de aceitar soluções de piora é utilizada para escapar das armadilhas dos ótimos locais; porém, esta tática pode fazer com que o algoritmo caia em ciclos, isto é, o algoritmo pode voltar a uma solução já gerada anteriormente. Por exemplo, quando a solução corrente é um ótimo local, seu melhor vizinho não será de melhora e será aceito. Na iteração seguinte, o ótimo local será novamente gerado por ser o melhor vizinho da solução corrente e o algoritmo ficará cíclico.

Para evitar ciclagem, o ideal seria armazenar todas as soluções já geradas pelo método em uma lista, denominada lista tabu. No entanto, isto é computacionalmente inviável. Uma alternativa é armazenar apenas as últimas $|LT|$ soluções geradas. Mas uma lista deste tamanho evitaria ciclos de até $|LT|$ iterações, além de poder ser inviável armazenar $|LT|$ soluções e testar se uma solução está ou não na lista tabu.

Na prática, o mais comum é criar uma lista tabu de movimentos proibidos. A lista tabu clássica contém os movimentos reversos aos últimos $|LT|$ movimentos realizados (onde $|LT|$ é um parâmetro do método) e funciona como uma fila FIFO (*First-In-First-Out*).

Apesar de reduzir o risco de ciclagem, uma lista tabu de movimentos pode ser muito restritiva (impede o retorno a uma solução já gerada anteriormente e também que novas soluções sejam geradas). Para tentar corrigir este problema, usa-se um critério de aspiração, que retira, sob certas condições, o *status* tabu de um movimento. Um exemplo de critério de aspiração é aceitar um movimento, mesmo que tabu, se ele melhorar o valor da função objetivo global. Esse é o chamado critério de aspiração por objetivo, sendo o mais usado em implementações Busca Tabu. Outra possibilidade é realizar o movimento tabu mais antigo se todos os possíveis movimentos forem tabus (aspiração por *default*). Uma relação completa de outros critérios de aspiração usualmente utilizados em Busca Tabu é encontrada em Glover & Laguna (1997).

Geralmente, o método é interrompido após $BTmax$ iterações sem melhora na melhor solução encontrada, sendo $BTmax$ um parâmetro do método.

O Algoritmo 7 descreve a metaheurística TS básica aplicada ao problema de minimização de uma função f . A vizinhança V , a função de aspiração A , o conjunto de soluções vizinhas testadas em cada iteração (SVT), o tamanho da lista tabu $|LT|$, o número máximo de iterações sem melhora $BTmax$ e a solução inicial s são as entradas do método.

Mais informações a respeito da TS podem ser encontradas em Glover (1986), Glover (1996), Glover & Laguna (1997), Gendreau (2003) e Souza (2008), entre outros trabalhos.

Simulated Annealing

O *Simulated Annealing* (SA) é uma técnica de busca local probabilística proposta por Kirkpatrick *et al.* (1983) e fundamentada na termodinâmica, ao simular o processo de recozimento (nome dado à operação de resfriamento de um conjunto de átomos aquecidos).

Para se obter metais mais estáveis, estruturalmente fortes e de menor energia interna, eles são aquecidos a uma temperatura suficientemente alta e depois são lentamente resfriados. Durante o recozimento, o material passa por vários estados possíveis.

Como o SA simula o processo de recozimento dos metais, são feitas as seguintes correspondências com um problema de otimização. As soluções do espaço de busca correspondem aos possíveis estados de um metal. O valor da função objetivo corresponde à energia interna em cada estado. O valor de um ótimo local (possivelmente global) corresponde à energia mínima.

```

procedimento Busca – Tabu( $f(\cdot)$ ,  $V(\cdot)$ ,  $A(\cdot)$ ,  $SVT$ ,  $|LT|$ ,  $BTmax$ ,  $s$ )
1  $s^* \leftarrow s$ ;                                /* Melhor solução obtida */
2  $Iter \leftarrow 0$ ;                            /* Contador do número de iterações sem melhora de  $s^*$  */
3  $LT \leftarrow \emptyset$ ;                        /* Lista tabu */
4 Inicialize a função de aspiração  $A$ ;
5 enquanto ( $Iter \leq BTmax$ ) faça
6    $Iter \leftarrow Iter + 1$ ;
7   Seja  $s' \leftarrow s \oplus m$  o melhor elemento de  $SVT \subset V(s)$  tal que o movimento  $m$  não
   é tabu ( $m \notin LT$ ) ou a condição de aspiração de  $s'$  ( $A(s')$ ) é atendida;
8   Atualize a  $LT$ ;
9    $s \leftarrow s'$ ;
10  se ( $f(s) < f(s^*)$ ) então
11     $s^* \leftarrow s$ ;
12     $Iter \leftarrow 0$ ;
13  fim-se
14  Atualize  $A$ ;
15 fim-enquanto;
16 Retorne  $s^*$ ;
fim Busca – Tabu;

```

Algoritmo 7: Busca Tabu.

O método parte de uma solução inicial qualquer e cada iteração consta basicamente de um *loop*. Em cada iteração deste *loop* interno é gerado, aleatoriamente, um único vizinho s' da solução corrente s . Seja $\Delta = f(s') - f(s)$ e considere que o problema seja de minimização. Se $\Delta < 0$ então a solução vizinha passa a ser a solução corrente. Caso contrário ($\Delta \geq 0$), o vizinho pode ser aceito com uma probabilidade $e^{-\frac{\Delta}{kT}}$, onde k é a constante de Boltzmann e T representa a temperatura corrente.

Inicialmente, a temperatura T assume um valor suficientemente alto T_0 . A cada iteração i do método, a temperatura é diminuída, ou seja, $T_i \leftarrow \alpha \cdot T_{i-1}$, onde $0 < \alpha < 1$. Os valores de T_0 e α são parâmetros do método.

Observe que quanto maior a temperatura, maior a probabilidade de se aceitar uma solução de piora. Portanto, a cada iteração do algoritmo a chance de uma solução de piora ser aceita é reduzida. As soluções de piora são aceitas para que se escape das armadilhas dos ótimos locais.

O método é interrompido quando a temperatura atinge um valor próximo de zero e somente vizinhos de melhora são aceitos, evidenciando o encontro de um ótimo local.

O Algoritmo 8 descreve o método SA básico para o problema de minimização de uma função f . A estrutura de vizinhança V , a razão de resfriamento α , o número de iterações em cada temperatura (SAm_{ax}), a temperatura inicial T_0 , um número real próximo de zero (ϵ) e uma solução inicial s são as entradas do método.

Mais detalhes do método podem ser encontrados em Kirkpatrick *et al.* (1983), Ribeiro (1996), Henderson *et al.* (2003) e Souza (2008), entre outros trabalhos.

```

procedimento  $SA(f(\cdot), V(\cdot), \alpha, SAmax, T_0, \epsilon, s)$ 
1  $s^* \leftarrow s$ ;                                /* Melhor solução obtida */
2  $IterT \leftarrow 0$ ;                            /* Número de iterações na temperatura  $T$  */
3  $T \leftarrow T_0$ ;                             /* Temperatura corrente */
4 enquanto ( $T > \epsilon$ ) faça
5   enquanto ( $IterT < SAmax$ ) faça
6      $IterT \leftarrow IterT + 1$ ;
7     Aleatoriamente, escolha  $s' \in V(s)$ ;
8      $\Delta \leftarrow f(s') - f(s)$ ;
9     se ( $\Delta < 0$ ) então
10       $s \leftarrow s'$ ;
11      se ( $f(s') < f(s^*)$ ) então  $s^* \leftarrow s'$ ;
12    senão
13      Aleatoriamente, escolha  $x \in [0, 1]$ ;
14      se ( $x < e^{\frac{-\Delta}{kT}}$ ) então  $s \leftarrow s'$ ;
15    fim-se
16  fim-enquanto;
17   $T \leftarrow \alpha \cdot T$ ;
18   $IterT \leftarrow 0$ 
19 fim-enquanto;
20 Retorne  $s^*$ ;
fim  $SA$ ;

```

Algoritmo 8: *Simulated Annealing.**Iterated Local Search*

O *Iterated Local Search* (ILS) é um método de busca local que faz uso de perturbações para escapar das armadilhas dos ótimos locais e, possivelmente, encontrar um ótimo global. Tais perturbações devem ser fortes o suficiente para permitir a fuga do ótimo local corrente e que a busca local explore diferentes regiões do espaço de busca. Ao mesmo tempo, elas devem ser suficientemente fracas para preservar características do ótimo local corrente e evitar reinícios aleatórios.

De acordo com Souza (2008), é necessário especificar quatro componentes para a aplicação do ILS:

1. Procedimento *GeraSolucaoInicial*, que gera uma solução inicial. Pode ser utilizado, por exemplo, uma HC (ver Seção 2.2.2);
2. Procedimento *BuscaLocal*, que retorna uma solução possivelmente melhorada. Por exemplo, pode ser utilizada uma heurística de refinamento (ver seção 2.2.2);
3. Procedimento *Perturbacao*, que modifica a solução corrente, guiando ela a uma solução intermediária;
4. Procedimento *CritérioAceitacao*, que decide em qual solução a próxima perturbação será aplicada.

Um critério de aceitação muito utilizado consiste em aceitar um ótimo local s'' somente se ele for melhor que o ótimo local corrente s . O Algoritmo 9 ilustra a aplicação do ILS básico a um problema de minimização de uma função f .

```

procedimento  $ILS(f(.))$ 
1  $s_0 \leftarrow GeraSolucaoInicial();$ 
2  $s \leftarrow BuscaLocal(s_0);$ 
3  $s^* \leftarrow s;$                                 /* Melhor solução obtida */
4 enquanto (critério de parada não é satisfeito) faça
5    $s' \leftarrow Perturbacao(s);$ 
6    $s'' \leftarrow BuscaLocal(s');$ 
7   se ( $f(s'') < f(s^*)$ ) então
8      $s^* \leftarrow s'';$ 
9      $s \leftarrow CriterioAceitacao(s, s'');$ 
10  fim-se
11 fim-enquanto;
12 Retorne  $s^*$ ;
fim  $ILS$ ;

```

Algoritmo 9: *Iterated Local Search.*

Geralmente, utiliza-se um número máximo de iterações sem melhora na melhor solução encontrada como critério de parada do algoritmo.

Mais detalhes do método podem ser encontrados em Souza (2008) e Lourenço *et al.* (2003).

Algoritmos Evolutivos

Algoritmos evolutivos (AE's) são metaheurísticas baseadas em busca populacional e que utilizam mecanismos inspirados na biologia, como mutação, *crossover*, seleção natural e condições de sobrevivência, para refinar, iterativamente, um conjunto de soluções candidatas (Weise, 2008).

Essencialmente, um AE parte de um conjunto de soluções candidatas, denominado população inicial, e, em cada iteração, todos os elementos (indivíduos) da população atual são avaliados pela função de avaliação. Caso o critério de parada seja satisfeito, o método é interrompido e os melhores elementos são listados. Caso contrário, com o auxílio de uma função de aptidão, que pode não depender somente da função de avaliação, os indivíduos são novamente avaliados. A função de aptidão tem o objetivo de auxiliar na seleção dos indivíduos da população que darão origem a novos indivíduos. A geração de novos indivíduos pode ser feita por meio de *crossover* ou mutação. No *crossover*, dois indivíduos pais dão origem a filhos que possuem características de ambos os pais. A mutação consiste em alterar uma característica do indivíduo. Terminada a etapa de geração de novos indivíduos, também com o auxílio da função de aptidão, é determinada a população sobrevivente, isto é, os indivíduos que farão parte da população da próxima iteração.

A Figura 2.4 ilustra a estrutura básica de um AE.

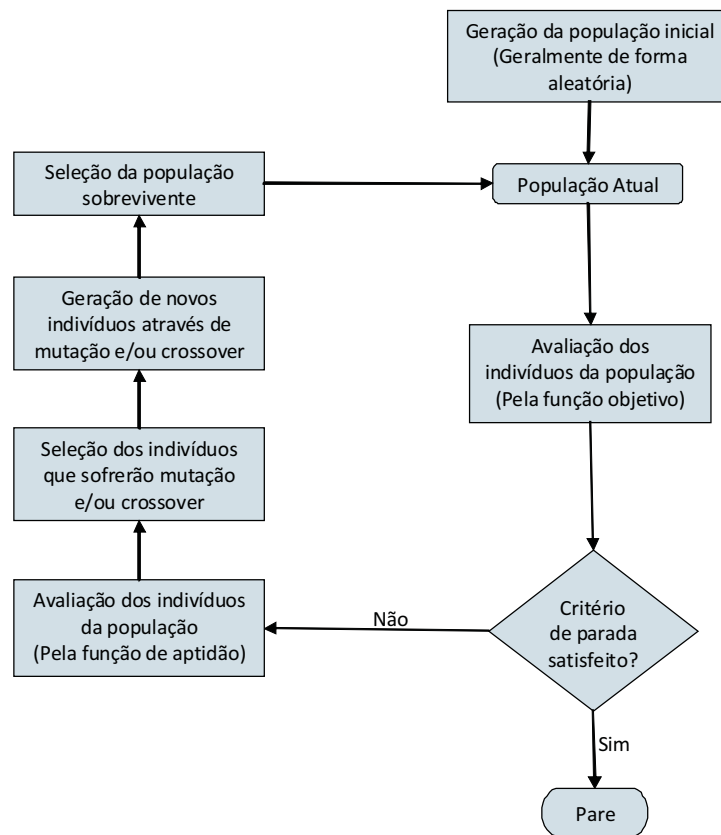


Figura 2.4: Estrutura básica de um Algoritmo Evolutivo.

De acordo com Souza (2008), os critérios de parada mais comuns são: quando um certo número de gerações é atingido, quando não há melhora na melhor solução encontrada após um certo número de iterações ou quando o desvio padrão da população é inferior a um certo valor (situação que evidencia uma homogeneidade da população).

Mais informações sobre os AE's podem ser encontradas, por exemplo, em Weise (2008).

Algoritmos Genéticos

Segundo Weise (2008), os Algoritmos Genéticos (AG's) compõem uma subclasse dos AE's, na qual os elementos do espaço de busca são vetores de componentes binários ou arranjos de tipos elementares.

É muito comum encontrar na literatura trabalhos que, por conveniência, consideram os espaços de soluções dos AG's formados por vetores de números inteiros. Neste caso, a definição de AG coincide com a definição de AE.

De qualquer forma, pode-se dizer que a Figura 2.4 ilustra a estrutura de um Algoritmo Genético básico.

Reeves (2003), Weise (2008) e Souza (2008) apresentam mais detalhes a respeito dos AG's.

Scatter Search

O *Scatter Search* (SS) é um método baseado em busca populacional que apresenta similaridades com os Algoritmos Genéticos; difere no tamanho da população e no uso de estratégias determinísticas, ao invés de probabilísticas, para atingir os objetivos de diversificação e intensificação (Belfiore *et al.*, 2006).

Ainda conforme Belfiore *et al.* (2006), o SS opera sobre um conjunto de soluções, chamado conjunto de referência, combinando estas soluções para criar novas soluções de modo a melhorar a soluções originais. É, portanto, um método evolutivo. O tamanho do conjunto de referência em *scatter search* é pequeno comparado à população nos Algoritmos Genéticos.

Segundo (Glover & Laguna, 2003), o processo do SS é organizado para:

1. Capturar informações não contidas separadamente nas soluções originais;
2. Tirar proveito de métodos de soluções heurísticas auxiliares (para avaliar as combinações produzidas e gerar ativamente novas soluções);
3. Utilizar estratégias específicas no lugar de métodos aleatórios de busca.

Uma abordagem mais detalhada do método é apresentada em Glover & Laguna (2003) e Belfiore *et al.* (2006).

2.2.4 Estratégias de Intensificação

As técnicas de intensificação são métodos que exploram com maior intensidade as soluções próximas de uma dada solução. A seguir serão apresentadas duas destas técnicas.

Princípio da Otimalidade Próxima

O Princípio da Otimalidade Próxima (POP) é uma técnica que refina uma parte da solução com o objetivo de melhorar toda a solução. Este método é baseado na ideia de que boas soluções em um nível estão próximas de boas soluções em um nível adjacente (Glover & Laguna, 1997).

Na prática, o POP é muito utilizado em métodos construtivos. Periodicamente, antes de se inserir um novo elemento na solução em construção, esta é submetida à uma busca local.

Reconexão por Caminhos

A Reconexão por Caminhos (*Path Relinking*, PR) é uma técnica de intensificação proposta em Glover (1996). O método parte de duas soluções dadas, sendo uma chamada de solução base e a outra de solução guia. Iterativamente, um atributo da solução guia é adicionado na solução base. A solução base sofre, então, uma busca local em que se fixam todos os atributos da solução guia que já foram incrementados à solução base. O método pára quando a solução base possui todos os atributos da solução guia.

O objetivo do método é, portanto, encontrar, durante as iterações, uma solução que possua os atributos bons contidos nas duas soluções dadas, isto é, uma solução melhor que as duas soluções iniciais.

O PR pode ser utilizado como um método de pós-otimização ou então para intensificar a busca em torno de ótimos locais. No primeiro caso, as melhores soluções distintas encontradas por algum procedimento heurístico são armazenadas em um conjunto e posteriormente aplica-se a PR em pares de soluções deste conjunto. Tal conjunto é chamado de conjunto elite e seu tamanho é um parâmetro do método. No segundo caso, sempre que um ótimo local é encontrado por um determinado procedimento, aplica-se a PR no par de soluções formado pelo ótimo local e a melhor solução até então encontrada pelo algoritmo.

Mais detalhes desta técnica podem ser encontrados em Glover (1996), Glover & Laguna (2003) e Souza (2008).

2.3 Trabalhos relacionados

Muitos trabalhos são encontrados na literatura com o objetivo de resolver o PSUMAA em suas diversas variantes. Uma breve descrição dos trabalhos mais relevantes são apresentados a seguir. Por questões didáticas, os trabalhos foram separados de acordo com as datas de entrega. Primeiramente são apresentados os trabalhos que tratam do PSUMAA com datas de entrega comum, em seguida os com datas de entrega distintas e, por último, os com janelas de entrega distintas.

2.3.1 PSUMAA com Datas de Entrega Comuns

Esta seção é destinada à apresentação de alguns trabalhos que tratam do PSUMAA com datas de entrega comuns.

Dentre as versões do PSUMAA, a mais estudada é a com datas de entrega comuns e tempos de *setup* independentes da sequência de produção. Para este caso, existem propriedades que podem ser utilizadas nos algoritmos de resolução, de modo a reduzir o espaço de busca e conseqüentemente o tempo computacional. As principais propriedades que são utilizadas por todos os autores são:

Propriedade 1 *Em uma sequência ótima não há tempo ocioso entre duas tarefas consecutivas.*

Propriedade 2 *Em uma sequência ótima a conclusão de uma tarefa coincide com a data de entrega ou o processamento da primeira tarefa inicia no instante zero.*

Propriedade 3 *Em uma sequência ótima, as tarefas antecipadas estão ordenadas de forma não-crescente pela razão entre o tempo de processamento e a penalidade por unidade de tempo de antecipação, enquanto as tarefas atrasadas estão ordenadas de forma não-decrescente pela razão entre o tempo de processamento e a penalidade por unidade de tempo de atraso.*

Se uma sequência satisfaz a Propriedade 3, ela é dita *V-shaped* (Biskup & Feldmann, 2001).

As justificativas para estas propriedades podem ser encontradas em Szwarc (1996). Neste trabalho, também foram estudadas propriedades da otimalidade no caso em que as

penalidades são dependentes das tarefas e as razões entre os tempos de processamento e as penalidades satisfazem certas condições. As propriedades estudadas foram utilizadas em três algoritmos baseados em programação dinâmica e de complexidade pseudo-polinomial, sendo que estes algoritmos refletiam o estado da arte até então. Experimentos em problemas-teste com até 100 tarefas evidenciaram a melhoria proporcionada pela utilização das propriedades apresentadas.

Lee & Kim (1995) resolveram o problema por meio de um algoritmo genético paralelo com representação binária, sendo este uma variante do AG convencional. Uma condição necessária para que um arranjo em série de genes seja *V-shaped* dentro de uma sequência é apresentada. Foram resolvidos problemas-teste com até 100 tarefas.

Alidaee & Dragan (1997) apresentaram um algoritmo de complexidade $O(n \log n)$ que resolve de forma ótima os casos em que as penalidades por antecipação e atraso são iguais para a mesma tarefa e proporcionais aos tempos de processamento das tarefas.

Biskup & Feldmann (2001) propuseram um gerador de problemas-teste. Foram utilizadas duas novas heurísticas para resolver e obter limites superiores para os problemas-teste, com até 1000 tarefas, gerados. As heurísticas se mostraram eficientes, pois encontraram resultados próximos dos ótimos conhecidos de problemas de menor porte. O gerador, junto com os limites superiores, são disponibilizados pelos autores para servirem como referências de trabalhos futuros sobre o PSUMAA.

Feldmann & Biskup (2003) apresentam uma nova e apropriada representação para o problema, baseada nas propriedades deste caso do PSUMAA. Foram utilizadas três metaheurísticas para resolver o problema: um Algoritmo Evolutivo, um *Simulated Annealing* e um *Threshold Accepting*, sendo este último uma variante do SA. Estas metaheurísticas foram aplicadas em problemas-teste propostos por Biskup & Feldmann (2001).

Uma construção gulosa e heurísticas baseadas em Busca Tabu e Algoritmos Genéticos foram propostas por Hino *et al.* (2005). Estratégias híbridas também foram analisadas com o objetivo de melhorar o desempenho destes métodos. Foram resolvidos problemas-teste com até 1000 tarefas.

Belfiore *et al.* (2006) aplicaram a metaheurística *Scatter Search* neste variante do PSUMAA. O algoritmo proposto foi utilizado para resolver todos os problemas-teste propostos por Biskup & Feldmann (2001).

Ying (2008) resolveu o problema por meio de um procedimento *Branch-and-Bound* que faz uso de um algoritmo *Recovering Beam Search* para caminhar na árvore de busca. O autor também utilizou a base de dados proposta por Biskup & Feldmann (2001) para realizar experimentos computacionais.

O PSUMAA com datas de entrega comum pode ser classificado em restrito ou irrestrito. Quando a soma dos tempos de processamento de todas as tarefas é maior que a data de entrega, diz-se que o problema é restrito. Caso contrário, o problema é dito irrestrito.

James (1997) testou diferentes formas de Busca Tabu para resolver o caso restrito do problema, inclusive uma que aplica as propriedades numa nova técnica que reduz o tamanho do espaço de busca de $n!$ para 2^n . Foram resolvidos problemas-teste com até 250 tarefas e a TS que faz uso da nova técnica mostrou-se mais eficiente que as demais.

Mondal & Sen (2001) consideraram que as penalidades por antecipação e atraso são independentes das tarefas. Foi apresentado um algoritmo com espaço de busca em grafo para encontrar a solução ótima de problemas-teste maiores do caso restrito. Ao contrário da programação dinâmica, o algoritmo encontrou soluções ótimas mesmo quando a me-

mória era limitada. O algoritmo se mostrou mais rápido que a programação dinâmica e *depth-first-branch-bound* (*depth-first* é uma heurística utilizada para caminhar na árvore de busca do método BB) e consegue resolver muitos problemas-teste maiores em tempo razoável. Também foram propostos e utilizados novos limites inferiores e superiores para o problema. Os problemas-teste resolvidos são com até 2000 tarefas.

Rabadi *et al.* (2004) consideraram o PSUMAA com datas de entrega comuns, penalidades iguais (o objetivo é minimizar a soma simples das antecipações e dos atrasos) e tempos de *setup* dependentes da sequência de produção. Foi proposto um algoritmo *Branch-and-Bound* para obter soluções ótimas do caso irrestrito. Os autores resolveram problemas-teste com até 25 tarefas em tempo computacional aceitável, o que representou um avanço para a época, já que os algoritmos exatos de então resolviam, na otimalidade, apenas problemas desta classe com até 8 tarefas.

2.3.2 PSUMAA com Datas de Entrega Distintas

Esta seção é destinada à apresentação de alguns trabalhos que tratam do PSUMAA com datas de entrega distintas.

M'Hallah (2007) considerou o problema com as penalidades iguais (o objetivo é minimizar a soma simples das antecipações e dos atrasos) e os tempos de *setup* independentes da sequência de produção, não sendo permitida a existência de tempo ocioso entre a execução de tarefas consecutivas. Foi proposto um método heurístico híbrido que combina heurísticas de busca local (regras de despacho, métodos de descida e SA) e um algoritmo evolutivo baseado em AG. Os efeitos de cada heurística no método são mostrados via experimentos com problemas-teste com até 4000 tarefas.

Chang (1999) tratou do problema com penalidades iguais e tempos de *setup* independentes da sequência de produção, sendo permitida a ociosidade de máquina. O autor utilizou um algoritmo *branch-and-bound* para resolver o problema. Um esquema para calcular diferentes limites inferiores baseados no procedimento de eliminação de sobreposição de uma sequência *Just-In-Time* também foi apresentado. Propriedades e teoremas deste procedimento de eliminação foram demonstradas. O algoritmo foi utilizado para resolver problemas-teste com até 45 tarefas. Ventura & Radhakrishnan (2003) abordaram o mesmo problema, para o qual também propuseram um algoritmo BB. Neste trabalho foram derivadas propriedades de uma sequência ótima. Um procedimento de relaxação lagrangeana que utiliza o algoritmo de subgradiente foi implementado para atingir soluções próximas da otimidade e ajustar os limites inferiores. O algoritmo foi testado em problemas-teste com até 100 tarefas e apresentou uma divergência relativa menor que 3% das soluções ótimas.

Li (1997) estudou o caso com tempos de *setup* independentes da sequência de produção, não sendo permitida a existência de tempos ociosos entre a execução de tarefas consecutivas. Neste trabalho, o problema foi decomposto em dois subproblemas com estruturas mais simples. O limite inferior do problema foi, então, dado pela soma dos limites inferiores dos dois subproblemas. O limite inferior de cada subproblema foi obtido por relaxação lagrangeana. Também foram desenvolvidos dois procedimentos de ajuste multiplicador, de complexidade $O(n \log n)$, para resolver os duais dos subproblemas. Um algoritmo *branch-and-bound* baseado nestes procedimentos foi apresentado e utilizado para resolver problemas-teste com até 50 tarefas, dobrando a dimensão dos problemas resol-

vidos, na otimidade, por algoritmos exatos até aquela data. O propôs, também, uma heurística baseada em busca local e resolveu problemas com até 3000 tarefas. Almeida & Centeno (1998) aplicaram um método heurístico composto, que combina TS, SA e técnicas de descida, em problemas-teste com até 50 tarefas. Liaw (1999), Valente & Alves (2005b) e Valente & Alves (2005a) também trataram este mesmo problema. O primeiro autor utilizou um algoritmo BB baseado em procedimentos que determinam limites inferiores e superiores eficientes. Foram derivadas regras de dominância para ajudar a eliminar nós não promissores da árvore de busca. O desempenho do algoritmo foi testado em problemas-teste com até 50 tarefas. No segundo trabalho foram introduzidas duas novas heurísticas, sendo uma regra de despacho e um procedimento guloso. Também foram utilizadas regras de dominância, visto que elas melhoram os resultados das heurísticas e demandam pouco tempo adicional. Os experimentos computacionais foram realizados em problemas-teste com até 500 tarefas. O terceiro trabalho propõe dois algoritmos BB. Cada BB proposto se utiliza de uma variante do algoritmo *beam-search* para caminhar na árvore de busca, sendo que um deles se utiliza de um *Filtered Beam Search* e o outro de um *Recovering Beam Search*. Os dois métodos foram testados em problemas-teste com até 500 tarefas.

Lee & Choi (1995) trataram do caso com tempos de *setup* independentes da sequência de produção, sendo permitida a ociosidade de máquina. Foi apresentado um algoritmo de complexidade polinomial que determina a data ótima de início de uma tarefa em uma dada sequência de produção. Também aplicou-se ao problema um algoritmo genético que faz uso deste algoritmo de data ótima. Foram resolvidos problemas-teste com até 80 tarefas. Mazzini & Armentano (2001) estudaram este mesmo problema. Neste trabalho, uma heurística construtiva que determina a sequência de produção e simultaneamente insere tempos ociosos foi proposta e utilizada para resolver problemas-teste com até 80 tarefas.

No trabalho de Bustamante (2007) foi estudado o problema com tempos de *setup* dependentes da sequência de produção, sendo permitida a existência de tempos ociosos de máquina. São desenvolvidos dois modelos de programação linear inteira mista, sendo que tais modelos exigem que os tempos de *setup* satisfaçam a desigualdade triangular. O autor também sugeriu alterações nos modelos que permitem resolver o caso com janelas de entrega distintas, mas os testes computacionais foram realizados apenas em problemas com datas de entrega distintas. Foram resolvidos, na otimalidade, problemas com até 10 tarefas por meio do software de otimização GLPK 4.8.

2.3.3 PSUMAA com Janelas de Entrega Distintas

Alguns trabalhos da literatura que tratam o PSUMAA com janelas de entrega distintas são apresentados a seguir.

Wan & Yen (2002) trataram o problema com tempos de *setup* independentes da sequência de produção, permitindo a existência de tempos ociosos de máquina. Primeiramente, foi apresentada uma formulação matemática do problema junto com várias propriedades importantes para sua resolução. Em seguida, foi proposto um procedimento de complexidade polinomial para determinar a data de conclusão ótima de processamento de cada tarefa em uma dada sequência, sendo este procedimento uma extensão do algoritmo proposto por Lee & Choi (1995). Por fim, uma TS, que faz uso do procedimento de datas

ótimas, foi proposta para resolver o problema. Foram realizados testes computacionais em problemas envolvendo até 80 tarefas.

Koulamas (1996) focou no problema com janelas de entrega distintas, porém com uma sensível distinção dos demais trabalhos da literatura em relação à antecipação. Em seu trabalho, há antecipação de uma tarefa quando seu processamento é iniciado antes do início de sua janela de entrega. O restante da literatura considera que há antecipação quando uma tarefa é concluída antes do início de tal janela. O autor também considerou tempos de *setup* independentes da sequência de produção, sendo permitida a ociosidade de máquina. O autor adaptou heurísticas já utilizadas em outros casos do PSUMAA e as aplicou nessa versão. Também foi proposto um algoritmo que insere tempos ociosos, de modo ótimo, em uma dada sequência de produção. Foram realizados experimentos computacionais em problemas-teste com até 200 tarefas.

Gomes Júnior *et al.* (2007), Penna (2009) e Ribeiro (2009) correspondem ao estado da arte em relação ao PSUMAA-JE-TPD. No primeiro trabalho foi proposto um modelo de programação matemática para representar o problema, no qual deixa de ser necessário que o problema satisfaça a desigualdade triangular. Também foi proposto um método heurístico baseado em GRASP, ILS e VND. Para cada sequência de tarefas gerada pela heurística proposta, é utilizado um algoritmo de complexidade polinomial que determina a data ótima de início de processamento das tarefas na sequência dada, sendo este algoritmo uma extensão daquele proposto por Wan & Yen (2002). Foram realizados experimentos computacionais em problemas-teste com os tempos de *setup* simétricos e com até 75 tarefas, utilizando a heurística, e até 12 tarefas, utilizando o modelo matemático. O segundo trabalho propõe uma heurística baseada em GRASP, VND, TS e PR, enquanto no terceiro trabalho é proposto um Algoritmo Genético Adaptativo para resolver o problema. Nos dois últimos trabalhos foram realizados experimentos nos mesmos problemas-teste utilizados por Gomes Júnior *et al.* (2007), sendo mostrada a superioridade dos algoritmos propostos em relação a este último. Ribeiro (2009) também realizou testes em outras duas bases de dados em que os tempos de *setup* não são necessariamente simétricos. Foram resolvidos problemas-teste com até 100 tarefas por meio do AG proposto.

Capítulo 3

Formulações de Programação Matemática

Como foi mencionado na Seção 2.2.1, os métodos exatos são de grande importância por serem capazes de encontrar a solução ótima do problema e auxiliarem na validação dos procedimentos heurísticos. Com esta finalidade, nesta seção são apresentadas três formulações de programação matemática para representar o PSUMAA-JE-TPD. A primeira delas foi proposta por Gomes Júnior (2007), enquanto as outras duas são propostas neste trabalho. A segunda formulação é um aperfeiçoamento da primeira e a última é indexada no tempo.

3.1 Modelo de Gomes Júnior *et al.* (2007) - MPLIM-G

A seguir é reproduzido o modelo de programação linear inteira mista (PLIM) proposto por Gomes Júnior (2007) para representar o PSUMAA-JE-TPD, doravante denotado por MPLIM-G.

Para auxiliar na modelagem, são utilizadas duas tarefas fictícias, 0 (zero) e $n + 1$, que devem ser sequenciadas necessariamente na primeira e na última posição, respectivamente. Admite-se que P_0 e P_{n+1} são iguais a zero e que $S_{0i} = S_{i0} = 0$ e $S_{i,n+1} = S_{n+1,i} = 0$, $\forall i \in I$. A data de início do processamento da tarefa $j \in I \cup \{0, n + 1\}$ e os tempos de antecipação e atraso da tarefa $i \in I$ são representados por s_j , e_i e t_i , respectivamente.

Sejam y_{ij} variáveis que determinam a sequência de produção, sendo $y_{ij} = 1$ se a tarefa j for sequenciada imediatamente após a tarefa i e $y_{ij} = 0$, caso contrário, $\forall i, j \in I \cup \{0, n + 1\}$. Considere, ainda, uma constante M de valor suficientemente grande. O

modelo MPLIM-G é representado pelas equações (3.1)-(3.10):

$$\min \quad Z = \sum_{i=1}^n (\alpha_i e_i + \beta_i t_i) \quad (3.1)$$

$$\text{s.a. } s_i + P_i + y_{ij}(M + S_{ij}) - M \leq s_j \quad \forall i \in I \cup \{0\}, \\ \forall j \in I \cup \{n+1\} \text{ e } i \neq j \quad (3.2)$$

$$\sum_{j=1, j \neq i}^{n+1} y_{ij} = 1 \quad \forall i \in I \cup \{0\} \quad (3.3)$$

$$\sum_{i=0, i \neq j}^n y_{ij} = 1 \quad \forall j \in I \cup \{n+1\} \quad (3.4)$$

$$s_i + P_i + e_i \geq E_i \quad \forall i \in I \quad (3.5)$$

$$s_i + P_i - t_i \leq T_i \quad \forall i \in I \quad (3.6)$$

$$s_i \geq 0 \quad \forall i \in I \cup \{0, n+1\} \quad (3.7)$$

$$e_i \geq 0 \quad \forall i \in I \quad (3.8)$$

$$t_i \geq 0 \quad \forall i \in I \quad (3.9)$$

$$y_{ij} \in \{0, 1\} \quad \forall i, j \in I \cup \{0, n+1\} \quad (3.10)$$

A função objetivo, representada pela equação (3.1), busca a minimização da soma ponderada das antecipações e atrasos. As restrições (3.2) garantem que existe tempo suficiente para executar uma tarefa i e preparar a máquina antes do início do processamento da tarefa seguinte j , ou seja, elas garantem que o processamento da tarefa j não será iniciado antes do término do processamento da tarefa i . As restrições (3.3) e (3.4) garantem que cada tarefa terá apenas uma tarefa imediatamente sucessora e uma tarefa imediatamente antecessora, respectivamente. As restrições (3.5) e (3.6) definem as antecipações e os atrasos de acordo com as respectivas janelas de entrega de cada tarefa. As restrições (3.7), (3.8), (3.9) e (3.10) dizem respeito ao tipo de variáveis.

3.2 Primeiro Modelo Proposto - MPLIM-BG

O modelo de programação linear inteira mista (PLIM) apresentado a seguir é baseado naquele proposto por Gomes Júnior *et al.* (2007) (ver Seção 3.1, página 29) e será denotado por MPLIM-BG. Diferentemente de tais autores, é utilizada apenas uma tarefa fictícia, denominada 0 (zero), para auxiliar na modelagem. Esta tarefa é sequenciada duas vezes, sendo uma na primeira posição e outra na última. Considera-se que $P_0 = 0$ e que $S_{0i} = S_{i0} = 0$, $\forall i \in I$.

Considerando s_i , e_i , t_i e y_{ij} tais como no MPLIM-G, tem-se que as variáveis $y_{n+1,i}$ e $y_{i,n+1}$, $\forall i \in I \cup \{0\}$, e s_{n+1} deixam de ser necessárias no MPLIM-BG. Sendo assim, se comparado ao MPLIM-G, no MPLIM-BG há redução de $2n+3$ variáveis e $n+1$ restrições, com n representando o número de tarefas a serem sequenciadas.

Finalmente, o MPLIM-BG é obtido do MPLIM-G substituindo-se as restrições (3.2), (3.3), (3.4), (3.7) e (3.10) pelas restrições (3.11), (3.12), (3.13), (3.14) e (3.15), respectivamente, apresentadas a seguir:

$$s_i + P_i + y_{ij}(M + S_{ij}) - M \leq s_j \quad \forall i \in I \cup \{0\}, \forall j \in I \text{ e } i \neq j \quad (3.11)$$

$$\sum_{j=0, j \neq i}^n y_{ij} = 1 \quad \forall i \in I \cup \{0\} \quad (3.12)$$

$$\sum_{i=0, i \neq j}^n y_{ij} = 1 \quad \forall j \in I \cup \{0\} \quad (3.13)$$

$$s_i \geq 0 \quad \forall i \in I \cup \{0\} \quad (3.14)$$

$$y_{ij} \in \{0, 1\} \quad \forall i, j \in I \cup \{0\} \quad (3.15)$$

3.3 Segundo Modelo Proposto - MPLIM-IT

A seguir é proposto um modelo de programação linear inteira mista (PLIM) indexado no tempo para representar o PSUMAA abordado. Tal modelo, denotado por MPLIM-IT, é baseado no trabalho de Paula (2008), que utilizou a discretização do tempo para modelar o problema de minimização dos atrasos em sequenciamento de máquinas paralelas com tempos de preparação dependentes da sequência. O MPLIM-IT é descrito como segue.

Seja $H = \{h_0, h_1, h_2, \dots, h_L\}$ o conjunto com as possíveis datas de início de processamento das tarefas de I e considere x_{ih} variáveis binárias, sendo $x_{ih} = 1$ se a tarefa i é programada para iniciar na data h e $x_{ih} = 0$, caso contrário, $\forall i \in I$ e $\forall h \in H$.

Se e_i e t_i representam, respectivamente, as unidades de tempo de antecipação e tempo de atraso da tarefa $i \in I$, a formulação MPLIM-IT pode ser escrita assim:

$$\min \quad Z = \sum_{i=1}^n (\alpha_i e_i + \beta_i t_i) \quad (3.16)$$

$$\text{s.a. } x_{ih} + \sum_{u=h}^{\min(h+P_i+S_{ij}-1, H_L)} x_{ju} \leq 1 \quad \forall i, j \in I, \forall h \in H \quad (3.17)$$

e $i \neq j$

$$\sum_{h=H_0}^{H_L} x_{ih} = 1 \quad \forall i \in I \quad (3.18)$$

$$\sum_{h=H_0}^{H_L} x_{ih} h + P_i + e_i \geq E_i \quad \forall i \in I \quad (3.19)$$

$$\sum_{h=H_0}^{H_L} x_{ih} h + P_i - t_i \leq T_i \quad \forall i \in I \quad (3.20)$$

$$e_i \geq 0 \quad \forall i \in I \quad (3.21)$$

$$t_i \geq 0 \quad \forall i \in I \quad (3.22)$$

$$x_{ih} \in \{0, 1\} \quad \forall i \in I \text{ e } \forall h \in H \quad (3.23)$$

A função objetivo representada pela equação (3.16) tem como critério de otimização a minimização da soma ponderada das antecipações e dos atrasos. As restrições (3.17) garantem que existe tempo suficiente para executar uma tarefa i e preparar a máquina antes do início do processamento da tarefa seguinte j , ou seja, elas garantem que o processamento da tarefa j não será iniciado antes do término do processamento da tarefa i . As restrições (3.18) garantem que cada tarefa seja executada uma única vez. As restrições (3.19) e (3.20) definem as antecipações e os atrasos de acordo com as respectivas janelas de entrega de cada tarefa. As restrições (3.21), (3.22) e (3.23) determinam os domínios das variáveis do problema.

É importante observar que, diferentemente do MPLIM-G e do MPLIM-BG, o MPLIM-IT somente é válido se o problema satisfazer a desigualdade triangular, ou seja, se as condições (3.24) a seguir forem verificadas:

$$S_{ik} \leq S_{ij} + S_{jk} + P_j \quad \forall i, j, k \in I, i \neq j, i \neq k \text{ e } j \neq k \quad (3.24)$$

No PSUMAA-JE-TPD abordado, a execução de qualquer tarefa pode ser iniciada em qualquer momento futuro ($H = [0, +\infty[$). Portanto, caso se tenha a garantia de que na sequência ótima o início do processamento da primeira tarefa não ocorre antes de h_{inf} e que o início do processamento da última tarefa é no máximo h_{sup} , pode-se discretizar o conjunto H no intervalo $[h_{inf}, h_{sup}]$, isto é, adotar $H = \{h_{inf}, h_{inf}+1, h_{inf}+2, \dots, h_{sup}\}$. Deste modo, como o MPLIM-IT é fortemente sensível à cardinalidade de H , quanto menor for o intervalo $[h_{inf}, h_{sup}]$, menor será o número de variáveis do modelo.

Capítulo 4

Algoritmo Proposto

Na prática, a aplicação das formulações de programação matemática apresentadas na Seção 3 é limitada a problemas de pequeno porte, ou seja, a problemas com poucas tarefas. Com o objetivo de resolver casos do PSUMAA-JE-TPD de maior dimensão (com um número de tarefas mais elevado), neste capítulo são propostos algoritmos heurísticos que combinam os procedimentos GRASP (ver Seção 2.2.3), Princípio da Otimalidade Próxima - POP (ver Seção 2.2.4) e Descida em Vizinhança Variável - VND (ver Seção 2.2.3).

4.1 GPV

Basicamente, a metodologia proposta – denotada por GPV – é composta de duas fases. Na primeira (linhas 1 a 12 do Algoritmo 10), gera-se uma solução com base na metaheurística GRASP, no POP e na metaheurística VND. Na segunda (linha 13 do Algoritmo 10), faz-se o pós-refinamento da solução proveniente da fase anterior por meio de outro VND. O GPV pode ser representado pelo Algoritmo 10. A função de avaliação f , a função adaptativa gulosa g , a vizinhança V , o número máximo de iterações sem melhora ($GRASP_{max}$) do GRASP, o número máximo de iterações sem melhora (VND_{max}) do VND, o número máximo de iterações sem melhora (MRD_{max}) do MRD e o parâmetro γ da fase de construção GRASP são as entradas do método.

4.2 Representação de uma Solução

Uma solução (sequência) para o PSUMAA-JE-TPD de n tarefas é representada por um vetor v de n posições, em que cada posição $i = 1, 2, \dots, n$ indica a ordem de produção da tarefa v_i . Por exemplo, dada a sequência $v = (5, 3, 2, 1, 4, 6)$, a tarefa 5 é a primeira a ser executada e a tarefa 6 a última.

4.3 Vizinhança de uma Solução

Para explorar o espaço de soluções são usados três tipos de movimentos: troca da ordem de processamento de duas tarefas da sequência de produção, realocação de uma

Algoritmo $GPV(f(\cdot), g(\cdot), V(\cdot), GRASP_{max}, VND_{max}, MRD_{max}, \gamma)$ <pre> 1 $f^* \leftarrow +\infty;$ 2 $Iter \leftarrow 0;$ 3 <u>enquanto</u> ($Iter < GRASP_{max}$) <u>faça</u> 4 $Iter \leftarrow Iter + 1;$ 5 $v_0 \leftarrow Constroi - Solucao - com - POP(f(\cdot), g(\cdot), N^R(\cdot), MRD_{max}, \gamma);$ 6 $v \leftarrow VND_1(v_0, V(\cdot), VND_{max});$ 7 <u>se</u> ($f(v) < f^*$) <u>então</u> 8 $v^* \leftarrow v;$ 9 $f^* \leftarrow f(v);$ 10 $Iter \leftarrow 0;$ 11 <u>fim-se</u> 12 <u>fim-enquanto</u>; 13 $v^* \leftarrow VND_2(v^*, V(\cdot));$ 14 Retorne $v^*;$ fim $GPV;$ </pre>

Algoritmo 10: Método GPV

tarefa para outra posição da sequência e realocação de um bloco de tarefas para outra posição da sequência. Esses movimentos definem, respectivamente, as vizinhanças N^T , N^R e N^{RB} .

Na vizinhança N^T , um exemplo de vizinho da solução $v = (5, 3, 2, 1, 4, 6)$ é a solução $v' = (5, 4, 2, 1, 3, 6)$, pois v' é obtida por meio da troca da ordem de processamento da segunda com a quinta tarefa de v . Dada uma sequência de produção de n tarefas, cada tarefa na sequência pode trocar de posição com qualquer uma das outras $n - 1$ tarefas. Ademais, trocar a ordem de processamento da i -ésima tarefa na sequência com a j -ésima tarefa é equivalente à troca da ordem de processamento da j -ésima tarefa na sequência com a i -ésima tarefa, para todo $i, j \in \{1, 2, \dots, n\}$. Sendo assim, para uma solução do PSUMAA-JE-TPD com n tarefas, há $n(n - 1)/2$ vizinhos distintos na vizinhança N^T .

A solução $v' = (5, 2, 1, 4, 3, 6)$ é um exemplo de vizinho da solução $v = (5, 3, 2, 1, 4, 6)$ na vizinhança N^R , pois é obtida pela realocação da tarefa 3, que está na segunda posição em v , para a quinta posição na sequência de produção. Dada uma sequência de produção de n tarefas, cada tarefa na sequência pode ser realocada para outras $n - 1$ posições distintas. Além disso, realocar a tarefa da i -ésima posição para a posição $i + 1$ na sequência é equivalente a realocar a tarefa da posição $i + 1$ para a i -ésima posição na sequência, para todo $i \in \{1, 2, \dots, n - 1\}$. Sendo assim, para uma solução do PSUMAA-JE-TPD com n tarefas, há $n(n - 1) - (n - 1) = (n - 1)^2$ vizinhos distintos na vizinhança N^R .

Já na vizinhança N^{RB} , um exemplo de vizinho da solução $v = (5, 3, 2, 1, 4, 6)$ é a solução $v' = (5, 4, 6, 3, 2, 1)$, pois é obtida pela realocação do bloco de três tarefas $\langle 3, 2, 1 \rangle$, que está sequenciado após a tarefa 5 em v , para depois da tarefa 6 na sequência de produção, ou seja, para duas posições sucessoras. Dada uma sequência de produção de n tarefas, existem $n - k + 1$ blocos distintos de k tarefas, para todo $k \in \{1, 2, \dots, n\}$, e cada um destes blocos podem ser realocados para $n - k$ posições distintas na sequência

de produção. Além disso, realocar um bloco de i tarefas para j posições sucessoras é equivalente à realocação de um bloco de j tarefas para i posições antecessoras, para todo $i, j \in \{1, 2, \dots, n\}$ e $i + j \leq n$. Sendo assim, para uma solução do PSUMAA-JE-TPD com n tarefas, há $\left(\sum_{k=1}^n (n-k+1)(n-k)\right)/2 = \left(\sum_{k=1}^n (n-k)^2 + \sum_{k=1}^n (n-k)\right)/2 = \left(n(2n-1)(n-1)/6 + n(n-1)/2\right)/2 = n(n-1)(n+1)/6$ vizinhos distintos na vizinhança N^{RB} .

Com o objetivo de evitar a visita a vizinhos já visitados e consequentemente reduzir o custo computacional, todas as observações mencionadas nos parágrafos acima foram consideradas na implementação computacional do algoritmo proposto.

4.4 Função de Avaliação

Como os movimentos utilizados não geram soluções inviáveis, uma sequência v é avaliada pela função f dada pela equação (4.1), a qual deve ser minimizada:

$$f(v) = \sum_{i=1}^n (\alpha_i e_i + \beta_i t_i) \quad (4.1)$$

Para determinar os valores de e_i e t_i , é utilizado o algoritmo de determinação das datas ótimas de início de processamento das tarefas (ADDOIPT) apresentado a seguir.

4.4.1 Algoritmo de Determinação das Datas Ótimas de Início de Processamento das Tarefas

Dada uma sequência de produção, é utilizado o algoritmo de determinação das datas ótimas de início de processamento das tarefas (ADDOIPT) proposto por Gomes Júnior *et al.* (2007) e baseado nos trabalhos de Wan & Yen (2002) e Lee & Choi (1995). O restante desta subseção é destinado à apresentação deste algoritmo.

Seja $V = (J_1, J_2, \dots, J_n)$ uma sequência do conjunto de tarefas I . Logo $J_i \in I$ e $J_i \neq J_j$, para todo $i, j \in \{1, 2, \dots, n\}$. Diz-se que uma sub-sequência de tarefas $B = (J_u, J_{u+1}, \dots, J_v) \subseteq V$ forma um bloco na sequência V se as tarefas em B são sequenciadas consecutivamente sem tempo ocioso entre elas e existe ociosidade entre as tarefas J_{u-1} e J_u e as tarefas J_v e J_{v+1} .

Considere que há l blocos em V , ou seja, $V = B_1 \cup B_2 \cup \dots \cup B_l$. Sejam $prim(j)$ e $ult(j)$ a primeira e a última tarefa no bloco B_j , respectivamente.

O processamento da primeira tarefa (J_1) da sequência V é programado para ser finalizado na data E_{J_1} se $P_{J_1} \leq E_{J_1}$ ou iniciado na data 0 (finalizando na data P_{J_1}) se $P_{J_1} > E_{J_1}$.

Seja C_k a data de conclusão da tarefa k .

As demais tarefas são programadas da seguinte forma:

- Se $C_{J_k} + S_{(J_k)(J_{k+1})} + P_{J_{k+1}} \geq E_{J_{k+1}}$, a tarefa J_{k+1} tem seu processamento programado para ser finalizado na data $C_{J_k} + S_{(J_k)(J_{k+1})} + P_{J_{k+1}}$ e passa a ser a última tarefa do bloco corrente.

- Caso contrário, se $C_{J_k} + S_{(J_k)(J_{k+1})} + P_{J_{k+1}} < E_{J_{k+1}}$, a tarefa J_{k+1} tem seu processamento programado para ser finalizado na data $E_{J_{k+1}}$ e um novo bloco é iniciado.

Após a determinação da data de conclusão do processamento de cada tarefa, os blocos são reposicionados em suas respectivas posições ótimas na subsequência corrente. Deste modo, após a inserção da última tarefa na sequência, todos os blocos estão otimamente alocados.

A antecipação e o atraso de uma tarefa i na sequência são dados, respectivamente por $\max\{0, E_i - C_i\}$ e $\max\{0, C_i - T_i\}$, respectivamente. Consequentemente, o custo gerado por antecipação e atraso de uma tarefa i concluída na data C_i pode ser dado pela função $g_i(C_i) = \alpha_i \cdot \max\{0, E_i - C_i\} + \beta_i \cdot \max\{0, C_i - T_i\}$.

Logo o custo para deslocar um bloco B_j , x unidades de tempo para a esquerda, pode ser expresso pela equação:

$$Custo_j(x) = \sum_{i \in B_j} g_i(C_i - x), \quad (4.2)$$

em que $C_i - x$ é a nova data de conclusão do processamento da tarefa i .

Claramente a função $g_i(C_i - x)$ é uma função linear convexa por partes em relação a x para todo $i \in \{1, 2, \dots, n\}$.

Lema 1 (Wan & Yen, 2002) *O somatório de duas funções lineares convexas por partes é também uma função linear convexa por partes.*

Baseado no Lema 1, tem-se:

Proposição 1 (Wan & Yen, 2002) *$Custo_j(x)$ é uma função linear convexa por partes em relação a x , para todo $j \in \{1, 2, \dots, l\}$.*

Devido a natureza linear convexa por partes da função custo, o custo mínimo do bloco B_j ocorre nos pontos extremos de sua função $Custo_j$, isto é, no início ou no final da janela de entrega de uma das tarefas no bloco. Tal custo mínimo pode ser facilmente obtido pela comparação dos somatórios das penalidades das tarefas no bloco, onde cada um destes somatórios é obtido ao programar uma tarefa do bloco para ser concluída no início ou no final de sua janela de entrega. Todos os possíveis somatórios são calculados e o menor deles corresponde ao ponto de custo mínimo do bloco.

Quando o ponto mínimo do bloco B_j é encontrado, todo o bloco é movido em direção ao ponto mínimo até que um dos três casos seguintes aconteça:

- (i) $s_{prim(j)} = 0$;
- (ii) O ponto mínimo é alcançado;
- (iii) $s_{prim(j)}$ torna-se igual a $C_{ult(j-1)} + S_{ult(j-1)(prim(j))}$.

Se ocorrer o caso (iii), o bloco B_j é unido ao bloco B_{j-1} , resultando em um novo bloco B_{j-1} . Logo o ponto mínimo do novo bloco B_{j-1} deve ser obtido.

O procedimento anterior deve ser realizado até que ocorra o caso (i) ou (ii) para cada bloco.

A seguinte proposição garante que o algoritmo apresentado determina as datas ótimas de início do processamento das tarefas de uma dada sequência.

Proposição 2 (Wan & Yen, 2002) *O custo total de uma dada sequência de tarefas alcança seu valor ótimo se cada bloco B_j na sequência alcançar seu ponto mínimo, exceto $s_{prim(1)}$, que pode ser igual a zero para o primeiro bloco.*

O Algoritmo 11, adaptado de Gomes Júnior (2007), descreve o procedimento ADDOIPT. Nele, o procedimento *MudaBloco* movimenta o bloco J até seu ponto de mínimo e, se necessário, combina-o com o bloco anterior.

```

procedimento Determina – Datas – Otimas – Inicio( $n, s$ )
1   $B \leftarrow 1$ ;
2   $prim(B) \leftarrow ult(B) \leftarrow 1$ ;
3   $s_1 \leftarrow \max(0, E_1 - P_1)$ ;
4   $C_1 \leftarrow \max(E_1, P_1)$ ;
5  para ( $i = 2, \dots, n$ ) faça
6    se ( $C_{i-1} + P_i + S_{(i-1)(i)} < E_i$ ) então
7       $B \leftarrow B + 1$ ;
8       $prim(B) \leftarrow ult(B) \leftarrow i$ ;
9       $s_i \leftarrow E_i - P_i + S_{(i-1)(i)}$ ;
10      $C_i \leftarrow E_i$ ;
11   senão
12     se ( $C_{i-1} + P_i + S_{(i-1)(i)} = E_i$ ) então
13        $ult(B) \leftarrow i$ ;
14        $s_i \leftarrow C_{i-1} + S_{(i-1)(i)}$ ;
15        $C_i \leftarrow E_i$ ;
16     senão
17        $ult(B) \leftarrow i$ ;
18        $s_i \leftarrow C_{i-1} + S_{(i-1)(i)}$ ;
19        $C_i \leftarrow s_i + P_i$ ;
20   fim-se
21 fim-se
22 repita (até todos os blocos estiverem no ponto mínimo ou  $s_{prim(1)} = 0$ )
23   MudaBloco( $B$ );
24 fim-repita;
25 fim-para;
26    $f(s) \leftarrow \sum_{i=1}^n \alpha_i \cdot \max\{0, E_i - C_i\} + \beta_i \cdot \max\{0, C_i - T_i\}$ ;
27 Retorne  $f(s)$ ;
fim Determina – Datas – Otimas – Inicio;

```

Algoritmo 11: Algoritmo para determinar as datas ótimas de início de processamento das tarefas.

Mais detalhes do ADDOIPT podem ser encontrados em Gomes Júnior (2007).

4.5 Construção de uma Solução

Nesta etapa do GPV (linha 5 do Algoritmo 10), uma solução é formada, tarefa por tarefa, de forma parcialmente gulosa, seguindo as ideias da fase de construção do algoritmo GRASP (ver Seção 2.2.3) e do POP (ver Seção 2.2.4). O Algoritmo 12 representa o procedimento proposto para construir uma solução.

procedimento *Constroi – Solucao – com – POP*($f(\cdot)$, $g(\cdot)$, $V(\cdot)$, $MRDmax$, γ)

```

1  $s \leftarrow \emptyset$ ;
2 Inicialize o conjunto  $C$  de candidatos;
3 enquanto ( $C \neq \emptyset$ ) faça
4    $g(t_{min}) \leftarrow \min\{g(t) \mid t \in C\}$ ;
5    $g(t_{max}) \leftarrow \max\{g(t) \mid t \in C\}$ ;
6    $LRC \leftarrow \{t \in C \mid g(t) \leq g(t_{min}) + \gamma(g(t_{max}) - g(t_{min}))\}$ ;
7   Aleatoriamente, selecione um elemento  $t \in LRC$ ;
8    $s \leftarrow s \cup \{t\}$ ;
9    $s \leftarrow MRD(s, N^R(\cdot), MRDmax)$ ;                                /* Aplica o POP */
10  Atualize o conjunto  $C$  de candidatos;
11 fim-enquanto;
12 Retorne  $s$ ;
fim Constroi – Solucao – com – POP;
```

Algoritmo 12: Construção de uma solução com POP.

Foram experimentadas 11 funções adaptativas gulosas para estimar o benefício da inserção de cada tarefa j ainda não sequenciada na iteração i , sendo elas:

- $g_1(j) = E_j$;
- $g_2(j) = T_j$;
- $g_3(j) = E_j + T_j$;
- $g_4(j) = T_j - P_j$;
- $g_5(j) = E_j - P_j$;
- $g_6(j) = \frac{(E_j - P_j)\alpha_j + (T_j - P_j)\beta_j}{\alpha_j + \beta_j}$;
- $g_7(j) = E_j + T_j - 2P_j$;
- $g_8(j) = \frac{2E_j}{\max\{E_k \mid k \in F_i\}} + \frac{2T_j}{\max\{T_k \mid k \in F_i\}} + \frac{\alpha_j}{\max\{\alpha_k \mid k \in F_i\}} - \frac{\beta_j}{\max\{\beta_k \mid k \in F_i\}} + \frac{P_j}{\max\{P_k \mid k \in F_i\}} + \frac{S_{i-1,j}}{\max\{S_{i-1,k} \mid k \in F_i\}}$, onde F_i representa o conjunto das tarefas ainda não sequenciadas até a i -ésima iteração (para $i = 1$, no lugar de $\frac{S_{i-1,j}}{\max\{S_{i-1,k} \mid k \in F_i\}}$, é utilizado $-\frac{\text{média}\{S_{kj} \mid k \in I \text{ e } k \neq j\}}{\max\{S_{kl} \mid k, l \in I \text{ e } k \neq l\}}$);

- $g_9(j) = \frac{E_j}{\max\{E_k | k \in F_i\}} + \frac{2T_j}{\max\{T_k | k \in F_i\}} + \frac{\alpha_j}{\max\{\alpha_k | k \in F_i\}} - \frac{\beta_j}{\max\{\beta_k | k \in F_i\}} + \frac{P_j}{\max\{P_k | k \in F_i\}} + \frac{S_{i-1,j}}{\max\{S_{i-1,k} | k \in F_i\}}$, onde F_i representa o conjunto das tarefas ainda não sequenciadas até a i -ésima iteração (para $i = 1$, no lugar de $\frac{S_{i-1,j}}{\max\{S_{i-1,k} | k \in F_i\}}$, é utilizado $-\frac{\text{média}\{S_{kj} | k \in I \text{ e } k \neq j\}}{\max\{S_{kl} | k, l \in I \text{ e } k \neq l\}}$);
- $g_{10}(j) = \frac{E_j}{\max\{E_k | k \in F_i\}} + \frac{T_j}{\max\{T_k | k \in F_i\}} + \frac{\alpha_j}{\max\{\alpha_k | k \in F_i\}} - \frac{\beta_j}{\max\{\beta_k | k \in F_i\}} + \frac{P_j}{\max\{P_k | k \in F_i\}} + \frac{S_{i-1,j}}{\max\{S_{i-1,k} | k \in F_i\}}$, onde F_i representa o conjunto das tarefas ainda não sequenciadas até a i -ésima iteração (para $i = 1$, no lugar de $\frac{S_{i-1,j}}{\max\{S_{i-1,k} | k \in F_i\}}$, é utilizado $-\frac{\text{média}\{S_{kj} | k \in I \text{ e } k \neq j\}}{\max\{S_{kl} | k, l \in I \text{ e } k \neq l\}}$);
- $g_{11}(j) = \frac{2E_j}{\max\{E_k | k \in F_i\}} + \frac{T_j}{\max\{T_k | k \in F_i\}} + \frac{\alpha_j}{\max\{\alpha_k | k \in F_i\}} - \frac{\beta_j}{\max\{\beta_k | k \in F_i\}} + \frac{P_j}{\max\{P_k | k \in F_i\}} + \frac{S_{i-1,j}}{\max\{S_{i-1,k} | k \in F_i\}}$, onde F_i representa o conjunto das tarefas ainda não sequenciadas até a i -ésima iteração (para $i = 1$, no lugar de $\frac{S_{i-1,j}}{\max\{S_{i-1,k} | k \in F_i\}}$, é utilizado $-\frac{\text{média}\{S_{kj} | k \in I \text{ e } k \neq j\}}{\max\{S_{kl} | k, l \in I \text{ e } k \neq l\}}$);

Cada função g_k , $k = 1, 2, \dots, 11$, foi utilizada para construir uma solução gulosa para cada um dos 608 problemas-teste com até 100 tarefas apresentados na Subseção 5.1. A qualidade da solução encontrada por cada função adaptativa gulosa g_k , aplicada a um determinado problema-teste, foi mensurada pela equação (4.3), onde f_{g_k} representa o valor da função objetivo da solução encontrada pela função adaptativa gulosa g_k e f_{g^*} denota o valor da função objetivo da melhor solução encontrada pelas 11 funções gulosas.

$$Desvio_{g_k} = \frac{f_{g_k} - f_{g^*}}{f_{g^*}} \quad (4.3)$$

Os resultados encontrados estão sintetizados na Tabela 4.1. Nesta tabela, a primeira linha indica as funções adaptativas gulosas testadas. Já na segunda linha, são apresentadas as médias dos desvios encontrados pelas respectivas funções aplicadas aos 608 problemas-teste.

Tabela 4.1: Resultados apresentados pelas funções adaptativas gulosas testadas.

Função	g_1	g_2	g_3	g_4	g_5	g_6	g_7	g_8	g_9	g_{10}	g_{11}
Desvio Médio	0,47	0,46	0,45	0,54	0,54	0,56	0,53	0,28	0,38	0,57	0,36

Na Tabela 4.1 é possível observar que as funções g_1 , g_2 , g_3 , g_8 , g_9 e g_{11} apresentaram desvios médios menores que 0,5. Sendo assim, optou-se por utilizar estas seis funções, uma em cada iteração do GRASP. Tais funções são utilizadas segundo a ordem crescente do

desvio médio encontrado, ou seja, na primeira iteração utiliza-se a função g_8 , na segunda utiliza-se a função g_{11} e assim por diante, até que todas as seis funções sejam utilizadas; repetindo-se a aplicação desta sequência a seguir.

Ainda na etapa de construção de uma solução, é aplicado o POP (linha 9 do Algoritmo 12). Assim, sempre que uma nova tarefa é inserida na solução em formação, esta é submetida a uma busca local. A busca local utilizada consiste em um método de descida, o qual é descrita a seguir.

Para determinar a vizinhança utilizada no método de descida relativo ao POP, foi realizado o seguinte procedimento em um problema-teste com 30 tarefas. Cada uma das seis funções adaptativas gulosas determinadas acima foi utilizada para construir uma solução para o problema. As soluções construídas foram, então, submetidas ao procedimento de Descida de Primeira Melhora com a vizinhança N^T . Este mesmo procedimento foi realizado com a vizinhança N^R . Como a ordem de visita aos vizinhos é determinante em um método de descida, o procedimento descrito foi realizado 20 vezes com cada vizinhança, sendo a ordem de visitas sempre escolhida aleatoriamente. No final de todas as execuções, as soluções obtidas com a vizinhança N^R apresentaram uma média 9,55% menor que a média das soluções obtidas com a vizinhança N^T . Além disso, a melhor solução obtida com a vizinhança N^R foi 7,63% melhor que a melhor solução obtida com a N^T . Estes resultados sugerem que a vizinhança N^R é capaz de encontrar soluções de melhor qualidade e com melhor média que as soluções obtidas com a vizinhança N^T e, por isso, foi a escolhida para o método de descida usado no POP.

Como o POP é aplicado sempre que uma tarefa é inserida na solução em formação, utilizou-se uma descida randômica com vizinhança N^R , visto que este método de descida demanda menos tempo computacional que a Descida de Primeira Melhora. Sendo assim, dada uma solução, aleatoriamente escolhe-se uma tarefa na sequência e uma nova posição para ela. Se a nova sequência produzir uma solução com um valor menor para a função de avaliação, a nova sequência é aceita e passa a ser a solução corrente; caso contrário, é testada outra realocação aleatória. A busca é interrompida após $MRDMax$ realocações consecutivas sem melhora na solução corrente, sendo $MRDMax$ um parâmetro do método.

4.6 VND₁

Para refinar as soluções contruídas na primeira fase do GPV (linha 6 do Algoritmo 10), utiliza-se um procedimento baseado em Descida em Vizinhança Variável (ver Seção 2.2.3), denotado por VND₁. Basicamente o VND₁ consiste em três passos:

Passo 1: Descida randômica com a vizinhança N^T ;

Passo 2: Descida randômica com a vizinhança N^R ;

Passo 3: Descida randômica com a vizinhança N^{RB} ;

No *Passo 1*, duas tarefas são escolhidas de forma aleatória e a ordem de seus processamentos na sequência de produção são trocadas. Se a nova sequência produzir uma solução com um valor menor para a função de avaliação, a nova sequência é aceita e passa a ser a solução corrente; caso contrário, é testada outra troca aleatória. O *Passo 1* termina

quando ocorrer $MRDmax$ trocas consecutivas sem melhora na solução corrente, sendo $MRDMax$ um parâmetro do método. Neste último caso, passa-se para o *Passo 2*.

No *Passo 2*, aleatoriamente escolhe-se uma tarefa na sequência e uma nova posição para ela. Se a nova sequência produzir uma solução com um valor menor para a função de avaliação, a nova sequência passa a ser a solução corrente e volta-se ao *Passo 1*; caso contrário, é testada outra realocação aleatória. O *Passo 2* é interrompido após $MRDMax$ realocações consecutivas sem melhora na solução corrente. Neste último caso, passa-se para o *Passo 3*.

No *Passo 3*, aleatoriamente escolhe-se um bloco de tarefas na sequência de produção (o tamanho do bloco também é escolhido de forma aleatória) e uma nova posição para ele. Se a nova sequência produzir uma solução com um valor menor para a função de avaliação, a nova sequência passa a ser a solução corrente e volta-se ao *Passo 1*; caso contrário, é testada outra realocação de bloco aleatória. O *Passo 3* é interrompido após $MRDMax$ realocações de blocos consecutivas sem melhora na solução corrente. Neste último caso, o VND_1 é interrompido e a melhor solução encontrada é retornada.

É interessante observar que a solução oriunda do VND_1 não é necessariamente um ótimo local em relação às vizinhanças adotadas, visto que nem toda a vizinhança é analisada em cada cada passo.

4.7 VND_2

Como a solução proveniente da primeira fase do GPV (linhas 1 a 12 do Algoritmo 10) não é necessariamente um ótimo local em relação às vizinhanças adotadas, ela é submetida a uma busca local mais efetiva, no caso, também baseada na Busca em Vizinhança Variável (VND_2). Nesta, a exploração do espaço de soluções é realizada de acordo com os seguintes passos:

Passo 1: Descida Primeiro de Melhora com a vizinhança N^T ;

Passo 2: Descida completa com a vizinhança N^{RB} ;

Se no *Passo 1*, todos os vizinhos com relação à vizinhança N^T não são de melhora, passa-se para o *Passo 2*.

No *Passo 2*, em que são realizadas realocações de blocos de tarefas, inicialmente testam-se todas as possíveis realocações de blocos com apenas uma tarefa e, quando não for mais possível melhorar a solução com um determinado tamanho de bloco, passa-se a explorar movimentos com blocos de tamanho imediatamente superior. Se uma solução de melhora é encontrada, volta-se ao *Passo 1*. O procedimento é interrompido quando um ótimo local com relação às vizinhanças N^T e N^{RB} é encontrado. Observa-se que a vizinhança N^R corresponde a realocar blocos com apenas uma tarefa. Logo, a vizinhança N^R está contida na vizinhança N^{RB} e um ótimo local em relação à esta última também é um ótimo local em relação à primeira.

4.8 Redução do Espaço de Busca - GPV-REB

A avaliação de todos os vizinhos gerados em uma busca local requer um grande e muitas vezes injustificável esforço computacional (França Filho, 2007). Por isso, é inte-

ressante utilizar estratégias que tenham por objetivo evitar a avaliação de soluções que apresentam determinadas características julgadas ruins e, assim, reduzir o esforço computacional. Como consequência, mais vizinhos podem ser explorados ou, então, problemas de dimensões maiores podem ser resolvidos.

No caso do PSUMAA-JE-TPD, se duas tarefas $i, j \in I$ satisfazem as condições (4.4), (4.5), (4.6) e (4.7):

$$E_i \leq E_j, \quad (4.4)$$

$$\alpha_i \leq \alpha_j, \quad (4.5)$$

$$T_i \leq T_j \text{ e} \quad (4.6)$$

$$\beta_i \geq \beta_j, \quad (4.7)$$

espera-se que boas soluções possuem a tarefa i sequenciada antes da tarefa j . Deste modo, se i é sequenciada antes de j na sequência corrente, é proposta a seguinte estratégia de redução do espaço de busca:

- Durante a exploração da vizinhança N^T , a solução produzida pela troca das ordens de processamentos da tarefa i com a tarefa j é rejeitada sem ser avaliada.
- Durante a exploração da vizinhança N^R , as realocações da tarefa i para posições sucessoras à tarefa j são descartadas sem serem avaliadas. Do mesmo modo, as realocações da tarefa j para as posições antecessoras à tarefa i também são descartadas.
- Durante a exploração da vizinhança N^{RB} , as realocações de blocos que levam a tarefa i a ser sequenciada depois da tarefa j na sequência de produção são descartadas sem serem avaliadas.

A Figura 4.1 ilustra a situação em que $E_3 \leq E_2$, $\alpha_3 \leq \alpha_2$, $T_3 \leq T_2$ e $\beta_3 \geq \beta_2$. Neste caso, as soluções obtidas pela troca das ordens de processamento da tarefa 3 com a tarefa 2 ou por realocações, de tarefas ou de blocos, que levam a tarefa 3 a ser sequenciada depois da tarefa 2 são descartadas antes de serem avaliadas.

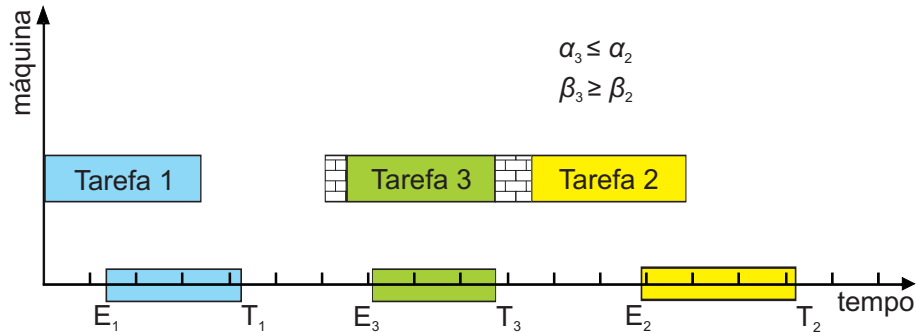


Figura 4.1: Exemplo de situação em que a estratégia de redução do espaço de busca é utilizada.

A estratégia de redução do espaço de busca proposta acima foi utilizada em um algoritmo derivado do algoritmo GPV, denotado por GPV-REB. O GPV-REB se diferencia

do GPV pela utilização da estratégia de redução do espaço de busca nas buscas locais da primeira fase do algoritmo, ou seja, na descida randômica aplicada nas soluções parciais da fase de construção de uma solução (linha 9 do Algoritmo 12) e no VND₁ (linha 6 do Algoritmo 10). A segunda fase dos dois algoritmos é a mesma.

Capítulo 5

Resultados Computacionais

Neste capítulo são apresentados os resultados computacionais obtidos pela aplicação das formulações de programação matemática propostas, bem como dos algoritmos heurísticos desenvolvidos. Primeiramente, são apresentados os conjuntos de problemas-teste utilizados. Na Seção 5.2 são apresentados os resultados obtidos com as formulações de programação matemática apresentadas no Capítulo 3. Na Seção 5.3 são apresentados os resultados obtidos com os algoritmos GPV e GPV-REB, bem como os procedimentos utilizados na determinação dos parâmetros de tais algoritmos.

Todos os experimentos computacionais foram realizados em um computador Intel(R) Core(TM) 2 Duo E8400 2.99 GHz, com 2 GB de memória RAM e sistema operacional Windows XP. Apesar de o processador deste equipamento possuir 2 núcleos, o algoritmo desenvolvido não foi otimizado para multiprocessamento.

5.1 Problemas-teste

Foram utilizados três conjuntos de problemas-teste. O primeiro deles, denotado por BDS, foi proposto por Gomes Júnior (2007), sendo baseado nos trabalhos de Wan & Yen (2002), Liaw (1999) e Mazzini & Armentano (2001). Esta base de dados é constituída por grupos de problemas-teste com 8, 9, 10, 11, 12, 15, 20, 25, 30, 40, 50 e 75 tarefas.

Dada uma tarefa i , o tempo de processamento (P_i), o custo por unidade de atraso (β_i) e o custo por unidade de antecipação (α_i) são valores inteiros selecionados aleatoriamente dentro dos intervalos $[1, 100]$, $[20, 100]$ e $[0, \beta_i]$, respectivamente.

O centro da janela de entrega de $i \in I$ é um valor aleatório dentro do intervalo $\left[\left(1 - FA - \frac{VRJ}{2}\right) \times TTP, \left(1 - FA + \frac{VRJ}{2}\right) \times TTP \right]$, onde TTP é o tempo total de processamento de todas as tarefas, FA é o fator de atraso e VRJ é a variação relativa da janela de entrega. São utilizados os valores 0,1; 0,2; 0,3 e 0,4 para FA e 0,8; 1,0 e 1,2 para VRJ .

Os tamanhos das janelas de entrega são uniformemente e discretamente distribuídos no intervalo $\left[1, \frac{TTP}{n}\right]$, onde n é o número de tarefas de I .

Para toda tarefa $j \in I$ e $j \neq i$, o tempo de *setup* (S_{ij}) é um inteiro aleatório dentro do intervalo $[0, 50]$. Utilizou-se tempos de preparação simétricos, ou seja, $S_{ij} = S_{ji}$.

Como são utilizados quatro valores distintos de FA e três valores distintos de VRJ , há 12 problemas-teste para cada número de tarefas, totalizando 144 problemas-teste. Doravante, esta base de dados será denotada por BDS.

O segundo conjunto de problemas-teste utilizado, denotado por BDNS, foi gerado considerando-se a metodologia apresentada nos trabalhos de Wan & Yen (2002) e Rabadi *et al.* (2004). Foram gerados grupos de problemas-teste com 6, 7, 8, 9, 10, 11, 12, 15, 20, 30, 40, 50, 75, 100 e 150 tarefas.

Dada uma tarefa $i \in I$, o tempo de processamento (P_i), o custo por unidade de atraso (β_i) e o custo por unidade de antecipação (α_i) são valores inteiros selecionados aleatoriamente dentro dos intervalos $[1, 100]$, $[1, 10]$ e $[1, \beta_i]$, respectivamente.

O centro da janela de entrega de i é um valor inteiro aleatório dentro do intervalo $\left[\left(1 - FA - \frac{VRJ}{2}\right) \times TTP, \left(1 - FA + \frac{VRJ}{2}\right) \times TTP \right]$, onde TTP é o tempo total de processamento de todas as tarefas, FA é o fator de atraso e VRJ é a variação relativa da janela de entrega. São utilizados os valores 0,1; 0,3; 0,5 e 0,8 para FA e 0,4; 0,7; 1,0 e 1,3 para VRJ .

O tamanho da janela de entrega de i é um número inteiro selecionado aleatoriamente dentro do intervalo $\left[0, \frac{TTP}{n} \right]$, onde n é o número de tarefas de I .

Para toda tarefa $j \in I$ e $j \neq i$, o tempo de *setup* (S_{ij}) é um inteiro aleatório dentro do intervalo $[5, 30]$. Nesta base de dados, os tempos de *setup* são genéricos e, assim, não são necessariamente simétricos como na base anterior.

Como são utilizados 4 valores distintos de TTP e 4 de VRJ , são 16 problemas-teste de cada tamanho, totalizando 240 problemas-teste. Esta nova base de dados será denotada por BDNS.

O terceiro conjunto de problemas-teste, denotado por BDDT, também foi baseado nos trabalhos de Rabadi *et al.* (2004) e Wan & Yen (2002), mas diferem do anterior na caracterização dos tempos de preparação, que, neste caso, satisfazem a desigualdade triangular. Os problemas-teste deste conjunto foram gerados do modo descrito a seguir.

Dada uma tarefa $i \in I$, o tempo de processamento (P_i), o custo por unidade de atraso (β_i) e o custo por unidade de antecipação (α_i) são números inteiros selecionados aleatoriamente dentro dos intervalos $[1, 40]$, $[1, 10]$ e $[1, \beta_i]$, respectivamente.

O centro da janela de entrega de i é um valor inteiro aleatório dentro do intervalo $\left[\left(1 - FA - \frac{VRJ}{2}\right) \times TTP, \left(1 - FA + \frac{VRJ}{2}\right) \times TTP \right]$, em que TTP é o tempo total de processamento de todas as tarefas, FA é o fator de atraso e VRJ é a variação relativa da janela de entrega.

O tamanho da janela de entrega é um valor inteiro selecionado aleatoriamente no intervalo $\left[0, \frac{TTP}{n} \right]$, sendo n o número de tarefas.

Para toda tarefa $j \in I$ e $j \neq i$, o tempo de *setup* (S_{ij}) é um número inteiro escolhido aleatoriamente dentro do intervalo $[5, 15]$. Assim, os tempos de *setup* também não são necessariamente simétricos.

Foram gerados conjuntos de problemas-teste com 6, 7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17, 18, 19 e 20 tarefas, respectivamente, sendo utilizados os valores 0,1; 0,3; 0,5 e 0,8 para FA e 0,4; 0,7; 1,0 e 1,3 para VRJ . Deste modo, há 16 problemas-teste em cada conjunto,

totalizando 240 problemas-teste.

Durante a geração dos problemas-teste desta última base de dados, sempre que um problema gerado não satisfazia a desigualdade triangular, ele era descartado e outro problema com os mesmos valores para FA e VRJ era gerado. Portanto, todos os problemas-teste desta base de dados satisfazem a desigualdade triangular.

5.2 Resultados Obtidos com as Formulações de Programação Matemática

Os modelos de programação matemática apresentados na Seção 3 foram implementados no modelador AMPL e resolvidos pelo otimizador CPLEX, versão 10.1, da ILOG, em sua configuração padrão. Foram realizados experimentos nos problemas-teste de pequeno porte das três bases de dados apresentadas na Seção 5.1, totalizando três baterias de testes.

Para os modelos MPLIM-G e MPLIM-BG, foi utilizado $M = 1000$, valor adotado por Gomes Júnior (2007). Para cada problema-teste, o conjunto H do MPLIM-IT foi determinado pelas expressões (5.1) e (5.2):

$$h_{inf} = \max \left(0, \min \left\{ T_i \mid i \in I \right\} + \min \left\{ S_{ij} \mid i, j \in I \text{ e } i \neq j \right\} - \sum_{i \in I} P_i - \sum_{i \in I} \max \left\{ S_{ij} \mid j \in I \text{ e } i \neq j \right\} \right) \quad (5.1)$$

$$\begin{aligned} h_{sup} = & \max \left\{ E_i \mid i \in I \right\} + \sum_{i \in I} P_i + \sum_{i \in I} \max \left\{ S_{ij} \mid j \in I \text{ e } i \neq j \right\} \\ & - \min \left\{ S_{ij} \mid i, j \in I \text{ e } i \neq j \right\} - \min \left\{ P_i \mid i \in I \right\} \\ & - \min \left(\min \left\{ P_i \mid i \in I \right\}, \max \left\{ E_i \mid i \in I \right\} \right) \end{aligned} \quad (5.2)$$

em que h_{inf} garante que todas as tarefas podem ser concluídas sem atraso (ou o processamento da primeira tarefa é iniciado na data 0) e h_{sup} assegura que todas as tarefas podem ser executadas sem antecipação. Portanto, tem-se a garantia de que as datas de início de processamento das tarefas na sequência ótima pertencem ao conjunto $H = \{h_{inf}, h_{inf} + 1, h_{inf} + 2, \dots, h_{sup}\}$.

Considerou-se, ainda, o limite de 3600 segundos para a resolução de cada problema-teste pelo CPLEX. Para os problemas em que este limite de tempo foi atingido, a solução retornada não é necessariamente ótima; porém, um limite inferior para o ótimo de tal problema é retornado pelo CPLEX. Este limite foi utilizado para medir a qualidade da solução retornada, que é obtida pela equação (5.3):

$$gap = \frac{(f^{CPLEX} - L)}{L} \times 100\% \quad (5.3)$$

em que f^{CPLEX} e L representam, respectivamente, o valor da função objetivo da solução e o limite inferior encontrados pelo otimizador. Deste modo, se uma solução ótima é encontrada, tem-se $gap = 0$.

Os resultados obtidos em cada bateria de testes são apresentados a seguir.

5.2.1 Primeira Bateria de Testes - BDS

A primeira bateria de testes foi realizada com os problemas-teste da base de dados BDS. Os problemas-teste dessa base de dados possuem os tempos de *setup* simétricos e não satisfazem, necessariamente, a desigualdade triangular. Deste modo, apenas os modelos MPLIM-G e MPLIM-BG foram utilizados nesta bateria, visto que o MPLIM-IT não permite resolver problemas em que a desigualdade triangular não é satisfeita.

Os resultados encontrados são resumidos na Tabela 5.1, onde a primeira coluna indica o número de tarefas dos problemas-teste de cada conjunto. Ademais, para cada conjunto de problemas, as colunas “Soluções Ótimas” mostram (em porcentagem) em quantos problemas o CPLEX encontrou uma solução ótima com os respectivos modelos, as colunas “ $\overline{gap}$ ” mostram as médias dos $\overline{gap}$ ’s retornados pelos respectivos modelos e as colunas “Tempo Médio” mostram as médias dos tempos (em segundos) demandados por tal otimizador na resolução dos problemas-teste com os respectivos modelos.

Tabela 5.1: Comparação MPLIM-G $\times$ MPLIM-BG aplicados à BDS.

# Tarefas	MPLIM-G			MPLIM-BG		
	Soluções Ótimas (%)	$\overline{gap}$ (%)	Tempo Médio (s)	Soluções Ótimas (%)	$\overline{gap}$ (%)	Tempo Médio (s)
08	100,00	0,00	1,52	100,00	0,00	1,61
09	100,00	0,00	35,12	100,00	0,00	31,71
10	100,00	0,00	86,68	100,00	0,00	77,16
11	66,67	15,78	1850,01	75,00	10,53	1974,97
12	50,00	30,51	1957,23	50,00	30,98	1962,74
Média	83,33	9,26	786,11	85,00	8,30	809,64

Na Tabela 5.1, observa-se que o MPLIM-G e o MPLIM-BG proporcionaram ao CPLEX desempenhos semelhantes para a base de dados em questão. Usando-se esses dois modelos, o CPLEX foi capaz de resolver na otimalidade todos os problemas-teste com até 10 tarefas e com tempos médios de processamentos bem próximos uns dos outros. Para os problemas-teste com 11 tarefas, o MPLIM-BG permitiu resolver 75,00% dos casos, contra 66,67% do MPLIM-G, e proporcionou $\overline{gap}$ inferior ao propiciado pelo MPLIM-G. Para estes problemas-teste, o tempo médio demandado pelo CPLEX com o MPLIM-G foi sensivelmente inferior ao despendido com o MPLIM-BG. Para os problemas-teste com 12 tarefas, ambos os modelos proporcionaram ao otimizador resolver 50,00% dos problemas, com $\overline{gap}$ ’s próximos de 31% e tempos médios próximos de 1960 segundos.

5.2.2 Segunda Bateria de Testes - BDNS

A segunda bateria de testes foi realizada com os problemas-teste da base de dados BDNS. Os problemas-teste dessa base de dados também não satisfazem, necessariamente, a desigualdade triangular. Do mesmo modo, apenas os modelos MPLIM-G e MPLIM-BG

foram utilizados nesta bateria, já que o MPLIM-IT não permite resolver problemas em que a desigualdade triangular não é satisfeita.

A Tabela 5.2 resume os resultados encontrados. Nela, a primeira coluna indica o número de tarefas em cada problema-teste dos conjuntos. Além disso, para cada conjunto de problemas, as colunas “Soluções Ótimas” mostram (em porcentagem) em quantos problemas o CPLEX encontrou uma solução ótima com os respectivos modelos, as colunas “ $\overline{gap}$ ” mostram as médias dos gap ’s com os respectivos modelos e as colunas “Tempo Médio” mostram as médias dos tempos (em segundos) demandados por tal otimizador na resolução dos problemas-teste com os respectivos modelos.

Tabela 5.2: Comparação MPLIM-G $\times$ MPLIM-BG aplicados à BDNS.

# Tarefas	MPLIM-G			MPLIM-BG		
	Soluções Ótimas (%)	$\overline{gap}$ (%)	Tempo Médio (s)	Soluções Ótimas (%)	$\overline{gap}$ (%)	Tempo Médio (s)
06	100,00	0,00	0,10	100,00	0,00	0,09
07	100,00	0,00	0,76	100,00	0,00	0,76
08	100,00	0,00	5,89	100,00	0,00	5,96
09	100,00	0,00	92,46	100,00	0,00	83,15
10	93,75	2,40	823,97	93,75	2,55	832,44
11	43,75	31,30	2275,84	50,00	27,53	2265,33
12	25,00	54,92	2887,26	25,00	54,53	2857,60
Média	80,36	12,66	869,47	81,25	12,09	863,62

Como pode ser observado na Tabela 5.2, os modelos MPLIM-G e MPLIM-BG também propiciaram ao CPLEX comportamentos semelhantes para a base de dados BDNS. Ambos os modelos permitiram resolver na otimalidade todos os problemas-teste com até 9 tarefas e, aproximadamente, a mesma proporção dos demais conjuntos de problemas. Os $\overline{gap}$ ’s e os tempos médios também foram bem próximos. Essencialmente, a diferença mais sensível ocorreu com os problemas-teste com 11 tarefas. Para este conjunto, o CPLEX conseguiu resolver 50,00% dos problemas por meio do MPLIM-BG, enquanto o MPLIM-G proporcionou ao otimizador resolver apenas 43,75% dos mesmos. Os tempos médios demandados com os dois modelos foram bem próximos uns dos outros, mas o $\overline{gap}$ com o MPLIM-BG foi um pouco menor que o $\overline{gap}$ com o MPLIM-G, 27,53% contra 31,30%.

5.2.3 Terceira Bateria de Testes - BDDT

Na terceira bateria de testes utilizou-se os problemas da base de dados BDDT. Como descrito anteriormente, os problemas-teste dessa base de dados satisfazem a desigualdade triangular. Sendo assim, os três modelos de programação matemática apresentados na Seção 3 foram utilizados nesta bateria.

Os resultados encontrados para o MPLIM-G, para o MPLIM-BG e para o MPLIM-IT são resumidos na Tabela 5.3, na qual a primeira coluna indica o número de tarefas dos problemas-teste de cada conjunto. Para cada conjunto de problemas, as colunas “Soluções Ótimas” mostram (em porcentagem) em quantos problemas o CPLEX encontrou uma

solução ótima com os respectivos modelos, as colunas “ $\overline{gap}$ ” mostram as médias dos gap ’s com os respectivos modelos e as colunas “Tempo Médio” mostram as médias dos tempos (em segundos) demandados por tal otimizador na resolução dos problemas-teste com os respectivos modelos.

Tabela 5.3: Comparação MPLIM-G $\times$ MPLIM-BG $\times$ MPLIM-IT aplicados à BDDT.

# Tarefas	MPLIM-G			MPLIM-BG			MPLIM-IT		
	Soluções Ótimas (%)	$\overline{gap}$ (%)	Tempo Médio (s)	Soluções Ótimas (%)	$\overline{gap}$ (%)	Tempo Médio (s)	Soluções Ótimas (%)	$\overline{gap}$ (%)	Tempo Médio (s)
06	100,00	0,00	0,08	100,00	0,00	0,07	100,00	0,00	3,10
07	100,00	0,00	0,88	100,00	0,00	0,87	100,00	0,00	7,15
08	100,00	0,00	8,07	100,00	0,00	9,19	100,00	0,00	11,47
09	100,00	0,00	102,79	100,00	0,00	108,71	100,00	0,00	24,21
10	93,75	2,88	1049,76	93,75	2,59	1019,47	100,00	0,00	57,82
11	43,75	32,42	2377,33	43,75	33,80	2327,45	100,00	0,00	128,14
12	12,50	59,79	3322,83	6,25	60,41	3377,30	100,00	0,00	278,23
13	18,75	66,27	2988,67	18,75	66,40	2999,26	87,50	0,58 ¹	1246,11
14	12,50	68,50	3198,65	12,50	69,40	3168,35	37,50	5,36 ²	2758,40
Média	64,58	25,54	1449,90	63,89	25,85	1445,63	91,67	0,66	501,63

¹ gap médio relativo aos 93,75% problemas-teste em que uma solução viável foi encontrada.

² gap médio relativo aos 50,00% problemas-teste em que uma solução viável foi encontrada.

Pela Tabela 5.3, observa-se que, novamente, o MPLIM-G e o MPLIM-BG apresentaram comportamentos bem semelhantes. Usando-se esses dois modelos, o CPLEX não foi capaz de resolver na otimalidade todos os problemas-teste com mais de 9 tarefas, para os quais o $\overline{gap}$ variou de 2,59% à 60,40%. Por outro lado, com o modelo MPLIM-IT, o CPLEX conseguiu resolver na otimalidade todos os problemas-teste com até 12 tarefas. Esta última formulação também exigiu, na maioria dos casos, um menor tempo médio computacional para se alcançar a otimalidade. De fato, com exceção dos três primeiros conjuntos de problemas-teste, a formulação MPLIM-IT demandou até 18 vezes, aproximadamente, menos tempo que as demais formulações.

Ainda na Tabela 5.3, nota-se que o CPLEX aplicado à formulação MPLIM-IT não foi capaz de encontrar, em uma hora de processamento, uma solução viável em 6,25% dos problemas-teste com 13 tarefas, bem como em 50,00% dos problemas-teste com 14 tarefas. Contudo, para os casos em que essa formulação conseguiu gerar uma solução viável, os $\overline{gap}$ ’s foram bem baixos (0,58% em problemas com 13 tarefas e 5,36% nos de 14 tarefas), se comparados com aqueles produzidos pelas formulações MPLIM-G e MPLIM-BG, que foram sempre superiores a 66%.

5.3 Resultados Obtidos com os Algoritmos GPV e GPV-REB

Os algoritmos heurísticos propostos (GPV e GPV-REB), bem como os procedimentos utilizados para a definição dos seus parâmetros, foram implementados em linguagem C,

usando-se o compilador Borland C++ Builder 5.0. Foram realizados experimentos com todos os problemas-teste das três bases de dados apresentadas na Seção 5.1, totalizando três baterias de testes.

Para auxiliar na apresentação e comparação dos resultados, utilizou-se duas medidas de desempenho, dadas pelas equações (5.4) e (5.5). Nestas equações, para um dado problema-teste i , f_i^* é o valor da função objetivo da melhor solução conhecida para o problema e $f_i^{X^*}$ e $\bar{f}_i^X$ são, respectivamente, o valor da função objetivo da melhor solução encontrada e a média dos valores da função objetivo aplicada nas soluções encontradas durante todas as execuções do algoritmo X para o problema-teste i .

$$imp_i^{best} = \frac{f_i^{X^*} - f_i^*}{f_i^*} \quad (5.4)$$

$$imp_i^{avg} = \frac{\bar{f}_i^X - f_i^*}{f_i^*} \quad (5.5)$$

A primeira equação calcula o desvio do valor da função objetivo da melhor solução encontrada pelo algoritmo X em torno da função objetivo da melhor solução conhecida para o respectivo problema-teste. Já a segunda equação calcula a variabilidade média dos valores da função objetivo das soluções encontradas pelo algoritmo X em torno da melhor solução conhecida para o respectivo problema.

A seguir são apresentados os procedimentos utilizados para definir cada parâmetro do GPV (utilizou-se os mesmos parâmetros no GPV-REB) e os resultados obtidos em cada bateria de testes.

5.3.1 Definição dos Parâmetros

Cada parâmetro do algoritmo GPV foi definido por meio de experimentos em um problema-teste de porte médio (30 tarefas). Definido um parâmetro, este era fixado nos procedimentos de definição dos demais parâmetros.

O primeiro parâmetro a ser definido foi o parâmetro γ da fase de construção de uma solução (ver Seção 4.5). Para isto, dado um valor para γ , foi construída uma solução com cada uma das seis funções adaptativas gulosas definidas na Seção 4.5 e a melhor destas soluções era retornada. Como o objetivo era determinar o valor de γ que proporciona construir soluções de melhor qualidade, não foi utilizado o POP durante as construções das soluções, ou seja, as soluções parciais não foram submetidas à uma busca local. Este procedimento foi realizado 20 vezes para cada um dos seguintes valores para γ : 0; 0,02; 0,04; 0,06; 0,08; 0,10; 0,12; 0,14; 0,16; 0,18; 0,20; 0,22; 0,24; 0,26; 0,28 e 0,30. Os resultados encontrados estão sintetizados no gráfico da Figura 5.1, onde é mostrado o menor valor e a média dos valores retornados para a função objetivo (fo) das soluções produzidas nas 20 execuções dos respectivos valores para γ .

Como pode ser observado no gráfico da Figura 5.1, o menor valor para fo foi encontrado para γ igual a 0,12. Já a menor média para fo foi encontrada para γ igual a 0. Desde modo, definiu-se γ igual a 0 na primeira vez em que cada uma das seis funções adaptativas gulosas é utilizada e γ igual a 0,12 nas demais.

O segundo parâmetro a ser definido foi o parâmetro $MRDmax$ da descida randômica (MRD) utilizada no POP (ver Seção 4.5). Para isto, dado um valor para $MRDmax$, foi

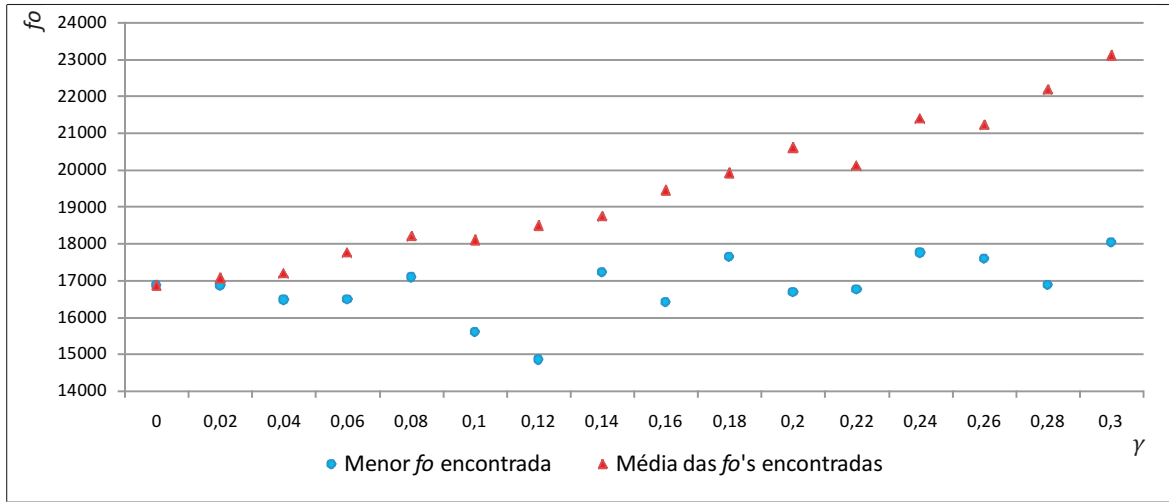


Figura 5.1: Resultados dos testes para definição de γ .

construída uma solução com cada uma das seis funções adaptativas gulosas definidas na Seção 4.5 e a melhor destas soluções era retornada. Como o objetivo era determinar o valor de $MRDmax$ que proporciona encontrar soluções de melhor qualidade a partir de uma mesma solução inicial, o valor de γ foi fixado em 0. Deste modo, a ordem de inclusão de uma tarefa na solução parcial é sempre a mesma para uma dada função adaptativa gulosa. Este procedimento foi realizado 20 vezes para cada valor de $MRDmax$. Os valores de $MRDmax$ testados são da forma $k \times n$, onde n representa o número de tarefas a serem sequenciadas e $k \in \{1, 2, 3, \dots, 15\}$. Os resultados encontrados estão apresentados no gráfico da Figura 5.2, onde é mostrado o menor valor e a média dos valores retornados para a função objetivo (f_0) das soluções encontradas nas 20 execuções dos respectivos valores de k .

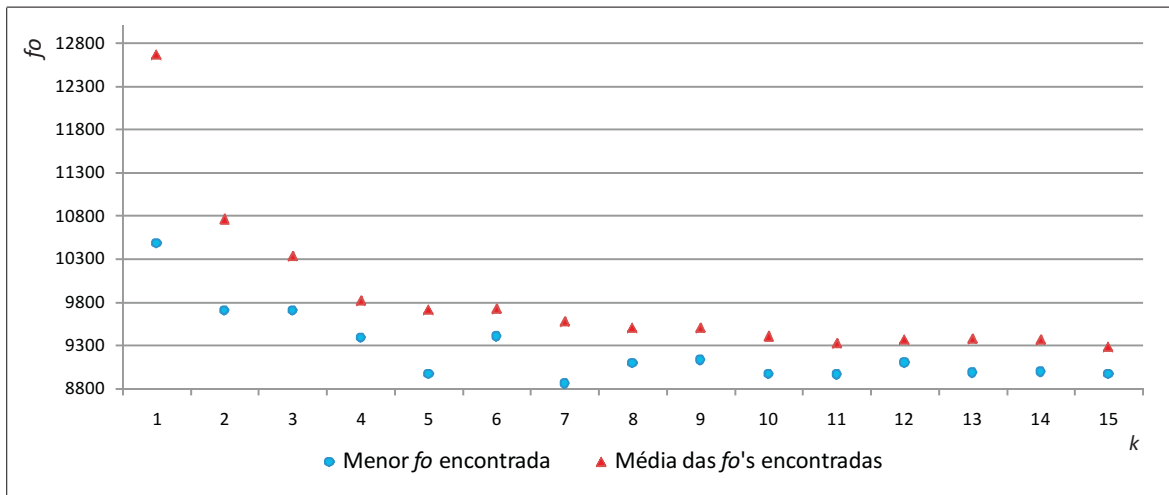


Figura 5.2: Resultados dos testes para definição de $MRDmax$.

Como pode ser observado no gráfico da Figura 5.2, o menor valor para f_0 foi encontrado

para k igual a 7, ou seja, para $MRDmax$ igual a $7 \times n$. Deste modo, definiu-se $MRDmax$ igual a $7 \times n$. Devido ao tempo computacional despendido nos problemas com 75 ou mais tarefas ($n \geq 75$), utilizou-se $MRDmax$ igual a $5 \times n$ para estes problemas.

Definidos os parâmetros da fase de construção de uma solução, o próximo parâmetro a ser definido foi o parâmetro $VNDmax$ do VND_1 (ver Seção 4.6). Para isto, dado um valor para $VNDmax$, foi construída uma solução com cada uma das seis funções adaptativas gulosas definidas na Seção 4.5 e, a seguir, cada solução construída foi submetida ao VND_1 . A melhor solução encontrada era retornada. Para construir as soluções, γ foi fixado em 0 e o $MRDmax$ foi $7 \times n$. Este procedimento foi realizado 20 vezes para cada valor de $VNDmax$. Os valores de $VNDmax$ testados também são da forma $k \times n$, onde n representa o número de tarefas a serem sequenciadas e $k \in \{1, 2, 3, \dots, 15\}$. Os resultados encontrados estão resumidos no gráfico da Figura 5.2, onde é mostrado o menor valor e a média dos valores retornados para a função objetivo (fo) das soluções encontradas nas 20 execuções dos respectivos valores de k .

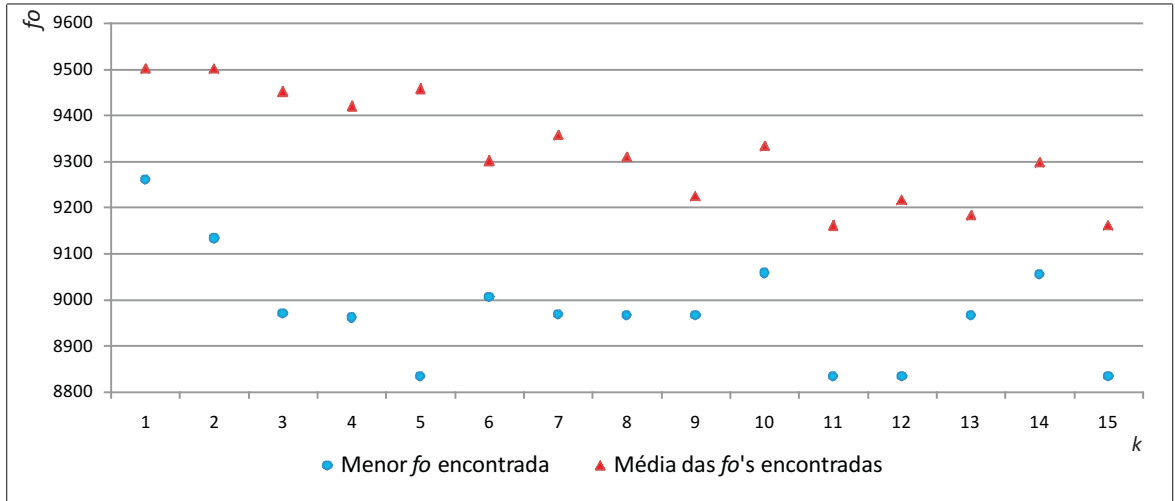


Figura 5.3: Resultados dos testes para definição de $VNDmax$.

Conforme o gráfico da Figura 5.3, o menor valor para fo foi encontrado para k igual a 5, 11, 12 e 15, ou seja, para $VNDmax$ igual a $5 \times n$, $11 \times n$, $12 \times n$ e $15 \times n$. Como a menor média para a fo foi encontrada para $VNDmax$ igual a $11 \times n$, definiu-se $VNDmax$ igual a $11 \times n$. Devido ao tempo computacional despendido para problemas com 75 ou mais tarefas ($n \geq 75$), nestes problemas utilizou-se $VNDmax$ igual a $5 \times n$.

Finalmente, o último parâmetro a ser definido foi o parâmetro $GRASPmax$ ¹ da primeira fase da metodologia proposta (ver Algoritmo 10). Para isto, dado um valor para $GRASPmax$, apenas a primeira fase do algoritmo GPV era aplicada ao problema e a melhor solução encontrada era retornada. Para construir as soluções, utilizou-se γ igual a 0 na primeira vez em que cada uma das seis funções adaptativas gulosas era utilizada e γ igual a 0,12 nas demais. $MRDmax$ e $VNDmax$ foram fixados em $7 \times n$ e $11 \times n$, respectivamente. Este procedimento foi realizado 20 vezes para cada valor de $GRASPmax$.

¹É importante lembrar que esta fase é baseada no procedimento GRASP, mas, diferentemente do GRASP tradicional, utiliza-se $GRASPmax$ iterações sem melhora na solução corrente e não simplesmente $GRASPmax$ iterações.

Os valores de $GRASP_{max}$ testados foram: 6, 12, 18, 24 e 30. Testou-se múltiplos de 6 para $GRASP_{max}$ por terem sido utilizadas seis funções adaptativas gulosas distintas para construir soluções. Os resultados encontrados estão sintetizados no gráfico da Figura 5.2, onde é mostrado o menor valor e a média dos valores retornados para a função objetivo (fo) das soluções encontradas durante as 20 execuções dos respectivos valores para $GRASP_{max}$.

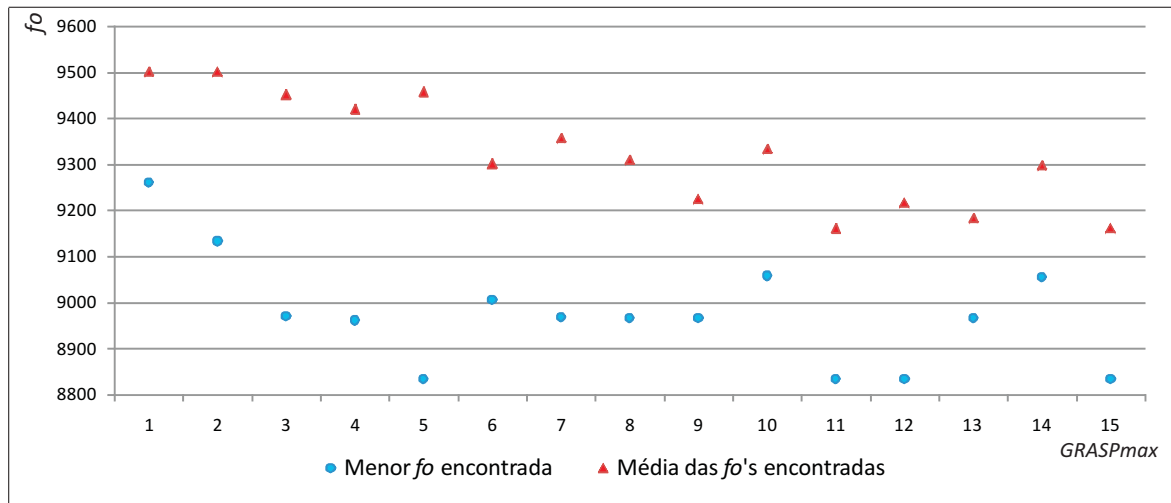


Figura 5.4: Resultados dos testes para definição de $GRASP_{max}$.

De acordo o gráfico da Figura 5.4, o menor valor para fo foi encontrado para $GRASP_{max}$ igual a 12, 18, 24 e 36. Como a menor média para a fo foi encontrada para $GRASP_{max}$ igual a 24, definiu-se $GRASP_{max}$ igual a 24. Devido ao tempo computacional despendido nos problemas com 75 ou mais tarefas ($n \geq 75$), utilizou-se $GRASP_{max}$ igual a 12 para estes problemas.

Os parâmetros utilizados estão resumidos na Tabela 5.4, onde n representa o número de tarefas a serem sequenciadas.

Tabela 5.4: Parâmetros utilizados.

Parâmetros	Valores
γ	0 e 0,12
MRD_{max}	$7 \times n$, se $n \leq 50$ e $5 \times n$, se $n \geq 75$
VND_{max}	$11 \times n$, se $n \leq 50$ e $5 \times n$, se $n \geq 75$
$GRASP_{max}$	24, se $n \leq 50$ e 12, se $n \geq 75$

5.3.2 Primeira Bateria de Testes - BDS

A primeira bateria de experimentos foi realizada com os problemas-teste da base de dados BDS. Os problemas-teste de tal base de dados possuem os tempos de *setup* simétricos. Cada problema-teste foi resolvido 30 vezes com cada algoritmo proposto, exceto aqueles com 75 tarefas, os quais foram resolvidos apenas 10 vezes em virtude do maior tempo computacional demandado. Os resultados obtidos pelo algoritmo GPV foram comparados com os alcançados em Gomes Júnior (2007), Penna (2009) e Ribeiro (2009), visto que estes autores também realizaram experimentos computacionais com esta base de dados.

Um resumo dos resultados obtidos pelo algoritmo GPV e pelos algoritmos da literatura que também fizeram experimentos com a base de dados em questão é apresentado na Tabela 5.5. Nesta tabela, a primeira coluna indica o número de tarefas nos problemas-teste de cada grupo. Nas colunas “Gomes Júnior (2007)”, “Penna (2009)”, “Ribeiro (2009)” e “GPV” são apresentados os resultados obtidos pelos respectivos algoritmos. Para um grupo de problemas-teste com i tarefas, as colunas “Tempo” mostram as médias dos tempos demandados (em segundos) pelos respectivos algoritmos, as colunas “ $\overline{imp}^{best}$ ” apresentam as médias dos imp_i^{best} s (em porcentagem) obtidos pelos respectivos algoritmos aplicados nos problemas deste conjunto e as colunas “ $\overline{imp}^{avg}$ ” apresentam as médias dos imp_i^{avg} s obtidos pelos respectivos algoritmos aplicados nos problemas-teste do conjunto.

Pela Tabela 5.5, observa-se que o algoritmo GPV e o algoritmo proposto por Gomes Júnior (2007) foram capazes de encontrar as melhores soluções conhecidas em todos os problemas-teste com até 25 tarefas. Além disso, para estes mesmos problemas, os $\overline{imp}^{avg}$ s obtidos pelo GPV foram sempre menores que os obtidos por tal autor, sendo o maior $\overline{imp}^{avg}$ apresentado pelo GPV igual a 1,53% e o maior obtido pelo algoritmo de Gomes Júnior (2007) igual a 2,32%. Já para os problemas-teste com mais de 25 tarefas, os $\overline{imp}^{best}$ s obtidos pelos dois algoritmos alternaram os menores valores e os $\overline{imp}^{avg}$ s apresentados pelo algoritmo de Gomes Júnior (2007) foram sempre superiores que os respectivos $\overline{imp}^{avg}$ s apresentados pelo GPV, sendo que, para os problemas com 75 tarefas, o $\overline{imp}^{avg}$ obtido pelo GPV foi de 13,49%, contra os 19,22% obtido pelo algoritmo deste autor. Além disso, para os problemas com mais de 30 tarefas os tempos médios demandados pelo algoritmo de Gomes Júnior (2007) foram de 30% a 74% maiores que os demandados pelo GPV, apesar dos testes terem sido realizados em máquinas diferentes.

Ainda na Tabela 5.5 é possível observar que os algoritmos propostos por Penna (2009) e Ribeiro (2009) obtiveram os menores valores para $\overline{imp}^{best}$ e $\overline{imp}^{avg}$ na maioria dos casos. Porém, para os problemas com mais de 20 tarefas, número de tarefas a partir do qual o tempo demandado passa a ser significativo, os dois algoritmos demandaram tempos médios superiores aos demandados pelo GPV nos respectivos conjuntos de problemas. Apesar de terem sido utilizadas máquinas diferentes, os tempos médios demandados pelo algoritmo de Penna (2009) e Ribeiro (2009) em tais conjuntos de problemas chegaram a ser aproximadamente 800% e 360%, respectivamente, superiores aos demandados pelo algoritmo GPV nos respectivos conjuntos de problemas.

Tabela 5.5: Comparação GPV $\times$ Gomes Júnior (2007) $\times$ Penna (2009) $\times$ Ribeiro (2009) aplicados à BDS.

# Tarefas	Gomes Júnior (2007)			Penna (2009)			Ribeiro (2009)			GPV		
	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ¹ (s)	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ² (s)	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ³ (s)	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ⁴ (s)
08	0,00	0,03	0,04	0,00	0,00	0,06	0,00	0,00	0,94	0,00	0,00	0,11
09	0,00	0,06	0,07	0,00	0,00	0,09	0,00	0,15	1,26	0,00	0,00	0,17
10	0,00	0,02	0,11	0,00	0,00	0,15	0,00	0,24	1,60	0,00	0,00	0,24
11	0,00	0,12	0,20	0,00	0,00	0,25	0,00	0,03	2,21	0,00	0,06	0,46
12	0,00	0,21	0,29	0,00	0,00	0,37	0,00	0,07	2,81	0,00	0,02	0,74
15	0,00	1,47	0,94	0,00	0,47	1,13	0,00	0,76	6,02	0,00	0,51	1,47
20	0,00	1,65	4,35	0,00	0,64	4,93	0,00	0,73	20,60	0,00	1,49	5,19
25	0,00	2,32	13,29	0,00	1,09	14,90	0,00	1,02	45,72	0,00	1,53	14,22
30	0,20	3,34	40,07	0,00	1,68	39,93	0,00	1,60	112,06	0,04	2,41	30,84
40	0,44	4,66	155,79	0,08	3,32	190,61	0,00	2,41	335,88	1,10	3,89	111,52
50	1,37	7,29	492,28	0,32	4,95	630,77	0,00	4,08	896,10	1,65	6,48	309,03
75	7,89	19,22	1368,08	0,82	7,46	6308,74	0,99	7,61	2005,05	6,08	13,49	787,90
Média	0,82	3,37	172,96	0,10	1,63	599,33	0,08	1,56	285,85	0,74	2,49	105,16

¹Testes realizados em um PC Athlon XP 64 Bits 3000+ (aproximadamente 2 GHz), com 1 GB de RAM.

²Testes realizados em um PC Intel Core 2 Quad (Q6600) 2,40 GHz, com 4 GB de RAM.

³Testes realizados em um PC Pentium Core 2 Duo 2,1 GHz, com 4 GB de memória RAM.

⁴Testes realizados em um PC Intel Core 2 Duo (E8400) 2.99 GHz, com 2 GB de memória RAM.

Os resultados encontrados nos testes computacionais realizados com o algoritmo GPV-REB estão sintetizados na Tabela 5.6. Na primeira coluna é indicado o número de tarefas nos problema-teste de cada conjunto. Nas colunas “GPV” e “GPV-REB” são apresentados os resultados obtidos pelos respectivos algoritmos. Para um grupo de problemas-teste com i tarefas, as colunas “Tempo”, “ $\overline{imp}^{best}$ ” e “ $\overline{imp}^{avg}$ ” são tais como na Tabela 5.5.

Conforme pode ser observado na Tabela 5.6, para os conjuntos de problemas-teste com até 50 tarefas, o algoritmo GPV-REB obteve valores de $\overline{imp}^{best}$ menores ou iguais a 2,13%. De fato, em dois destes conjuntos o $\overline{imp}^{best}$ foi próximo de 1,60% e nos demais, menores que 1,00%. Para os problemas-teste com 75 tarefas o $\overline{imp}^{best}$ foi de 8,63%, valor próximo daquele obtido com o algoritmo GPV. Já os $\overline{imp}^{avg}$ s obtidos por meio do GPV-REB foram sempre superiores aos apresentados pelo GPV nos conjuntos correspondentes. No entanto, os valores dos $\overline{imp}^{avg}$ s obtido pelo GPV-REB aproximam-se dos obtidos com o GPV à medida em que considera-se conjuntos de problemas com mais tarefas. Finalmente, para os problemas em que o tempo demandado é mais significativo, isto é, problemas com mais de 20 tarefas, os tempos demandados pelo GPV foram de 60% a 200%, aproximadamente, superiores aos demandados pelo GPV-REB nos conjuntos de problemas correspondentes.

Tabela 5.6: Comparação GPV $\times$ GPV-REB aplicados à BDS.

# Tarefas	GPV			GPV-REB		
	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo (s)	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo (s)
08	0,00	0,00	0,11	0,23	0,23	0,06
09	0,00	0,00	0,17	0,00	0,00	0,13
10	0,00	0,00	0,24	0,00	0,04	0,13
11	0,00	0,06	0,46	0,15	0,19	0,37
12	0,00	0,02	0,74	0,00	0,98	0,52
15	0,00	0,51	1,47	1,48	2,53	0,90
20	0,00	1,49	5,19	0,04	3,77	3,15
25	0,00	1,53	14,22	1,57	3,86	6,96
30	0,04	2,41	30,84	0,72	4,55	19,26
40	1,10	3,89	111,52	1,77	6,73	60,36
50	1,65	6,48	309,03	2,13	8,70	151,74
75	6,08	13,49	787,90	8,63	16,34	467,71
Média	0,74	2,49	105,16	1,39	3,99	59,27

5.3.3 Segunda Bateria de Testes - BDNS

Na segunda bateria de testes utilizou-se os problemas da base de dados BDNS. Os problemas-teste dessa base de dados não possuem, necessariamente, os tempos de *setup* simétricos e satisfazendo a desigualdade triangular, sendo assim, mais genéricos. Assim como na primeira bateria de testes, cada problema foi resolvido 30 vezes, exceto aqueles com mais de 50 tarefas, os quais foram resolvidos apenas 10 vezes devido ao maior tempo computacional demandado. Os resultados obtidos pelos algoritmos GPV e GPV-REB

foram comparados com os resultados do algoritmo proposto por Ribeiro (2009), dado que este é o único autor da literatura que realizou experimentos com esta base de dados.

Uma síntese dos resultados encontrados pelos algoritmos GPV e GPV-REB e pelo algoritmo de Ribeiro (2009) é apresentada na Tabela 5.7. Nela, a primeira coluna indica o número de tarefas dos problemas de cada grupo de problemas-teste. Nas colunas “Ribeiro (2009)”, “GPV” e “GPV-REB” são apresentados os resultados obtidos pelos respectivos algoritmos. Para um grupo de problemas-teste com i tarefas, as colunas “Tempo” mostram as médias dos tempos demandados (em segundos) pelos respectivos algoritmos, as colunas “ $\overline{imp}^{best}$ ” apresentam as médias dos imp_i^{best} ’s (em porcentagem) obtidos pelos respectivos algoritmos aplicados nos problemas deste grupo e as colunas “ $\overline{imp}^{avg}$ ” apresentam as médias dos imp_i^{avg} ’s obtidos pelos respectivos algoritmos aplicados nos problemas do grupo.

Tabela 5.7: Comparação GPV $\times$ GPV-REB $\times$ Ribeiro (2009) aplicados à BDNS.

# Tarefas	Ribeiro (2009)			GPV			GPV-REB		
	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ¹ (s)	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ² (s)	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ² (s)
006	0,00	0,00	0,76	0,00	0,00	0,04	0,00	0,00	0,03
007	0,00	0,00	0,76	0,00	0,00	0,07	0,00	0,00	0,03
008	0,00	0,00	1,03	0,00	0,00	0,30	0,00	0,76	0,25
009	0,00	0,00	1,35	0,00	0,00	0,37	0,01	0,18	0,29
010	0,00	0,00	1,61	0,00	0,20	0,45	0,47	0,86	0,34
011	0,00	0,00	2,20	0,00	0,01	0,40	0,11	0,16	0,40
012	0,00	0,00	3,32	0,00	0,00	0,52	0,00	0,43	0,42
015	0,00	0,01	7,49	0,00	0,01	1,26	0,00	0,54	0,73
020	0,00	0,30	23,88	0,00	0,10	4,52	0,49	1,83	2,85
030	0,19	1,40	172,51	0,01	0,58	36,06	0,51	1,99	18,17
040	0,18	1,17	801,67	0,06	0,79	132,94	0,36	1,51	67,04
050	0,30	1,45	1575,11	0,06	1,28	440,49	0,65	2,57	215,60
075	5,90	5,90	4978,02	3,10	5,64	1187,76	4,14	6,66	636,86
100	2,36	2,62	22299,79	0,21	2,32	4612,79	0,57	3,79	3015,16
150	-	-	-	0,14	2,53	49524,76	0,83	2,62	32833,61
Média	0,64	0,92	2133,54	0,24	0,90	3729,52	0,54	1,59	2452,79

¹Testes realizados em um PC Pentium Core 2 Duo 2,1 GHz, com 4 GB de memória RAM.

²Testes realizados em um PC Intel Core 2 Duo (E8400) 2.99 GHz, com 2 GB de memória RAM.

De acordo com a Tabela 5.7, o GPV obteve valores de $\overline{imp}^{best}$ sempre menores ou iguais que os respectivos valores obtidos pelo algoritmo de Ribeiro (2009). Isso significa que as soluções produzidas pelo GPV foram, em média, melhores ou iguais às produzidas pelo algoritmo de tal autor. Ambos os algoritmos encontraram $\overline{imp}^{avg}$ ’s próximos de zero nos conjuntos de problemas-teste com até 20 tarefas. Já nos conjuntos de problemas com mais de 20 tarefas, os valores dos $\overline{imp}^{avg}$ ’s obtidos pelo GPV foram relativamente inferiores aos respectivos valores encontrados pelo algoritmo de Ribeiro (2009). Ademais, os tempos médios demandados pelo algoritmo de Ribeiro (2009) para resolver os problemas de cada

conjunto de problemas-teste foram sempre significativamente maiores que os demandados pelo GPV nos problemas dos respectivos conjuntos. Para o conjunto de problemas-teste com 100 tarefas, por exemplo, o tempo médio demandado pelo algoritmo proposto por Ribeiro (2009) foi aproximadamente 4,8 vezes superior ao demandado pelo GPV.

Na Tabela 5.7 também é possível observar que o algoritmo GPV obteve valores de $\overline{imp}^{best}$ menores que 1,00% para todos os conjuntos de problemas-teste, exceto para o conjunto de problemas com 75 tarefas. Mesmo assim, para este último conjunto, o $\overline{imp}^{best}$ foi menor que o obtido pelo algoritmo de Ribeiro (2009). Os $\overline{imp}^{avg}$ s foram relativamente baixos, sendo o maior valor igual a 6,66% e, para os problemas com 150 tarefas, próximo do valor obtido pelo GPV. Os tempos médios demandados também foram inferiores os demandados pelos outros dois algoritmos. Se comparado aos tempos médios demandados pelo GPV, os demandados pelo GPV-REB foram de 34% a 50% menores.

5.3.4 Terceira Bateria de Testes - BDDT

A terceira bateria de experimentos foi realizada com os problemas-teste da base de dados BDDT. Os tempos de processamento e de *setup* dos problemas-teste de tal base de dados satisfazem a desigualdade triangular. Cada problema-teste foi resolvido 30 vezes. Novamente os resultados obtidos pelos algoritmos GPV e GPV-REB foram comparados com os resultados de Ribeiro (2009), pois este último é o único autor da literatura que realizou experimentos com a base de dados em questão.

Os resultados obtidos pelos algoritmos GPV e GPV-REB e por Ribeiro (2009) estão resumidos na Tabela 5.8, onde a primeira coluna indica o número de tarefas nos problemas-teste de cada conjunto de problemas. Nas colunas “Ribeiro (2009)”, “GPV” e “GPV-REB” são apresentados os resultados obtidos pelos respectivos algoritmos. Para um grupo de problemas-teste com i tarefas, as colunas “Tempo” mostram as médias dos tempos demandados (em segundos) pelos respectivos algoritmos, as colunas “ $\overline{imp}^{best}$ ” apresentam as médias dos imp_i^{best} s (em porcentagem) obtidos pelos respectivos algoritmos aplicados nos problemas deste grupo e as colunas “ $\overline{imp}^{avg}$ ” apresentam as médias dos imp_i^{avg} s obtidos pelos respectivos algoritmos aplicados nos problemas-teste do grupo.

Pela Tabela 5.8, observa-se que o algoritmo GPV, bem como o algoritmo proposto por Ribeiro (2009), foi capaz de encontrar todas as respectivas melhores soluções conhecidas para os problemas-teste da base de dados BDDT. Apesar de o mesmo não ter ocorrido com o algoritmo GPV-REB, o maior $\overline{imp}^{best}$ obtido com este algoritmo foi de 0,35%. Os três algoritmos obtiveram valores de $\overline{imp}^{avg}$ menores que 1,00%. Em relação aos tempos médios demandados em cada conjunto de problemas-teste, novamente o algoritmo de Ribeiro (2009) demandou tempos médios maiores que os respectivos tempos médios demandados pelo algoritmo GPV, que por sua vez demandou tempos médios maiores que os despendidos pelo algoritmo GPV-REB. Tais tempos médios foram, no máximo, 3,43 segundos para o algoritmo GPV-REB, 5,26 segundos para o algoritmo GPV e 30,32 segundos para o algoritmo de Ribeiro (2009).

Dado um determinado problema-teste da base de dados BDDT, a melhor solução produzida para ele foi utilizada para determinar o conjunto $H = \{h_{inf}, h_{inf} + 1, h_{inf} + 2, h_{sup}\}$ necessário na formulação de programação matemática MPLIM-IT (ver Seção 3.3). Os valores de h_{inf} e de h_{sup} utilizados são dados, respectivamente, pela data de início da primeira tarefa e pela data de conclusão da última tarefa na melhor sequência encontrada

Tabela 5.8: Comparação GPV $\times$ GPV-REB $\times$ Ribeiro (2009) aplicados à BDDT.

# Tarefas	Ribeiro (2009)			GPV			GPV-REB		
	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ¹ (s)	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ² (s)	$\overline{imp}^{best}$ (%)	$\overline{imp}^{avg}$ (%)	Tempo ² (s)
06	0,00	0,00	0,67	0,00	0,00	0,04	0,00	0,00	0,02
07	0,00	0,00	0,86	0,00	0,00	0,17	0,00	0,06	0,04
08	0,00	0,00	1,13	0,00	0,00	0,29	0,00	0,27	0,06
09	0,00	0,00	1,47	0,00	0,00	0,34	0,00	0,02	0,09
10	0,00	0,00	1,78	0,00	0,00	0,41	0,00	0,03	0,15
11	0,00	0,00	2,43	0,00	0,04	0,32	0,00	0,15	0,19
12	0,00	0,00	3,88	0,00	0,00	0,49	0,01	0,09	0,31
13	0,00	0,00	4,87	0,00	0,01	0,69	0,00	0,25	0,44
14	0,00	0,00	5,75	0,00	0,04	0,91	0,00	0,11	0,56
15	0,00	0,02	8,50	0,00	0,04	1,25	0,00	0,19	0,71
16	0,00	0,06	11,64	0,00	0,14	2,36	0,04	0,61	1,20
17	0,00	0,04	1,47	0,00	0,04	2,34	0,09	0,36	1,40
18	0,00	0,42	15,23	0,00	0,05	3,03	0,05	0,37	1,95
19	0,00	0,04	22,16	0,00	0,01	3,80	0,35	0,89	2,59
20	0,00	0,02	30,32	0,00	0,14	5,26	0,10	0,88	3,43
Média	0,00	0,04	7,48	0,00	0,03	1,45	0,04	0,29	0,88

¹Testes realizados em um PC Pentium Core 2 Duo 2,1 GHz, com 4 GB de memória RAM.

²Testes realizados em um PC Intel Core 2 Duo (E8400) 2.99 GHz, com 2 GB de memória RAM.

pelo algoritmo GPV para o problema. Este procedimento foi utilizado para determinar o conjunto H de cada problema-teste da base de dados e, em seguida, utilizou-se o CPLEX para resolver o MPLIM-IT relativo à cada problema.

Considerou-se o limite de 3600 segundos para a resolução de cada problema-teste pelo CPLEX e a qualidade das soluções retornadas foram medidas via a expressão (5.3) apresentada na Seção 5.2.

Os resultados encontrados são resumidos na Tabela 5.9. Nela, a primeira coluna indica o número de tarefas nos problemas de cada conjunto de problemas-teste. Para cada conjunto de problemas, a coluna “Soluções Ótimas” mostra (em porcentagem) em quantos problemas o CPLEX encontrou uma solução ótima com o modelo MPLIM-IT, a coluna “ $\overline{gap}$ ” mostra as médias dos gap ’s e a coluna “Tempo Médio” mostra as médias dos tempos demandados (em segundos) por tal otimizador na resolução dos problemas-teste com o MPLIM-IT.

Segundo a Tabela 5.9, o modelo MPLIM-IT associado ao algoritmo GPV permitiu ao CPLEX resolver na otimalidade todos os problemas-teste com até 13 tarefas. Nos conjuntos de problemas-teste com 14 e 15 tarefas, o maior $\overline{gap}$ encontrado foi igual a 1,61%.

Ainda na Tabela 5.9, observa-se que o CPLEX aplicado à formulação MPLIM-IT associada ao algoritmo GPV não foi capaz de encontrar, em uma hora de processamento, uma solução viável em 12,50% dos conjuntos de problemas-teste com 16 e 17 tarefas. Contudo, para os casos em que essa formulação conseguiu gerar uma solução viável, os

Tabela 5.9: Resultados do MPLIM-IT associado ao GPV.

# Tarefas	MPLIM-IT		
	Soluções Ótimas (%)	$\overline{gap}$ (%)	Tempo Médio (s)
06	100,00	0,00	1,93
07	100,00	0,00	5,06
08	100,00	0,00	10,44
09	100,00	0,00	16,99
10	100,00	0,00	41,70
11	100,00	0,00	105,65
12	100,00	0,00	297,89
13	100,00	0,00	795,70
14	87,50	0,50	944,55
15	68,75	1,61	1593,53
16	50,00	1,92 ¹	2563,40
17	37,50	5,86 ¹	2911,82
Média	86,98	0,82	774,05

¹ $\overline{gap}$ médio relativo aos 87,50% dos problemas-teste em que uma solução viável foi encontrada.

$\overline{gap}$'s também foram baixos (1,92% nos problemas com 16 tarefas e 5,86% nos de 17 tarefas).

5.4 Teste de Probabilidade Empírica

O comportamento do algoritmo GPV também foi comparado com os algoritmos propostos por Gomes Júnior (2007), Penna (2009) e Ribeiro (2009) via um teste de probabilidade empírica proposto por Aiex *et al.* (2000). Para realizar tal teste, escolheu-se um problema com 40 tarefas da base de dados BDS. Em seguida, determinou-se o valor da função objetivo da melhor solução conhecida para o problema como sendo o alvo a ser atingido (portanto, um alvo difícil). Cada um destes algoritmos foi executado 100 vezes no problema escolhido e cada execução foi interrompida somente quando o alvo foi atingido. Neste último momento, registrou-se o tempo despendido. Para cada algoritmo, os tempos iguais foram eliminados e o restante foi ordenado de forma crescente. Ao tempo da posição i , t_i , associou-se a probabilidade $p_i = \frac{i - 0,5}{N}$, onde N é o número de tempos não repetidos.

Os gráficos com os pontos (t_i, p_i) relativos aos algoritmos analisados são apresentados na Figura 5.5. Segundo este gráfico, os algoritmos propostos por Gomes Júnior (2007), Penna (2009) e o GPV apresentam comportamentos semelhantes para o problema-teste em questão, com vantagem dos dois primeiros, que requerem um menor tempo de processamento para obter um dado nível de probabilidade de alcançar o valor alvo. Com os algoritmos GPV e o proposto por Penna (2009), a partir de 140 segundos, aproximadamente, tem-se mais de 90% de probabilidade de atingir o alvo para o problema-teste escolhido. Já o algoritmo de Gomes Júnior (2007) possui tal probabilidade somente a

partir de 190 segundos, aproximadamente. O algoritmo proposto por Ribeiro (2009) é o mais lento entre os algoritmos analisados, pois, para qualquer tempo dado, ele possui a probabilidade de alcançar o alvo menor que os demais algoritmos.

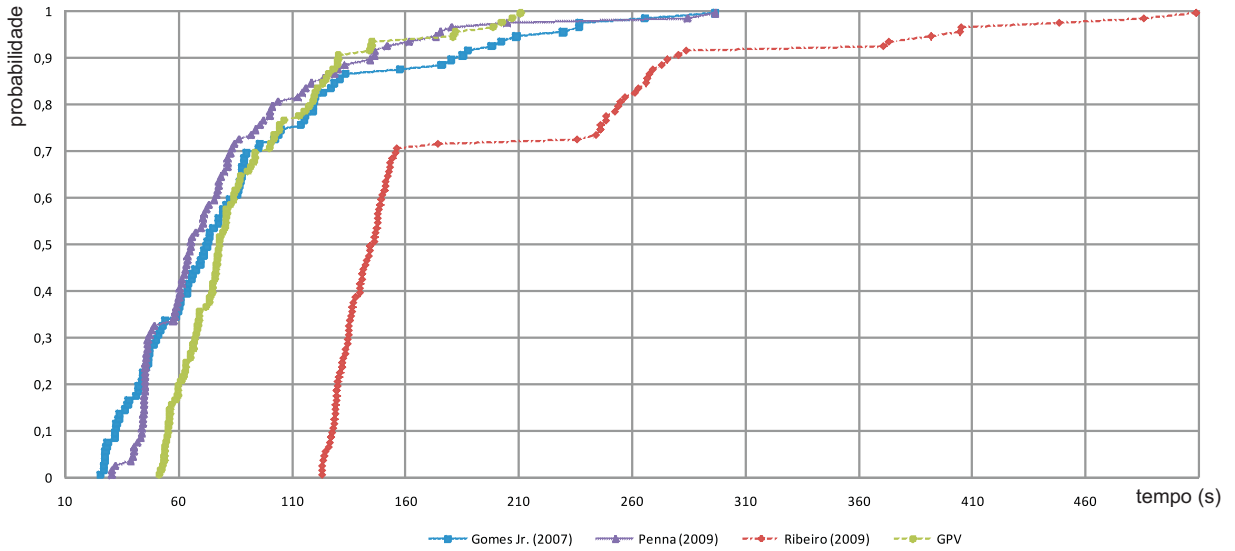


Figura 5.5: Teste de probabilidade empírica.

O teste de probabilidade empírica também foi realizado com outros dois problemas com 30 tarefas da base de dados BDS. Os resultados encontrados foram semelhantes aos apresentados na Figura 5.5.

Capítulo 6

Conclusão

Este trabalho tratou o problema de sequenciamento de tarefas em uma máquina com penalidades por antecipação e atraso da produção, considerando janelas de entrega distintas e tempo de preparação da máquina dependente da sequência, denotado por PSUMAA-JE-TPD.

Foram propostas, inicialmente, duas novas formulações de programação matemática para representar o problema. A primeira delas, denotada por MPLIM-BG, consiste em um aperfeiçoamento da formulação proposta por Gomes Júnior *et al.* (2007) (MPLIM-G), pois exige menos variáveis e restrições que a formulação de tais autores. A segunda formulação proposta, denotada por MPLIM-IT, é indexada no tempo.

Propôs-se, também, um algoritmo heurístico para determinar a melhor sequência de produção. Tal algoritmo, denominado GPV, é composto de duas fases. Na primeira, gera-se uma solução com base na metaheurística GRASP, no Princípio da Otimalidade Próxima (POP) e na Busca em Vizinhança Variável (*Variable Neighborhood Descent* - VND). Na segunda, faz-se o pós-refinamento da solução proveniente da fase anterior por meio de outro VND. Para cada sequência gerada pelo algoritmo, um procedimento de complexidade polinomial encontrado na literatura é acionado para determinar as dadas ótimas de conclusão das tarefas. Com o objetivo de reduzir o custo computacional, adicionalmente é proposta uma estratégia de redução do espaço de busca. Essa estratégia foi utilizada durante a primeira fase do GPV, originando, assim, o algoritmo denominado GPV-REB.

Visto que os problemas-teste da única base de dados (BDS) existente na literatura referente ao problema tratado possuem os tempos de preparação da máquina simétricos, foram geradas outras duas novas bases de dados. Em uma delas, BDDT, os problemas-teste envolvem até 20 tarefas e satisfazem a desigualdade triangular com relação aos tempos de preparação da máquina. Na outra, BDNS, os problemas-teste envolvem até 150 tarefas, e os tempos de preparação da máquina não satisfazem a uma propriedade particular, sendo, portanto, mais genérica.

O otimizador CPLEX 10.1 foi utilizado para resolver as três formulações matemáticas em tela, aplicadas em problemas-teste com até 14 tarefas. Os resultados computacionais mostram que as formulações MPLIM-G e MPLIM-BG proporcionam resultados semelhantes e que a formulação MPLIM-IT é mais eficiente que estas duas na resolução de problemas com até 12 tarefas, dado que esta última proporcionou ao CPLEX encontrar soluções com garantia de melhor qualidade e com menor tempo médio de processamento

que as demais. Para os problemas-teste com mais de 12 tarefas, devido ao limite de tempo estabelecido para o CPLEX retornar uma solução, os resultados obtidos não permitem uma comparação efetiva entre os três modelos.

Os algoritmos GPV e GPV-REB foram aplicados em todos os problemas-teste disponíveis, divididos em três baterias de testes. A primeira bateria foi realizada com a base de dados BDS, que é composta de problemas-teste com até 75 tarefas. Os resultados encontrados pelo GPV foram comparados com os resultados de Gomes Júnior (2007), Penna (2009) e Ribeiro (2009). O GPV apresentou desempenho superior ao algoritmo proposto por Gomes Júnior (2007), tanto em relação à qualidade das soluções encontradas, quanto em relação ao tempo computacional médio demandado. Apesar de os resultados encontrados pelo GPV serem um pouco inferiores aos encontrados pelo algoritmo de Penna (2009) e pelo algoritmo de Ribeiro (2009), os tempos computacionais médios demandados pelo GPV foram significativamente inferiores aos tempos médios demandados pelos algoritmos de tais autores.

Na segunda e na terceira bateria de testes foram utilizadas, respectivamente, as bases de dados BDNS e BDDT. Os resultados encontrados pelo GPV foram comparados com os de Ribeiro (2009), que foi o único autor a utilizar também esta base de dados. No conjunto de problemas-teste BDNS, o GPV mostrou-se mais eficiente que o algoritmo de tal autor, encontrando soluções de melhor qualidade e com tempo computacional bastante inferior. Já com a base de dados BDDT, as soluções encontradas por ambos os algoritmos foram semelhantes, porém o GPV exigiu menor tempo computacional. Ainda na terceira bateria de testes, utilizou-se o GPV para determinar o horizonte de planejamento do MPLIM-IT, o que possibilitou ao CPLEX encontrar soluções bem próximas das ótimas em todos os problemas-teste com até 15 tarefas.

Nas três baterias de testes, os tempos computacionais demandados pelo algoritmo GPV-REB foram sempre significativamente inferiores aos respectivos tempos despendidos pelo GPV. Apesar de as soluções encontradas pelo GPV-REB serem um pouco inferiores às aquelas encontradas pelo GPV, esta diferença diminui à medida em que se aumenta o número de tarefas do problema. Isto sugere que vale a pena investir em técnicas e propriedades que permitam reduzir o espaço de busca.

Enfim, conclui-se que a metodologia proposta é competitiva com aquelas existentes na literatura, tendo o baixo custo computacional como o seu grande aliado. Desta forma, foi possível resolver problemas-teste com até 150 tarefas em tempo computacional viável.

Publicações

A seguir são listados os trabalhos provenientes da presente pesquisa que foram apresentados e aceitos para publicação nos anais dos respectivos eventos.

Título: Princípio da Otimalidade Próxima e GRASP para o problema de sequenciamento de tarefas em uma máquina com penalidades por antecipação e atraso da produção.

Autores: Bruno Ferreira Rosa e Marcone Jamilson Freitas Souza.

Evento: IV Semana de Ciência & Tecnologia do CEFET-MG.

Local: Belo Horizonte/MG - Brasil.

Período: 20 a 26 de outubro de 2008.

Título: Uma nova formulação de programação matemática indexada no tempo para uma classe de problemas de sequenciamento em uma máquina.

Autores: Bruno Ferreira Rosa e Marcone Jamilson Freitas Souza.

Evento: XLI SBPO - Simpósio Brasileiro de Pesquisa Operacional.

Local: Porto Seguro/BA - Brasil.

Período: 1 a 4 de setembro de 2009.

Título: Formulações de programação matemática para o problema de sequenciamento em uma máquina com janelas de entrega distintas e tempo de preparação dependente da sequência de produção.

Autores: Bruno Ferreira Rosa, Marcone Jamilson Freitas Souza e Sérgio Ricardo de Souza.

Evento: XXXII CNMAC - Congresso Nacional de Matemática Aplicada e Computacional.

Local: Cuiabá/MT - Brasil.

Período: 8 a 11 de setembro de 2009.

Título: Princípio da Otimalidade Próxima aplicado ao problema de sequenciamento de tarefas em uma máquina com penalidades por antecipação e atraso da produção.

Autores: Bruno Ferreira Rosa, Marcone Jamilson Freitas Souza e Sérgio Ricardo de Souza.

Evento: IX SBAI - Simpósio Brasileiro de Automação Inteligente.

Local: Brasília/DF - Brasil.

Período: 20 a 23 de setembro de 2009.

Referências Bibliográficas

- Aiex, Renata M., Resende, Mauricio G. C., Ribeiro, Celso C., Celso, & Ribeiro, C. 2000. Probability Distribution Of Solution Time In Grasp: An Experimental Investigation. *Journal of Heuristics*, **8**, 200–2.
- Alidaee, B., & Dragan, I. 1997. A note on minimizing the weighted sum of tardy and early completion penalties in a single machine: A case of small common due date. *European Journal of Operational Research*, **96**, 559–563.
- Allahverdi, A., Gupta, J. N.D., & Aldowaisan, T. 1999. A review of scheduling research involving setup considerations. *Omega: The International Journal of Management Science*, **27**, 219–239.
- Almeida, M. T., & Centeno, M. 1998. A composite heuristic for the single machine early/tardy job scheduling problem. *Computers & Operations Research*, **25**(7/8), 625–635.
- Baker, K. R., & Scudder, G. D. 1990. Sequencing with Earliness and Tardiness Penalties: A Review. *Operations Research*, **38**, 22–36.
- Belfiore, P. P., da Silva, V. F., & Ronconi, D. P. 2006. Scatter Search para programação de tarefas em uma única máquina com penalidades de adiantamento e atraso e data de entrega comum. In: *Anais do XXVI Encontro Nacional de Engenharia de Produção*. Associação Brasileira de Engenharia de Produção, Fortaleza, CE, Brasil.
- Biskup, D., & Feldmann, M. 2001. Benchmarks for scheduling on a single machine against restrictive and unrestrictive common due dates. *Computers & Operations Research*, **28**, 787–801.
- Bustamante, L. M. 2007. *Minimização do custo de antecipação e atraso para o problema de sequenciamento de uma máquina com tempo de preparação dependente da sequência: aplicação em uma usina siderúrgica*. M.Phil. thesis, UFMG, Belo Horizonte.
- Chang, P. C. 1999. A Branch and Bound Approach for Single Machine Scheduling with Earliness and Tardiness Penalties. *Computers & Mathematics with Applications*, **37**, 133–144.
- Christofoletti, L. M. 2002. *Métodos de Reinício Aplicados ao Sequenciamento em Uma Máquina com Tempos de Preparação e Datas de Entrega*. M.Phil. thesis, UNICAMP, Campinas.

- de Paula, Mateus Rocha. 2008. *Heurísticas para a minimização dos atrasos em seqüenciamento de máquinas paralelas com tempos de preparação dependentes da seqüência*. M.Phil. thesis, UFMG, Belo Horizonte.
- Feldmann, M., & Biskup, D. 2003. Single-machine scheduling for minimizing earliness and tardiness penalties by meta-heuristic approaches. *Computers & Industrial Engineering*, **44**, 307–323.
- Feo, T. A., & Resende, M. G. C. 1995. Greedy randomized adaptive search procedures. *Journal of Global Optimization*, **6**, 109–133.
- França Filho, M. F. 2007. *GRASP e Busca Tabu aplicados a problemas de programação de tarefas em máquinas paralelas*. Ph.D. thesis, UNICAMP, Campinas.
- Gendreau, M. 2003. An Introduction to Tabu Search. 37-54 2 of: Glover, F., & Kochenberger, G. A. (eds), *Handbook of Metaheuristics*. Kluwer Academic Publishers.
- Glover, F. 1986. Future paths for integer programming and links to artificial intelligence. *Computers & Operations Research*, **13**(5), 533–549.
- Glover, F. 1996. *Tabu Search and adaptive memory programming - advances, applications and challenges*. Interfaces in computer science and operations research edn. Kluwer Academic Publishers. Pages 1–75.
- Glover, F., & Laguna, M. 1997. *Tabu Search*. Boston: Kluwer Academic Publishers.
- Glover, F., & Laguna, M. 2003. Scatter Search and Path Relinking: Advances and Applications. Chap. 1, pages 1–35 of: Glover, F., & Kochenberger, G. A. (eds), *Handbook of Metaheuristics*. Kluwer Academic Publishers.
- Gomes Júnior, A. C. 2007. *Problema de Seqüenciamento em uma Máquina com Penalidades por Antecipação e Atraso: Modelagem e Resolução*. M.Phil. thesis, UFMG, Belo Horizonte.
- Gomes Júnior, A. C., Carvalho, C. R. V., Munhoz, P. L. A., & Souza, M. J. F. 2007. Um método heurístico híbrido para a resolução do problema de sequenciamento em uma máquina com penalidades por antecipação e atraso da produção. Pages 1649–1660 of: *Anais do XXXIX Simpósio Brasileiro de Pesquisa Operacional*. SOBRAPO.
- Gray, P., Hart, W., Painton, L., Phillips, C., Trahan, M., & Wagner, J. 1997. *A Survey of Global Optimization Methods*. Tech. rept. Sandia National Laboratories. Disponível em < <http://www.cs.sandia.gov/opt/survey> >. Acesso em: 27 nov. 2008.
- Gupta, S. R., & Smith, J. S. 2006. Algorithms for single machine total tardiness scheduling with sequence dependent setups. *European Journal of Operational Research*, **175**, 722–739.
- Hansen, P., & Mladenović, N. 2003. Variable Neighborhood Search. 145-184 6 of: Glover, F., & Kochenberger, G. A. (eds), *Handbook of Metaheuristics*. Kluwer Academic Publishers.

- Henderson, D., Jacobson, S. H., & Johnson, A. W. 2003. The Theory and Practice of Simulated Annealing. *Chap. 10, pages 287–319 of: Glover, F., & Kochenberger, G. A. (eds), Handbook of Metaheuristics*. Kluwer Academic Publishers.
- Hino, C. M., Ronconi, D. P., & Mendes, A. B. 2005. Minimizing earliness and tardiness penalties in a single-machine problem with a common due date. *European Journal of Operational Research*, **160**, 190–201.
- James, R. J. W. 1997. Using tabu search to solve the common due date early/tardy machine scheduling problem. *Computers & Operations Research*, **24**(3), 199–208.
- Kirkpatrick, S., Gelatt, C. D., & Vecchi, M. P. 1983. Optimization by simulated annealing. *Science*, **220**, 671–680.
- Koulamas, C. 1996. Single-machine scheduling with time windows and earliness/tardiness penalties. *European Journal of Operational Research*, **91**, 190–202.
- Lee, C. Y., & Choi, J. Y. 1995. A Genetic Algorithm for Job Sequencing Problems with Distinct Due Dates and General Early-Tardy Penalty Weights. *Computers and Operations Research*, **22**(8), 857–869.
- Lee, C. Y., & Kim, S. J. 1995. Parallel genetic algorithms for the earliness-tardiness job scheduling problem with general penalty weights. *Computers & industrial engineering*, **28**(2), 231–243.
- Li, G. 1997. Single machine earliness and tardiness scheduling. *European Journal of Operational Research*, **96**, 546–558.
- Liaw, C.-F. 1999. A branch-and-bound algorithm for the single machine earliness and tardiness scheduling problem. *Computers & Operations Research*, **26**, 679–693.
- Lourenço, H. R., Martin, O. C., & Stützle, T. 2003. Iterated Local Search. *Chap. 11, pages 321–353 of: Glover, F., & Kochenberger, G. A. (eds), Handbook of Metaheuristics*. Kluwer Academic Publishers.
- Mazzini, R., & Armentano, V. A. 2001. A heuristic for single machine scheduling with early and tardy costs. *European Journal of Operational Research*, **128**, 129–146.
- M’Hallah, R. 2007. Minimizing total earliness and tardiness on a single machine using a hybrid heuristic. *Computers & Operations Research*, **34**, 3126–3142.
- Mladenović, N., & Hansen, P. 1997. Variable Neighborhood Search. *Computers and Operations Research*, **24**(11), 1097–1100.
- Mondal, S. A. 2002. Minimization of squared deviation of completion times about a common due date. *Computers & Operations Research*, **29**, 2073–2085.
- Mondal, S. A., & Sen, A. K. 2001. Single machine weighted earliness tardiness penalty problem with a common due date. *Computers & Operations Research*, **28**, 649–669.

- Panwalkar, S. S., Dudek, R. A., & Smith, M. L. 1973. Sequencing research and the industrial scheduling problem. *Pages 29–38 of: Elmaghraby, S. E. (ed), Symposium on the Theory of Scheduling and its Applications.*
- Penna, Puca Huachi Vaz. 2009. *Um algoritmo heurístico híbrido para minimizar os custos com a antecipação e o atraso da produção em ambientes com janelas de entrega e tempos de preparação dependentes da sequência.* M.Phil. thesis, UFOP, Ouro Preto.
- Rabadi, G., Mollaghasemi, M., & Anagnostopoulos, G. C. 2004. A branch-and-bound algorithm for the early/tardy machine scheduling problem with a common due-date and sequence-dependent setup time. *Computers & Operations Research*, **31**, 1727–1751.
- Reeves, C. 2003. Genetic Algorithms. *Chap. 3, pages 55–82 of: Glover, F., & Kochenberger, G. A. (eds), Handbook of Metaheuristics.* Kluwer Academic Publishers.
- Resende, M. G.C., & Ribeiro, C. C. 2003. Greedy Randomized Adaptive Search Procedures. *219–249 of: Glover, F., & Kochenberger, G. A. (eds), Handbook of Metaheuristics.* Kluwer Academic Publishers.
- Ribeiro, C. C. 1996. Metaheuristics and Applications. In *Advanced School on Artificial Intelligence*, Estoril, Portugal.
- Ribeiro, Fábio Fernandes. 2009. *Um algoritmo genético adaptativo para a resolução do problema de sequenciamento em uma máquina com penalização por antecipação e atraso da produção.* M.Phil. thesis, CEFET-MG, Belo Horizonte.
- Sourd, F. 2005. Optimal timing of a sequence of tasks with general completion costs. *European Journal of Operational Research*, **165**, 82–96.
- Souza, M. J. F. 2008. *Inteligência Computacional para Otimização.* Notas de aula. Disponível em < [http : //www.decom.ufop.br/prof/marcone/ Disciplinas/InteligenciaComputacional/InteligenciaComputacional.htm](http://www.decom.ufop.br/prof/marcone/ Disciplinas/InteligenciaComputacional/InteligenciaComputacional.htm) >. Acesso em: 20 nov. 2008.
- Souza, M. J. F., Penna, P. H. V., Gonçalves, F. A. C. A., & Ochi, L. S. 2008. Uma heurística híbrida para minimizar custos com antecipação e atraso em sistemas de produção com janelas de entrega e tempos de preparação dependentes da sequência. *Pages 1–16 of: Anais do XI Simpósio de Pesquisa Operacional e Logística da Marinha.*
- Szwarc, W. 1996. The weighted common due date single machine scheduling problem revisited. *Computers and Operations Research*, **23**(3), 255–262.
- Tavakkoli-Moghaddam, R., Moslehi, G., Vasei, M., & Azaron, A. 2006. A branch-and-bound algorithm for a single machine sequencing to minimize the sum of maximum earliness and tardiness with idle insert. *Applied Mathematics and Computation*, **174**, 388–408.
- Valente, J. M.S., & Alves, R. A. F. S. 2005a. Filtered and recovering beam search algorithms for the early/tardy scheduling problem with no idle time. *Computers & Industrial Engineering*, **48**, 363–375.

- Valente, J. M.S., & Alves, R. A. F. S. 2005b. Improved heuristics for the early/tardy scheduling problem with no idle time. *Computers & Operations Research*, **32**, 557–569.
- Ventura, J. A., & Radhakrishnan, S. 2003. Single machine scheduling with symmetric earliness and tardiness penalties. *European Journal of Operational Research*, **144**, 598–612.
- Wan, G., & Yen, B. P.-C. 2002. Tabu search for single machine scheduling with distinct due windows and weighted earliness/tardiness penalties. *European Journal of Operational Research*, **142**, 271–281.
- Weise, T. 2008. *Global Optimization Algorithms: Theory and Application*. 2ª edição edn. Disponível em: < [http : //www.it – weise.de/](http://www.it-weise.de/) >. Acesso em: 01 dez. 2008.
- Weng, X., & Ventura, J. A. 1996. Scheduling about a given common due date to minimize mean squared deviation of completion times. *European Journal of Operational Research*, **88**, 328–335.
- Ying, K.-C. 2008. Minimizing earliness-tardiness penalties for common due date single-machine scheduling problems by a recovering beam search algorithm. *Computers & Industrial Engineering*, doi:10.1016/j.cie.2008.01.008.
- Ziviane, N. 2004. *Projeto de Algoritmos com implementações em PASCAL e C*. 2ª edição revista e ampliada edn. Thomson.