



**UNIVERSIDADE FEDERAL RURAL DO SEMI-ÁRIDO
UNIVERSIDADE ESTADUAL DO RIO GRANDE DO NORTE
PROGRAMA DE PÓS-GRADUAÇÃO EM CIÊNCIAS DA
COMPUTAÇÃO**



KAIO HENRIQUE FONSECA DANTAS

**OTIMIZAÇÃO DO CORTE GUILHOTINADO
NÃO-ESTAGIADO ATRAVÉS DE APRENDIZAGEM POR
REFORÇO COM FUNÇÃO RECOMPENSA TARDIA**

MOSSORÓ

2025

KAIO HENRIQUE FONSECA DANTAS

**OTIMIZAÇÃO DO CORTE GUILHOTINADO
NÃO-ESTAGIADO ATRAVÉS DE APRENDIZAGEM POR
REFORÇO COM FUNÇÃO RECOMPENSA TARDIA**

Dissertação apresentada ao Programa de Pós-Graduação em Ciências da Computação - associação ampla entre a Universidade Estadual do Rio Grande do Norte e a Universidade Federal Rural do Semi-Árido, para a obtenção do título de Mestre em Ciência da Computação.

Orientador: Prof. Dr. Dario José Aloise-
UERN

Coorientador: Prof. Dr. Thiago Henrique Freire
de Oliveira- IFRN

MOSSORÓ

2025

© Todos os direitos estão reservados a Universidade do Estado do Rio Grande do Norte. O conteúdo desta obra é de inteira responsabilidade do(a) autor(a), sendo o mesmo, passível de sanções administrativas ou penais, caso sejam infringidas as leis que regulamentam a Propriedade Intelectual, respectivamente, Patentes: Lei nº 9.279/1996 e Direitos Autorais: Lei nº 9.610/1998. A mesma poderá servir de base literária para novas pesquisas, desde que a obra e seu(a) respectivo(a) autor(a) sejam devidamente citados e mencionados os seus créditos bibliográficos.

Catálogo da Publicação na Fonte.
Universidade do Estado do Rio Grande do Norte.

D192o Dantas, Kaio Henrique Fonseca
OTIMIZAÇÃO DO CORTEGUILHOTINADO NÃO-
ESTAGIADO ATRAVÉS DE APRENDIZAGEM POR
REFORÇO COM FUNÇÃO RECOMPENSA TARDIA. /
Kaio Henrique Fonseca Dantas. - Mossoró, 2025.
89p.

Orientador(a): Prof. Dr. Dario José Aloise.

Coorientador(a): Prof. Dr. Thiago Henrique Freire de Oliveira.

Dissertação (Mestrado em Programa de Pós-Graduação em Ciência da Computação). Universidade do Estado do Rio Grande do Norte.

1. CSP. 2. Corte Guilhotinado. 3. Q-Learning. 4. Aprendizagem de Máquina. I. Aloise, Dario José. II. Universidade do Estado do Rio Grande do Norte. III. Título.

RESUMO

O Problema de Corte de Estoque (CSP - *Cutting Stock Problem*) é um clássico desafio de otimização combinatória com ampla aplicação industrial, cuja principal meta é reduzir o desperdício de material durante o processo de corte. No entanto, métodos tradicionais demonstram dificuldades em se adaptar a cenários dinâmicos e parametrizações complexas. Este estudo propõe uma reavaliação das estratégias de resolução do problema e o desenvolvimento de uma versão aprimorada do algoritmo Q-Learning, adaptada para o Problema de Corte Guilhotinado de Estoque (GCSP - *Guillotine Cutting Stock Problem*).

A pesquisa introduz uma nova função de recompensa, capaz de avaliar a diferença entre áreas remanescentes sucessivas em relação à área da chapa, penalizando ações que demandam a utilização de novas chapas. Além disso, a taxa de aprendizado (α) é ajustada de forma adaptativa conforme a taxa de exploração atinge um limiar predefinido, promovendo uma convergência mais estável.

Os experimentos demonstraram que a configuração da função de recompensa exerce forte influência sobre o comportamento do agente, podendo inclusive induzir respostas escalares que reduzem a relevância do aprendizado de máquina. Logo esse trabalho demonstra que as novas parametrizações e as estratégias propostas contribuem para o estado da arte e servirão como referência ou *benchmark* para futuras pesquisas no domínio do GCSP.

Palavras-chave: CSP. Corte Guilhotinado. Q-Learning. Aprendizagem de Máquina.

ABSTRACT

The Cutting Stock Problem (CSP) is a classic combinatorial optimization challenge with broad industrial applications, whose main goal is to minimize material waste during the cutting process. However, traditional methods often struggle to adapt to dynamic scenarios and complex parameter configurations.

This study proposes a re-evaluation of problem-solving strategies and the development of an enhanced version of the Q-Learning algorithm, adapted to the Guillotine Cutting Stock Problem (GCSP).

The research introduces a novel reward function capable of assessing the difference between successive remaining areas relative to the sheet area, penalizing actions that require the use of new sheets. Furthermore, the learning rate (α) is adaptively adjusted as the exploration rate reaches a predefined threshold, promoting more stable convergence.

Experimental results show that the reward function design has a strong influence on the agent's behavior, potentially inducing scalar responses that diminish the relevance of machine learning. Therefore, this work demonstrates that the proposed parameterizations and strategies contribute to the state of the art and may serve as a reference or benchmark for future research in the GCSP domain.

Keywords: CSP. Guillotin cutting. Q-Learning. Machine Learning.

LISTA DE FIGURAS

Figura 1 – Metodologia de Pesquisa de Acordo com a <i>Design Science</i>	13
Figura 2 – Diagrama de um Processo de Decisão de Markov (estados, ações, transições e recompensas).	25
Figura 3 – Cortes possíveis	27
Figura 4 – Casos especiais	28
Figura 5 – HHD Heuristic	29
Figura 6 – Metodologia Híbrida	30
Figura 7 – Obtendo delta das chapas e da área das sobras	32
Figura 8 – Fluxograma GRASP-Learning com Busca local	35
Figura 9 – Matriz-Q final de dois treinamentos com 200 episódios para o método xyz	39
Figura 10 – Matriz Q resultante da convergência com x para a próxima a ação e y para a ação atual	41
Figura 11 – Iteração 1 - $S_0 \rightarrow S_{23}$	41
Figura 12 – Iteração 2 - $S_{23} \rightarrow S_{32}$	42
Figura 13 – Iteração 3 - $S_{32} \rightarrow S_{14}$	42
Figura 14 – Iteração 4 - $S_{14} \rightarrow S_{25}$	43
Figura 15 – Layout de corte parcial	43
Figura 16 – Áreas de itens do layout parcial de corte	43
Figura 17 – Matriz Q resultante com função de recompensa tardia	44

LISTA DE TABELAS

Tabela 1 – Descrição dos símbolos utilizados na formulação	17
Tabela 2 – Parâmetros utilizados por Rebouças (2021)	38
Tabela 3 – Parâmetros utilizados tentativa de convergência	40
Tabela 4 – Comparativo entre os resultados de todos os algoritmos	44
Tabela 5 – Comparativo entre os resultados da Classe 10 para os algoritmos pro- postos e resultados da literatura	45
Tabela 6 – Resultados experimentais para diferentes instâncias	68
Tabela 7 – Resultados experimentais para diferentes instâncias	86
Tabela 8 – Resultados experimentais para diferentes instâncias	88

LISTA DE ABREVIATURAS E SIGLAS

QC	Questões Conceituais
QGP	Questão Geral de Pesquisa
QP	Questões Práticas
QT	Questões Tecnológicas

SUMÁRIO

1	INTRODUÇÃO	10
1.1	Questões de Pesquisa	11
1.2	Metodologia de pesquisa	12
1.3	Organização do Documento	13
2	REFERENCIAL TEÓRICO	15
2.1	O Problema de Corte e Empacotamento	15
2.2	O Problema de Corte Bidimensional Guilhotinado	16
2.3	Métodos para Resolver o Problema	16
2.4	Formulação Matemática	17
2.5	Revisão de Literatura	18
3	ABORDAGENS HEURÍSTICAS	21
3.1	Meta-heurísticas	21
3.1.1	<i>Grasp</i>	21
3.1.1.1	<i>Fase Construtiva</i>	22
3.1.1.2	<i>Busca Local</i>	23
3.1.2	<i>Q-Learning</i>	23
3.1.2.1	<i>Q-Learning: Fundamentos e Fundamentação</i>	24
3.1.2.1.1	<i>Processo de Decisão de Markov</i>	24
3.1.2.2	<i>Q-Learning: definição e elementos</i>	25
3.1.3	<i>HHDHeuristic</i>	26
4	ABORDAGEM COMPUTACIONAL	30
4.1	Função de recompensa tardia	31
4.2	Base de dados	33
4.3	Avaliação	33
4.3.1	<i>GRASP-Learning com busca local simples</i>	34
4.3.1.1	<i>Leitura da Instância e Inicialização do Problema</i>	34
4.3.1.2	<i>Treinamento com Q-Learning</i>	35
4.3.1.3	<i>Convergência</i>	36
4.3.1.4	<i>Fase Construtiva</i>	36
4.3.2	<i>Retorno da Solução Final</i>	37

5	AVALIAÇÃO DAS ABORDAGENS COMPUTACIONAIS	38
5.1	Experimentos Computacionais	38
5.1.1	<i>Ambiente de Desenvolvimento</i>	38
5.1.2	<i>Análise dos Resultados</i>	38
6	CONCLUSÕES	46
6.1	Trabalhos futuros	46
	REFERÊNCIAS	48
	APÊNDICES	51
	APÊNDICE A – Resultados - Grasp-Learning-Delayed	51
	APÊNDICE B – Resultados - Grasp-Learning-Delayed (BLA)	69
	APÊNDICE C – Resultados - Grasp-Learning-Delayed (ILSA)	87
	ANEXOS	88

1 INTRODUÇÃO

O Problema de Corte de Estoque (Cutting Stock Problem - CSP) é um clássico problema de otimização combinatória que surge em diversos contextos industriais, onde o objetivo é cortar materiais em estoque de grandes dimensões em peças menores demandadas, minimizando o desperdício. Exemplos comuns incluem o corte de chapas metálicas, rolos de papel ou tábuas de madeira. Algoritmos exatos e heurísticas tradicionais têm sido amplamente empregados para resolver este problema (LODI *et al.*, 2002), mas frequentemente apresentam dificuldades em termos de escalabilidade e adaptabilidade em cenários mais dinâmicos ou complexos.

Nos últimos anos, o Aprendizado por Reforço (AR) tem ganhado atenção como uma abordagem viável para lidar com problemas de otimização combinatória (BENGIO *et al.*, 2021). Em particular, o algoritmo Q-learning—um algoritmo de AR livre de modelo—tem sido explorado como um meio para derivar políticas adaptativas para resolver problemas como o CSP. Aplicações anteriores do Q-learning em CSPs demonstraram resultados promissores, evidenciando sua capacidade de aprender sequências de corte eficientes e reduzir o desperdício ao longo do tempo (PITOMBEIRA *et al.*, 2021).

Neste estudo, seguimos a estratégia híbrida proposta por (REBOUÇAS *et al.*, 2023) para resolver o Problema de Corte de Estoque Guilhotinado (Guillotine Cutting Stock Problem - GCSP). No entanto, introduzimos uma melhoria ao algoritmo Q-learning, que é um dos componentes utilizados no processo de solução. Essa melhoria consiste em uma nova função de recompensa que considera tanto os estados passados quanto futuros das ações, além de penalizar o uso de novas chapas. Essas penalidades atuam como incentivo de aprendizado que guia o agente para soluções mais eficientes.

Além disso, para direcionar o processo de aprendizado à convergência, introduzimos também uma programação adaptativa da taxa de aprendizado. Uma vez que o parâmetro ϵ na política ϵ -greedy decai abaixo de um limite definido pelo usuário, a taxa de aprendizado α é gradualmente reduzida. Esse ajuste estabiliza as atualizações dos valores Q à medida que o agente transita da exploração para a exploração, melhorando assim a convergência para uma política eficaz (TOKIC; PALM, 2011).

1.1 Questões de Pesquisa

Seguindo as ideias de Wieringa (2014), foi formulado um problema de pesquisa que utiliza a design science como base para sua investigação e resolução. Nesse modelo a partir de uma Questão Geral de Pesquisa (QGP), são derivadas Questões Conceituais (QC) para explorar o conhecimento teórico e conceitual relacionado ao problema, Questões Tecnológicas (QT) para desenvolver soluções práticas e inovadoras e Questões Práticas (QP) para implementar e avaliar essas soluções no mundo real. Essas abordagens englobam diferentes aspectos da pesquisa, desde a teoria até a aplicação prática dos resultados (Wieringa, 2009). Para esse trabalho foi definida a seguinte questão geral de pesquisa (QGP):

Como otimizar a utilização de material na realização de cortes guilhotinados?

E a partir dela, foram definidas as seguintes questões de específicas :

1. QC: Como ocorre o processo de corte guilhotinado nas indústrias

- a) Quais são requisitos os necessários para que uma sequência de cortes guilhotinados, seja considerada bem sucedida?
- b) Quais os principais recursos gastos durante o processo e os possíveis desperdícios envolvidos na operação?
- c) Quais fatores são levados em conta para o empenho de valores para arcar com o custo da operação?
- d) Como funcionam os equipamentos que realizam efetivamente os cortes guilhotinados?

2. QT: Quais ferramentas e técnicas de otimização podem ser utilizadas na otimização da utilização de materiais no processo de corte guilhotinado?

- a) Nesse contexto, existe uma formulação matemática que modele o problema de corte guilhotinado?
- b) Através da execução de técnicas em algoritmos, pressupondo a utilização instâncias de pequeno, médio e grande porte, seria possível encontrar resultados equiparados ou melhores que o estado da arte, do problema?
- c) Nesse contexto quais meta-heurísticas poderiam ser empregadas para resolver o problema com soluções próximas do ótimo? Seria possível identificar suas limitações e apresentar melhorias para elas
- d) No contexto da aprendizagem de máquina, uma função de recompensa pode

ser conceitualmente melhorada?

- e) Parametrizações mais efetivas das técnicas existentes podem ser encontradas?
- f) Outros componentes da abordagem anterior podem ser substituídos para melhorar os resultados?

3. **QP: Como implantar os algoritmos e técnicas de otimização em equipamentos que realizem corte guilhotinado?**

- a) Quais possíveis cenários poderão ser melhor aproveitados pelas partes interessadas?
- b) Existem equipamentos que processariam os algoritmos desenvolvidos nativamente?
- c) Estatisticamente qual é o nível de exatidão e velocidade que a execução do algoritmo pode agir durante uma tomada de decisão das partes interessadas?
- d) Seria possível utilizar a mesma abordagem desenvolvida para resolver problemas similares em outras áreas?

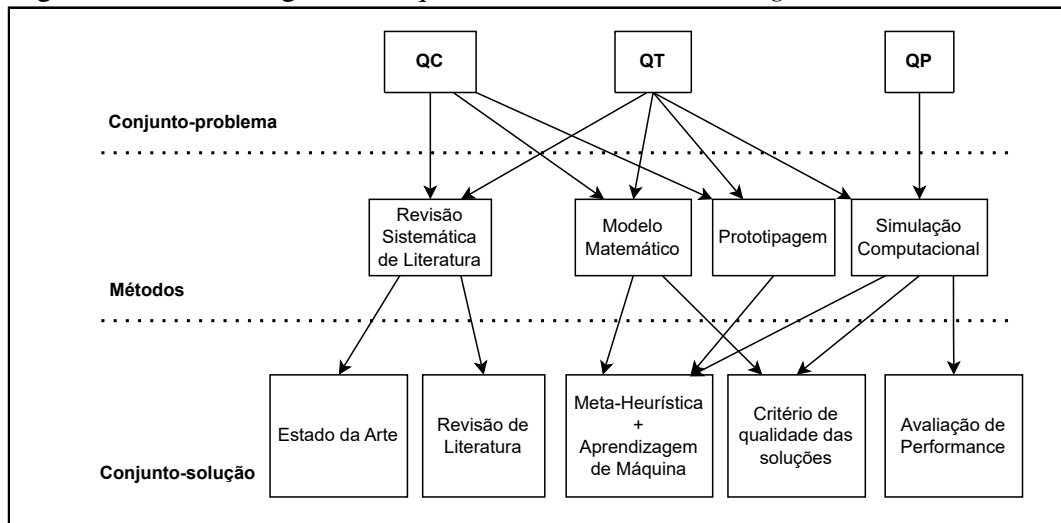
1.2 Metodologia de pesquisa

A metodologia de pesquisa empregada neste trabalho é baseada em métodos específicos que conectam o conjunto de problemas de pesquisa ao conjunto de soluções em forma de artefatos. Cada método é aplicado em um dos ciclos da *Design Science* - investigação do problema, projeto de tratamento e validação do tratamento - para fornecer respostas aos problemas de pesquisa (WIERINGA, 2009). A Figura 1, representa como os métodos aplicados ao conjunto-problema podem criar artefatos no conjunto-solução.

A seguir são descritos o que pretende-se alcançar com cada um dos métodos mencionados e a forma como serão executados:

1. **Revisão de Literatura e Estado da Arte:** O intuito da revisão de literatura é explorar o tema das questões de conhecimento e entender o atual cenário do corte guilhotinado, a fim de explorar e conhecer quais requisitos necessários para que uma sequência de cortes seja considerada bem avaliada, assim como os principais problemas envolvidos na atividade. Por ventura, utiliza-se a temática das questões técnicas para conhecer, principalmente, as técnicas de otimização utilizadas para solucionar problemas de corte guilhotinado em diferentes cenários, assim como elencar possíveis brechas, deixadas por trabalhos anteriores, que possam ser exploradas para enriquecer o

Figura 1 – Metodologia de Pesquisa de Acordo com a *Design Science*.



Fonte: Autoria própria.

estado da arte;

2. **Meta-Heurística + Aprendizagem de Máquina:** Implementar a partir das questões conceituais e técnicas visto a necessidade de conhecer o problema, assim como utilizar o modelo matemático a fim de obter ou adaptar uma meta-heurística em forma de algoritmo, que respeite as restrições do problema apresentado e minimize ou maximize valores próximo ao ótimo, obtido por técnicas de otimização linear associadas a aprendizagem de máquina; e
3. **Critério de qualidade das soluções:** Definir os critérios de qualidade para a solução encontrada, uma vez obtidos o modelo matemático e o algoritmo da meta-heurística. Será necessário realizar o ajuste fino dos parâmetros para execução em computador, da solução desenvolvida.
4. **Avaliação de Performance:** Avaliar se o artefato obtido, Meta-Heurística associada a aprendizagem de máquina, superam métodos existentes. Para tal, é necessário realizar simulações em laboratório em busca da obtenção de estatística em relação a desempenho e obtenção de valores próximo em uma análise comparativa com os métodos documentados em literatura e trabalhos anteriores.

1.3 Organização do Documento

Este trabalho encontra-se organizado em mais quatro capítulos além dessa introdução. O Capítulo 2 traz um referencial teórico atualizada que caracteriza o estado da arte para o contexto deste estudo. No Capítulo 3 é apresentada a proposta de solução e

metodologia adotada para atingir os critérios de qualidade das soluções. No Capítulo 4 são apresentados os resultados decorrente da prototipação e análise de performance da solução. Finalmente o Capítulo 5 traz uma reflexão dos resultados obtidos, a relevância do trabalho na área, assim como as principais dificuldades encontradas no processo de execução da pesquisa e os respectivos trabalhos futuros realizados que podem surgir ao referenciar esta dissertação.

2 REFERENCIAL TEÓRICO

2.1 O Problema de Corte e Empacotamento

Os Problemas de Corte e Empacotamento (PCE) constituem uma classe abrangente dentro da área de otimização combinatória, caracterizando-se por um objetivo comum: dividir materiais de grandes dimensões — ou espaços disponíveis — em partes menores, denominadas itens, cujas dimensões e formatos são previamente definidos. Tais problemas aparecem com frequência em contextos práticos. Por exemplo, uma indústria moveleira precisa cortar chapas de madeira de maneira a minimizar o desperdício de material. Da mesma forma, uma empresa de transporte de cargas necessita distribuir os volumes a serem transportados de modo a aproveitar ao máximo o espaço interno dos veículos.

Nos problemas de corte, os objetos maiores são materiais sólidos que devem ser divididos em itens menores. Esses problemas podem ser classificados de acordo com suas dimensões: unidimensionais, como o corte de barras metálicas ou bobinas de papel, que ocorre em apenas uma dimensão; bidimensionais, como o corte de placas de madeira, tecidos ou chapas de aço; e tridimensionais, como o corte de blocos de espuma utilizados na fabricação de colchões e travesseiros. Em todos os casos, o objetivo principal é minimizar o desperdício de material, dado que esse fator impacta diretamente a eficiência e o custo do processo produtivo.

Os problemas de empacotamento apresentam estrutura análoga aos problemas de corte, porém, com uma interpretação distinta. Nesse caso, os objetos maiores representam espaços vazios — como os interiores de veículos, contêineres ou paletes — que devem ser ocupados por itens menores. Assim como nos problemas de corte, pode-se compreender o empacotamento como uma forma de “recorte” do espaço vazio em volumes que serão preenchidos pelos itens menores (DYCKHOFF, 1990). O objetivo, nesse contexto, é minimizar o espaço ocioso.

Apesar da diversidade de aplicações, os PCE compartilham uma estrutura lógica comum. Em geral, parte-se de dois conjuntos: o primeiro contém os objetos maiores, que representam a matéria-prima ou o espaço disponível; o segundo reúne os itens menores, que representam a demanda a ser alocada ou cortada. O desafio está em realizar essa alocação de forma que:

- todos os itens sejam alocados integralmente dentro do objeto correspondente;

- os itens não apresentem sobreposição entre si.

Com base nisso, define-se uma função objetivo, que pode consistir na minimização do desperdício, da quantidade de objetos utilizados ou do espaço ocioso. Essa estrutura é capaz de modelar diferentes variações dos problemas de corte e empacotamento.

Visando organizar e padronizar a classificação dos PCE, Dyckhoff (1990) propôs uma tipologia que se tornou referência na área. Posteriormente, diante de limitações encontradas para representar problemas mais recentes, essa tipologia foi atualizada por Wäscher, Haußner e Schumann (2007). Logo os problemas foram classificados pela dimensionalidade: seja unidimensional, bidimensional ou tridimensional, e consequentemente por restrições aplicadas a sua natureza.

- **Problemas unidimensionais:** Cortar barras ou rolos de materiais.
- **Problemas bidimensionais:** Cortar placas retangulares de madeira, vidro ou metal.
- **Problemas tridimensionais:** Empacotar caixas em contêineres ou cortar blocos de espuma.

2.2 O Problema de Corte Bidimensional Guilhotinado

Neste tipo específico de problema, os cortes têm uma restrição adicional: devem ser feitos sempre atravessando todo o material. Essa limitação torna o problema mais complicado de resolver, já que não podemos simplesmente cortar as peças como quisermos. Este tipo de problema foi extensivamente estudado por autores como Gilmore e Gomory (1965) e Martello e Vigo (1998), que propuseram modelos matemáticos e algoritmos que servem como base para pesquisas até hoje.

2.3 Métodos para Resolver o Problema

- **Métodos Exatos:** como Programação Linear e o método Branch-and-Bound, que garantem a solução ótima, mas podem ser lentos para grandes instâncias.
- **Heurísticas e Metaheurísticas:** métodos Heurísticos como o HHD-Heuristic proposto por Nascimento (1999) o GRASP por Resende (2005) e o BRKGA também por Resende em (2011) são mais rápidos e geralmente oferecem boas soluções, mas não garantem a solução perfeita.
- **Aprendizagem por Reforço (Q-Learning):** uma técnica que permite ao algoritmo

aprender com a experiência, escolhendo decisões que melhoram continuamente as soluções iniciais.

Este trabalho especificamente explora a combinação de GRASP com Q-Learning, com o intuito de utilizar a abordagem híbrida proposta por Rebouças (2021) aperfeiçoando a função de recompensa para validar o Q-learning empregado na técnica e, consequentemente, que o mesmo gere melhores resultados.

2.4 Formulação Matemática

Segundo Silva (2017), da década de 60 para os dias atuais, nenhuma formulação para o corte guilhotinado não estagiado foi encontrada, até Messaoud, Chu e Espinouse (2008) apresentarem um modelo para a variação do problema proveniente do problema de uma dimensão aberta ou problema de faixa (*Strip Packing Problem*). Mas, somente Furini, Malaguti e Thomopulos (2016) apresentaram um modelo capaz de tratar o caso derivado do problema de corte e empacotamento. Portanto, o Problema de Corte Bidimensional Guilhotinado pode ser modelado, segundo Furini, Malaguti e Thomopulos (2016), como:

Tabela 1 – Descrição dos símbolos utilizados na formulação

Símbolo	Descrição
O	Conjunto de possíveis orientações de cortes vertical e horizontal. $O = \{h, v\}$.
J	Conjunto de nós (por exemplo, nós de roteamento, recursos ou pontos intermediários).
$\bar{J}$	Subconjunto de J cujas demandas são fixas (i.e., nós de demanda).
$Q(j, o)$	Conjunto de colunas (ou sequências de decisões) viáveis para atender à ordem o no nó j .
$a_{q,k,j}^o$	Coeficiente que indica quantas unidades da coluna q , oriunda do nó $k \in J$, são encaminhadas ao nó j para atender à ordem o .
d_j	Demanda fixa (ou volume requerido) no nó $j \in \bar{J}$.
$x_{q,j}^o$	Variável inteira não negativa que indica a quantidade de colunas $q \in Q(j, o)$ usadas para atender à ordem o no nó j .
y_j	Quantidade de recurso adicional (por exemplo, número de veículos, novas unidades ou cortes extras) que se deve introduzir no nó j .

$$\min \sum_{o \in O} \sum_{q \in Q} x_{q,0}^o + y_0 \quad (2.1)$$

$$\sum_{k \in J} \sum_{o \in O} \sum_{q \in Q(k,o)} a_{q,k,j}^o x_{q,k}^o - \sum_{o \in O} \sum_{q \in Q(j,o)} x_{q,j}^o - y_j \geq 0, \forall j \in \bar{J}, j \neq 0 \quad (2.2)$$

$$\sum_{k \in J} \sum_{o \in O} \sum_{q \in Q(k,o)} a_{q,k,j}^o x_{q,k}^o - \sum_{o \in O} \sum_{q \in Q(j,o)} x_{q,j}^o \geq 0, \forall j \in J \setminus \bar{J} \quad (2.3)$$

$$y_j = d_j, \forall j \in \bar{J} \quad (2.4)$$

$$x_{q,j}^o \in \mathbb{Z}^+, \forall j \in J, o \in O, q \in Q(j,o) \quad (2.5)$$

$$x_{q,j}^o \in \mathbb{Z}^+, \forall j \in \bar{J} \quad (2.6)$$

A função objetivo está definida em (2.1), a qual tenta minimizar o número de contêineres utilizados. A restrição (2.2) impõe que o número de partes j que estão em cada retângulo ou itens individuais não exceda o número de retângulos j obtidos através de corte de algum outro retângulo. A restrição (2.3) é equivalente à restrição (2.4) para os retângulos $j \notin \bar{J}$ (consequentemente, para quando as variáveis correspondentes y_j não estejam definidas). A restrição (2.3) é equivalente à restrição (2.4) para os retângulos $j \notin \bar{J}$ (consequentemente, para quando as variáveis correspondentes y_j não estejam definidas). A restrição (2.3) impõe que a demanda associada a cada item seja satisfeita e as restrições (2.4) e (2.5) definem o escopo das variáveis. O principal problema dessa abordagem é a quantidade de variáveis geradas e isso impacta fortemente na complexidade em solucionar o problema. Como dito pelos próprios autores, um exemplo do uso desse modelo para solucionar o caso da mochila, onde poucos contêineres são utilizados (geralmente um), gerou cerca de sessenta e seis milhões de variáveis para 50 itens.

2.5 Revisão de Literatura

Na dissertação de mestrado apresentada por Rebouças (2021), foi realizada uma revisão de literatura até o ano de 2020, em busca de contribuições para entender o estado da

arte e conhecer que técnicas foram utilizadas para resolver o problema de corte guilhotinado e empacotamento.

Esse capítulo propõe complementar a pesquisa realizada, a partir de 2021, em busca de novas abordagens utilizadas para solucionar o problema.

Em 2021, Jiang et al. (2021) propuseram um agente DRL multimodal para o problema de empacotamento 3D, que recebe uma sequência, orientação e posição de cada item e que consegue solucionar o problema, com uma abordagem de multicamadas processando até 100 caixas, superando abordagens anteriores em qualidade de empacotamento e escalabilidade nesse tipo de problema.

Também em 2021, Zhang, Zi Ge (2021) apresentam o *Attend2Pack*, um modelo fim-a-fim com atenção e controle com sobreposição priorizada, para alcançar resultados aprimorados no empacotamento offline e também no modo online.

Hang Zhao et al. (2022) introduzem o *Packing Configuration Tree (PCT)*, combinando árvores hierárquicas com *Deep Reinforcement Learning (DRL)* para empacotamento online 3D. Na pesquisa, o agente opera em um espaço de ações contínuo e chega a superar métodos heurísticos em ambientações reais.

Murdivien Um (2023) propõem o *BoxStacker*, combinando DRL com visualização via game engine para empacotamento 3D sequencial em logística. Demonstram eficácia em cenários realistas e aprendizado robusto em tempo real.

Pan et al. (2023), por sua vez, apresentam o AR2L, um framework de RL robusto para empacotamento online 3D, que otimiza o retorno médio e o worst-case via adversarial RL. Os resultados destacaram apontaram que a política obtida é confiável e eficaz, segundo os autores.

Xiong et al. (2024) propõem o *Generalizable Online 3D Bin Packing via Transformer-based Deep Reinforcement Learning (GOPT)*, usando transformadores para generalização em empacotamento online 3D com diferentes dimensões de caixa, o processo inclui um módulo de geração de subespaços e execução robótica, o qual, objetive forte desempenho generalizado.

Pugliese et al. (2025) abordam o problema 2D+1 com restrições de altura, usando agente DRL treinado em OpenAI Gym. Evitando heurísticas explícitas de forma a aprimoram a decisão espacial em ambientes simulados.

Finalmente, Lee Nam (2025) introduzem um framework hierárquico para empaco-

tamento 2D online/semi-online, que combina heurísticas, DRL e manipuladores duplos, suportando repacking e alcançando soluções quase ótimas.

3 ABORDAGENS HEURÍSTICAS

3.1 Meta-heurísticas

De acordo com Osman e Kelly (1996), meta-heurísticas são métodos iterativos que orientam a construção e modificação de soluções por meio da combinação inteligente de diferentes estratégias, com o objetivo de explorar e explorar eficientemente o espaço de busca. Tais métodos estruturam informações e utilizam mecanismos adaptativos que favorecem a obtenção de soluções de alta qualidade, mesmo em problemas de grande complexidade.

Diferentemente dos algoritmos de otimização exata, as meta-heurísticas não garantem a obtenção de soluções ótimas globais. Além disso, em contraste com algoritmos de aproximação, não fornecem limites formais sobre a qualidade relativa das soluções encontradas. Entretanto, destacam-se por sua aplicabilidade ampla e pela boa performance em contextos onde métodos exatos tornam-se inviáveis do ponto de vista computacional.

As meta-heurísticas podem ser encaradas como esquemas gerais de resolução, os quais podem ser adaptados a diferentes classes de problemas com poucas modificações estruturais. Dentre os exemplos mais conhecidos estão: Busca Tabu, Algoritmos Genéticos, Recozimento Simulado (Simulated Annealing), Otimização por Colônia de Formigas (*Ant Colony Optimization*) e Busca Local Iterada (*Iterated Local Search*).

Segundo Blum e Roli (2003), dois princípios fundamentais regem o funcionamento eficaz de uma metaheurística: a diversificação e a intensificação. A diversificação está associada à exploração ampla do espaço de busca, permitindo a investigação de diferentes regiões da solução. Já a intensificação visa aprofundar a busca em áreas promissoras, aproveitando o conhecimento acumulado ao longo das iterações anteriores. Em resumo, a diversificação promove a variedade de soluções visitadas, enquanto a intensificação refina os resultados obtidos, contribuindo para a convergência a soluções de alta qualidade.

3.1.1 *Grasp*

O *Greedy Randomized Adaptive Search Procedure* (GRASP) é uma metaheurística de busca multi-inicialização, amplamente empregada em problemas de otimização combinatória. Proposta originalmente por Feo e Resende (1995), sua estrutura é composta por duas técnicas, primeiramente uma heurística construtiva para na sua fase inicial,

gerando possíveis soluções e finalmente uma heurística de busca local, com o intuito de melhorar-las.

3.1.1.1 Fase Construtiva

Na fase construtiva, uma solução viável é construída de forma iterativa, adicionando componentes selecionados a partir de uma *Lista Restrita de Candidatos* (LRC). Esta lista é formada com base em um critério de ganância (greedy), que avalia os candidatos disponíveis segundo uma função de custo. Para garantir a diversidade das soluções geradas, o GRASP insere um mecanismo de aleatoriedade controlada na seleção dos elementos da LRC. Essa abordagem equilibra a busca entre intensificação e diversificação.

A fase construtiva é responsável por, a cada iteração, gerar uma solução inicial viável para o GRASP. Esta solução é construída selecionando-se um elemento por vez, de forma iterativa, e inserindo-o à solução parcial, até que esta esteja completa. A seleção dos elementos que vão compor a solução é feita com base em uma estratégia gulosa/aleatória.

Para formar esta solução, a fase construtiva utiliza-se de dois elementos denominados: Lista Restrita de Candidatos (LRC) e a Lista de Candidatos (LC). A LRC contém os elementos a compor o problema que possuem o melhor custo incremental, dentre a Lista de Candidatos (LC) que contém todos os elementos possíveis a integrar uma solução completa do problema.

Cada novo elemento a ser inserido na solução parcial é selecionado de forma aleatória da LRC. A LRC pode ser formada, segundo Lima Júnior (2009), de duas formas: baseada na cardinalidade da lista ou baseada na qualidade dos elementos. Para compreender melhor essas duas estratégias, considere um problema de minimização e um conjunto $LC = \{e_1, e_2, \dots, e_n\}$, onde $N = |E|$ e E é o conjunto formado por todos os possíveis elementos candidatos a fazerem parte da solução em construção (uma lista de candidatos). Seja $c(e_i)$ uma função gulosa que designa o custo incremental associado à incorporação do elemento e_i na solução corrente parcial. Representa-se por c_{min} e c_{max} , o menor e o maior custo incremental, respectivamente. Considere α um parâmetro no intervalo $[0, 1]$ que delimita o tamanho da LRC.

- **Estratégia de construção da LRC baseada na cardinalidade da lista**

A LRC é formada pelos primeiros p elementos com os melhores custos incrementais, selecionados a partir da lista de candidatos (LC) ordenada em forma decrescente, de

acordo com a função gulosa $c(e)$, em que p é um parâmetro calculado pela equação:

$$p = 1 + \alpha(N - 1) \quad (3.1)$$

no qual N é o número total de elementos candidatos.

- **Baseada na qualidade dos elementos**

Neste caso, a LRC é construída considerando os valores associados a cada elemento $c \in LC$, que respeitem a condição imposta pela equação 3.2:

$$LRC = \{e \in LC \mid c(e) \leq c_{min} + \alpha(c_{max} - c_{min})\} \quad (3.2)$$

possuindo, desta forma, tamanho variável.

3.1.1.2 Busca Local

Uma vez construída a solução inicial, a fase de busca local é aplicada com o objetivo de encontrar um mínimo local. Essa busca explora a vizinhança da solução corrente por meio de operadores de movimento, selecionando iterativamente soluções adjacentes que promovam melhoria na função objetivo.

O processo completo é repetido por um número determinado de iterações, ou até que um critério de parada seja satisfeito, como tempo limite ou estagnação das soluções. Ao final das repetições, a melhor solução encontrada durante o processo é retornada como resultado.

O desempenho do GRASP depende diretamente do equilíbrio entre os mecanismos de exploração (aleatoriedade na fase construtiva) e exploração (refinamento pela busca local). Tal equilíbrio permite que a metaheurística escape de mínimos locais e explore regiões promissoras do espaço de busca, sendo especialmente eficaz em problemas nos quais o espaço de soluções viáveis é grande e a função objetivo apresenta múltiplos ótimos locais.

3.1.2 Q-Learning

Como mencionado anteriormente, o método de aprendizado por reforço *Q-learning* é utilizado neste trabalho com o objetivo de permitir que uma meta-heurística inicie seu

processamento a partir de soluções iniciais de alta qualidade. O *Q-learning*, introduzido em (WATKINS, 1989), é um método livre de modelo, baseado em um processo de decisão de Markov com tempo finito e discreto, e em mapeamentos de estados para ações, ou seja, em políticas ou regras para decidir o que fazer a partir do conhecimento disponível sobre um determinado estado. Uma política deve ser definida sobre todo o espaço de estados, a fim de especificar qual ação deve ser tomada em cada situação possível.

Seguindo essa formalização, cada item demandado d_i é representado por um estado e_i a ser visitado por um agente de aprendizado, com $i = 1, \dots, n$. Uma ação caracteriza a escolha de cortar um item a partir de uma chapa, implicando na transição do estado atual para o estado associado ao item em questão. A matriz de valores- Q (*Q-values*) tem sempre dimensão $n \times n$. Além disso, o *Q-learning* atualiza iterativamente os valores- Q para os pares estado-ação por meio da seguinte expressão:

$$Q[e_i, a_{ij}] = (1 - \alpha) \cdot Q[e_i, a_{ij}] + \alpha \cdot \left(r_j + \gamma \cdot \left(\max_{\substack{1 \leq k \leq n \\ k \neq i, j}} Q[e_j, a_{jk}] \right) \right), \quad (3.1)$$

em que e_i representa o estado atual, a_{ij} é a ação realizada do estado e_i para o estado e_j , r_j é a recompensa imediata recebida pela execução da ação a_{ij} , a_{jk} representa as ações possíveis no estado e_j , γ é o fator de desconto ($0 \leq \gamma < 1$) e α é o fator de aprendizado ($0 < \alpha < 1$).

O *Q-learning* é considerado um dos avanços mais importantes no aprendizado por reforço, pois sua convergência não depende da política aplicada (WATKINS, 1989). Foi o primeiro método de aprendizado por reforço a apresentar forte evidência de convergência. Watkins (WATKINS, 1989) demonstrou que, se cada par estado-ação (e_i, a_{ij}) for visitado um número infinito de vezes (usando um valor suficientemente pequeno para o parâmetro α), o valor $Q[e_i, a_{ij}]$ converge com certeza para o valor ótimo Q^* .

3.1.2.1 *Q-Learning: Fundamentos e Fundamentação*

3.1.2.1.1 Processo de Decisão de Markov

O Q-Learning fundamenta-se no conceito de *Processo de Decisão de Markov* (MDP), definido pelo tuplo:

$$\mathcal{M} = (\mathcal{S}, \mathcal{A}, P, R, \gamma),$$

onde:

- $\mathcal{S}$ é o conjunto de estados;
- $\mathcal{A}$ é o conjunto de ações;
- $P(s' | s, a)$ é a probabilidade de transição do estado s para s' dado a ;
- $R(s, a, s')$ é a recompensa imediata após a transição (s, a, s') ;
- $\gamma \in [0, 1)$ é o fator de desconto, ponderando recompensas futuras em relação às imediatas.

A propriedade de Markov assegura que a transição e recompensa dependem apenas do estado atual e da ação, não de históricos anteriores (OLIVEIRA, 2021; LODI *et al.*, 2002). O objetivo é aprender uma política $\pi : \mathcal{S} \rightarrow \mathcal{A}$ que maximize o retorno descontado esperado:

$$G_t = \sum_{k=0}^{\infty} \gamma^k R_{t+k+1}.$$

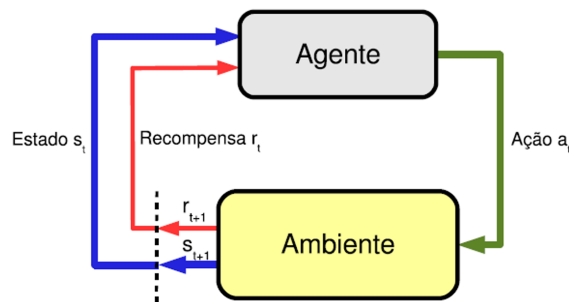


Figura 2 – Diagrama de um Processo de Decisão de Markov (estados, ações, transições e recompensas).

3.1.2.2 *Q-Learning: definição e elementos*

O Q-Learning é um algoritmo *off-policy* de aprendizado por reforço, que estima a função valor-ação ótima:

$$Q^*(s, a) = \max_{\pi} \mathbb{E}[G_t | S_t = s, A_t = a, \pi].$$

A cada transição observada, aplica-se a atualização:

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha [r_{t+1} + \gamma \max_{a'} Q(s_{t+1}, a') - Q(s_t, a_t)],$$

onde:

- α é a taxa de aprendizado;
- r_{t+1} é a recompensa imediata;
- $\max_{a'} Q(s_{t+1}, a')$ representa o valor ótimo estimado do próximo estado.

A política de seleção de ações, comumente ϵ -greedy, pode ser diferente daquela derivada de Q (off-policy), auxiliando na exploração e aprendizado conjunto.

O algoritmo Q-learning utiliza alguns parâmetros essenciais e estratégias de exploração, fundamentais para o treinamento, são eles:

- **Taxa de aprendizado (α)** - Controla a influência de novas observações sobre Q . Valores inadequados podem resultar em convergência lenta ou instabilidade.
- **Fator de desconto (γ)** - Regula a importância de recompensas futuras. Valores próximos de 0 focam no imediato; próximos de 1 valorizam resultados de longo prazo.

O balanço entre explorar e explorar com a política ϵ -greedy é fundamental. Além disso, estratégias adaptativas como a VDBE ajustam ϵ dinamicamente com base nas variações dos valores Q (TOKIC; PALM, 2011).

O Q-learning se mostra eficiente em otimização combinatória, na tese de Oliveira (2021), são aplicados algoritmos baseados em Q-Learning para problemas de otimização multiobjetivo, destacando sua versatilidade em cenários reais complexos. Além disso, a abordagem de aprendizado por reforço é utilizada na modelagem dos processos de seleção de soluções, estruturando estados, ações e recompensas que respeitam restrições do problema (ex.: capacidade de mochilas).

3.1.3 *HHDHeuristic*

A *HHDHeuristic* apresentada por Nascimento et al. 1999 é uma heurística construtiva voltada à geração de soluções viáveis para o Problema de Corte Bidimensional Guilhotinado (PCBG). Ele tem sido amplamente utilizada como base para construção de soluções iniciais em métodos híbridos envolvendo meta-heurísticas como GRASP e algoritmos genéticos, dada sua simplicidade e eficácia computacional.

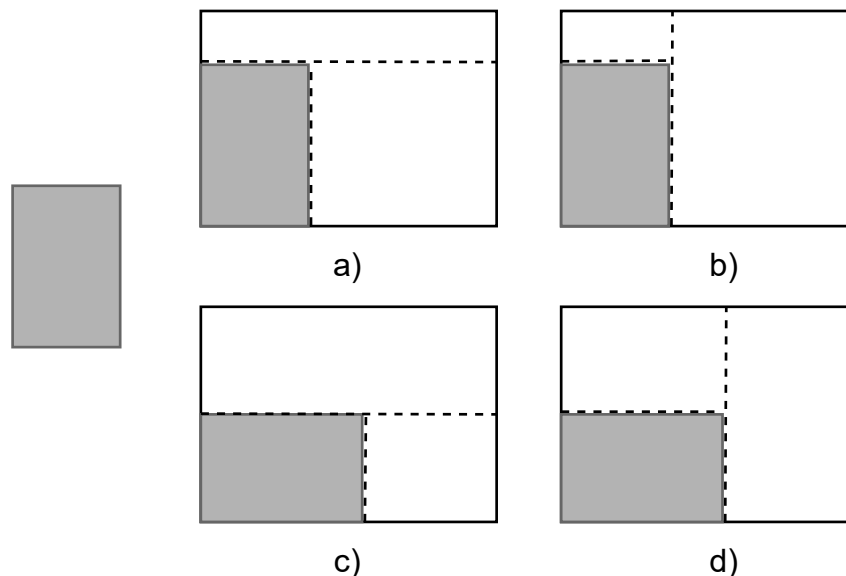
Seu funcionamento baseia-se em uma abordagem sequencial de alocação de peças retangulares sobre chapas de dimensões padronizadas, respeitando restrições de corte do tipo guilhotinado.

Para isso, o algoritmo mantém duas listas ordenadas: uma contendo os itens a

serem cortados, organizados em ordem não crescente de área, denotada por $\sigma(D)$, e outra com as sobras disponíveis, ordenadas em ordem não decrescente de área denotada por $\sigma(S)$. Esta estratégia visa favorecer o aproveitamento das sobras e priorizar o atendimento dos itens de maior área.

Antes de executar qualquer corte, o HHDHeuristic analisa todas as possibilidades de cortes possíveis de uma sobra para um item. Na posição original do item, o mesmo é colocado em uma das extremidades da esquerda, superior ou inferior. E em seguida é verificado que dimensões as sobras geradas teriam, para um corte na horizontal seguido de um na vertical (Figura 3a), gerando duas sobras; ou um corte na vertical e outro na horizontal (Figura 4b), gerando outras duas sobras. Se habilitado o algoritmo também pode rotacionar um item em 90° para gerar outras duas possibilidades, um corte horizontal e outro vertical com o item rotacionado (Figura 3c); e finalmente um corte vertical e outro na horizontal, também com o item rotacionado (Figura 3d)

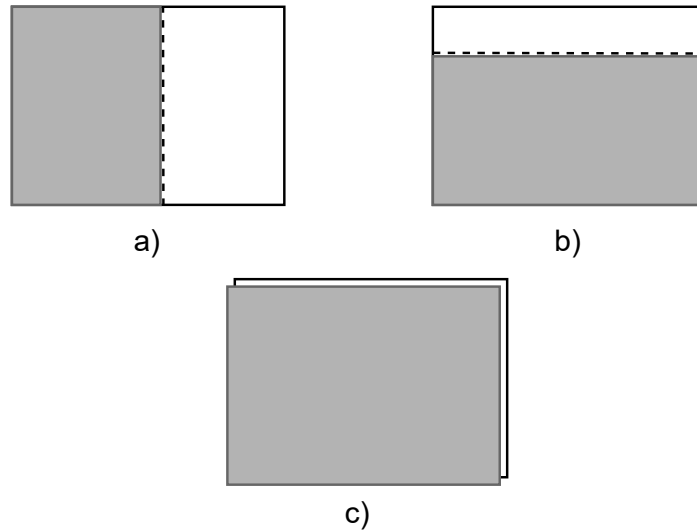
Figura 3 – Cortes possíveis



Fonte: Autoria própria. (2025)

Ainda existem três casos especiais: quando o item possui a altura igual à sobra, nesse caso é realizado apenas um corte na vertical (Figura 4a); já se a largura do item for igual à sobra, apenas um corte na horizontal é necessário (Figura 4b), ambas as situações gerarão apenas uma sobra. Finalmente, se o item tiver a mesma largura e altura da sobra analisada, nesse caso não há sobras (Figura 4c).

Figura 4 – Casos especiais



Fonte: Autoria própria. (2025)

Como existem até quatro possibilidades de cortes possíveis, incluindo as rotações, Nascimento et al. 1999 descrevem no artigo que apresentou o HHDHeuristic três estratégias de corte que geram os melhores cenários de corte, ou seja, quais cortes necessários para entregar as melhores sobras, as estratégias são:

- Maximizar a área da menor sobra resultante (se não nula);
- Maximizar a área da maior sobra resultante (se não nula);
- Maximizar o menor lado (largura ou altura) das sobras resultantes (se não nulas).

A Figura 5, representa o comportamento da heurística em três iterações, de forma que a cada iteração, a menor sobra disponível em $\sigma(S)$ é associada ao maior item disponível em $\sigma(D)$. Para esse exemplo, foi utilizada a estratégia de maximizar a área da maior sobra resultante.

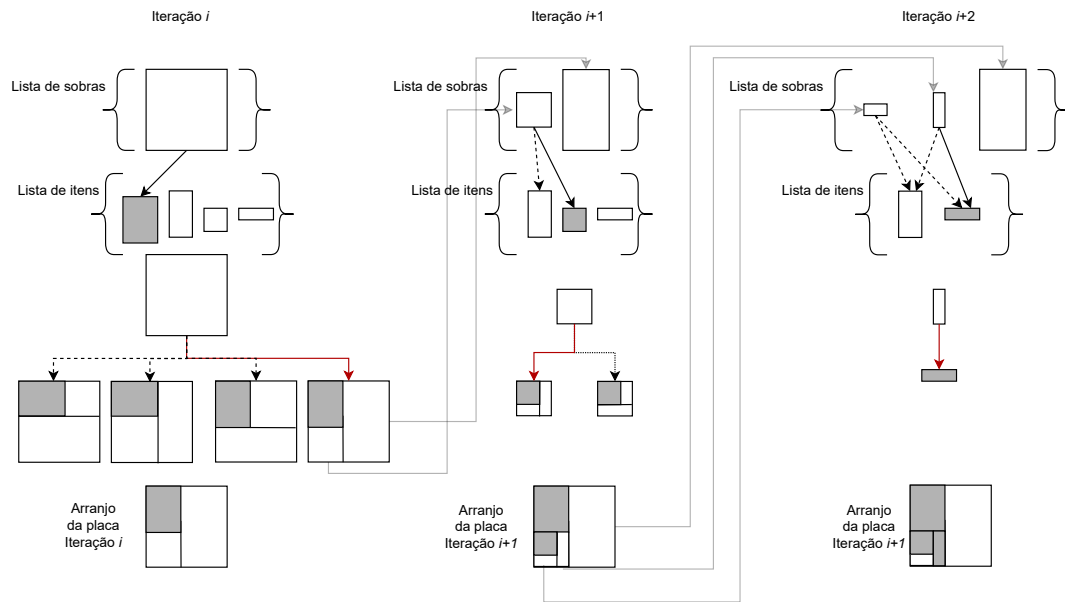
Nesse caso, cada corte gera até duas novas sobras, as quais são reinseridas em $\sigma(S)$ em posições apropriadas, mantendo a ordenação crescente por área das sobras. Quando não há item que se encaixe na sobra iterada, o teste é realizado com a próxima sobra e os demais itens. Caso a lista de sobras $\sigma(S)$, não contenha sobras maiores que algum dos itens restantes a serem cortados, uma nova chapa é inserida em $\sigma(S)$

A heurística se encerra quando $\sigma(D)$ é vazio, ou seja, todos os itens foram cortados.

A saída do algoritmo é o conjunto de padrões de corte $\mathcal{C} = \langle c_1, c_2, \dots, c_r \rangle$, onde cada c_i é definido como uma quádrupla $\langle d, s, rot, orient \rangle$, indicando:

- d – item cortado;

Figura 5 – HHD Heuristic



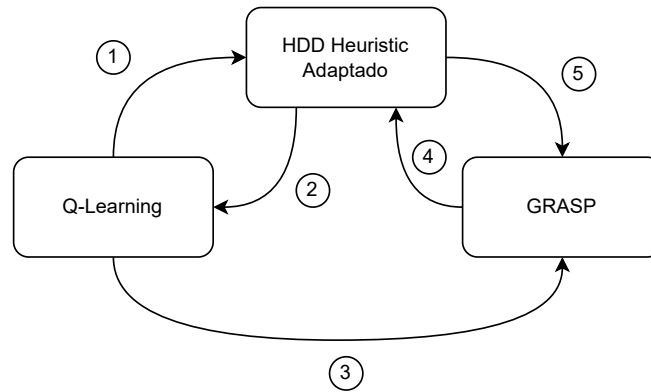
Fonte: Adaptado de Silva(2009).

- s – sobra utilizada para cortar d ;
- $rot \in \{\text{True}, \text{False}\}$ – flag indicando se d foi rotacionado;
- $orient \in \{\text{none}, h, v, hv, vh\}$ – tipo de corte necessário.

4 ABORDAGEM COMPUTACIONAL

A abordagem computacional proposta neste trabalho fundamenta-se na continuidade da abordagem híbrida desenvolvida por Rebouças (2023), que integra o algoritmo de Aprendizado por Reforço Q-learning com uma metaheurística base nesse caso o GRASP para a resolução do Problema de Corte Bidimensional Guilhotinado Não-Estagiado (PCBGNE). Nesta nova proposta, a ideia é manter a estrutura geral da abordagem anterior, conforme representado na Figura , porém foi introduzida uma modificação significativa na função de recompensa utilizada pelo Q-learning, com o objetivo de aprimorar o aprendizado da política de corte.

Figura 6 – Metodologia Híbrida



Fonte: Adaptado de Rebouças(2023).

No processo de implementação do *Q-learning* também será reavaliada a função de retorno ou recompensa apresentada por Rebouças (2023), abaixo:

$$r_i = \frac{\text{area do Item}_i}{\text{area da placa}}$$

A hipótese é que a mesma poderia estar gerando uma matriz Q de forma escalar, uma vez que a razão apresentada tende a ter valores maiores quando itens possuam áreas próximas das placas utilizadas no problema. Logo, é possível que a abordagem do trabalho anterior priorize somente os estados com os maiores itens, o que seria um comportamento guloso e inviabilizaria, tornando irrelevante, a utilização do Q-learning na metodologia da solução.

Recapitulando, como o problema do corte bidimensional guilhotinado, em questão, é tratado como um problema de decisão sequencial, o mesmo pode ser representado como uma quintupla $M = \{T, S, A, R, P\}$. Sendo:

- T : Conjunto de instantes de decisão, correspondente aos itens que farão parte da solução;
- S : O conjunto de possíveis estados, no qual cada decisão sequencial pode levar;
- A : O conjunto de possíveis ações durante o processo de decisão sequencial, nesse caso toda ação $A_{x,y}$ realizada a partir de um estado S_x para um estado S_y desabilita o estado S_x , por exemplo;
- R : Função de recompensa proposta a seguir na seção 4.1 .
- P : A probabilidade de transição entre estados, para o problema assumiu-se que a mesma é sempre 1, ou seja, ela sempre acontece sendo determinística.

4.1 Função de recompensa tardia

A nova função de recompensa foi desenhada para refletir mais adequadamente o objetivo do problema, priorizando estratégias que maximizem o aproveitamento das chapas e minimizem a geração de desperdício. Acredita-se que esse ajuste pode impactar positivamente na qualidade das soluções iniciais produzidas pela fase construtiva, influenciando diretamente o desempenho global dos algoritmos.

- Recompensa Simples

$$r_{simples} = -\Delta_{chapas} + \Delta_{Area\ sobras}$$

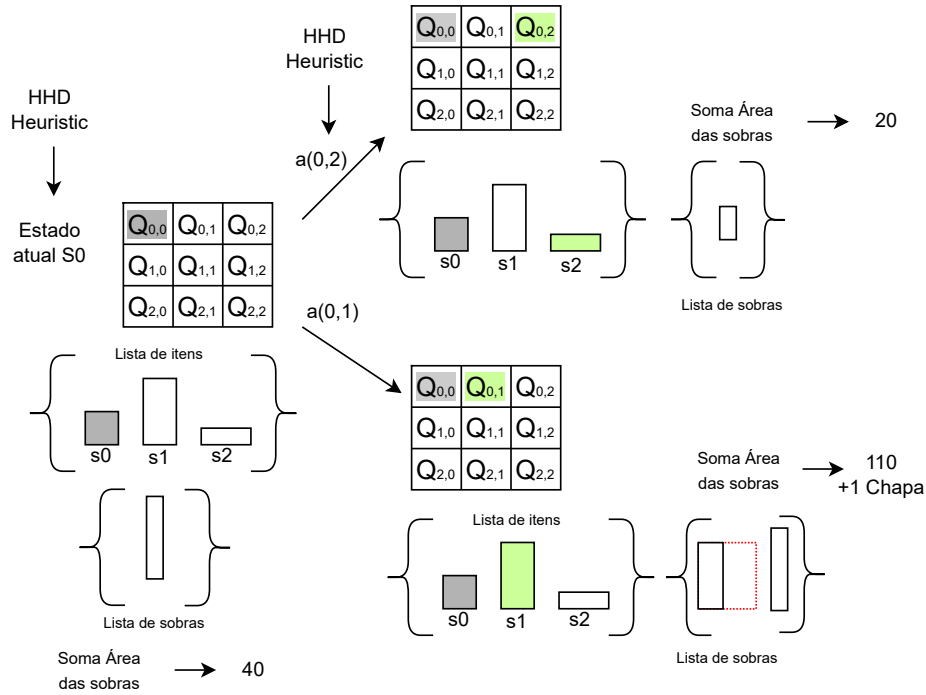
As iterações de cada episódio são realizadas primeiro, gerando um possível *layout* de corte no estado atual, usando o HHD Heuristic adaptado. Dessa forma obtêm-se as sobras e se seria necessário a utilização de alguma nova chapa para realizar um possível corte. Para a próxima ação o HHD Heuristic é utilizado novamente para que se possa obter também um novo *layout* de corte, assim como as sobras e as chapas

A ideia é que haja uma penalização maior na função de recompensa, caso haja utilização de novas chapas, por isso o Δ de chapas negativo. O mesmo é calculado da seguinte forma:

$$\Delta_{chapas} = Total_{Atual} - Total_{futura}$$

Já o Δ da área das sobras é calculado entre a diferença da soma das áreas das sobras resultantes de uma ação futura menos a soma das áreas das sobras que se tem atualmente,

Figura 7 – Obtendo delta das chapas e da área das sobras



Fonte: Autoria própria (2025).

conforme abaixo:

$$\Delta_{Area\ sobras} = \sum_{n=1}^N Area_n - \sum_{m=1}^M Area_m$$

$$n \in \sigma(S) \text{ futuras} ; m \in \sigma(S) \text{ atuais}$$

Finalmente no final de cada construção de episódio, quando a última ação do Q-Learning é executada, ou seja quando todos os itens são atendidos no plano de corte a recompensa muda, para:

- Recompensa Terminal

$$r_{terminal} = -total_{chapas}$$

Nesse caso, uma penalização maior é aplicada, caso o episódio utilize mais de uma chapa para acomodar os itens no *layout* do corte.

Além disso durante o treinamento esperasse seguir uma metodologia incremental semelhante a apresentada por Even-Dar (2004), a qual durante a execução do algoritmo

uma vez que o ε do ε_{greed} decaia, em um certo valor de ε , o α passe a decair mais rápido, forçando o refinamento da solução, com menor exploração.

4.2 Base de dados

Quanto a base de dados, foi utilizado um conjunto clássico de instâncias amplamente adotado na literatura do Problema de Corte Bidimensional. As instâncias foram originalmente propostas por Berkey e Wang (1987), e posteriormente organizadas e estendidas por Lodi, Martello e Vigo (1999). Esse conjunto compreende dez classes, cada uma contendo cinquenta instâncias, totalizando 500 problemas distintos.

Cada classe é subdividida em cinco grupos de acordo com o número de itens, variando entre 20, 40, 60, 80 e 100 itens por instância. As seis primeiras classes foram geradas aleatoriamente com base no método descrito por Berkey e Wang (1987), enquanto as quatro classes restantes foram produzidas conforme os critérios estabelecidos por Martello e Vigo (1998).

Alem disso, para caráter comparativo, o estado da arte aproveita a proposta de Dell’Amico, Martello e Vigo (2002), que utilizaram um método pseudo-polinomial para a obtenção de limites inferiores aplicáveis ao Problema de Corte Bidimensional. A qual baseia-se na conversão dos itens retangulares em quadrados por meio de cortes sucessivos. Esses quadrados são então agrupados, e o número mínimo de placas necessárias para acomodar todos os itens é utilizado como limite inferior da instância.

Posteriormente, Clautiaux, Jouglet e El Hayek (2007) aprimoraram esse procedimento e desenvolveram novas estimativas de limites inferiores, conforme descrito em 2007. Tais métodos serão utilizados neste trabalho como referência para avaliação comparativa dos algoritmos propostos.

4.3 Avaliação

Para avaliar o impacto da função de recompensa e das estratégias de busca adotadas, serão testadas três variantes do algoritmo então batizado por Rebouças (2023), GRASP-Learning, todas utilizando o Q-learning como componente da fase construtiva:

- **GRASP-Learning com busca local simples:** aplica uma estratégia de melhoria baseada em vizinhança básica com 2-opt, com movimentos elementares de troca

ou inserção, uma vez encontrado uma melhor solução na busca local o algoritmo é encerrado;

- **GRASP-Learning com busca local aperfeiçoada:** semelhante a anterior, porém buscando melhorias em um intervalo específico de iterações durante a busca local;
- **GRASP-Learning com *Iterated Local Search* (ILS):** utiliza uma busca local incorporada a uma estratégia de perturbação sistemática, que permite escapar de ótimos locais e intensificar a exploração de múltiplas regiões da paisagem de soluções.

Todos os métodos serão aplicados a um conjunto padrão de instâncias do PCBGNE, permitindo comparações diretas quanto à qualidade das soluções geradas, à robustez dos métodos e ao tempo computacional requerido. As execuções serão repetidas diversas vezes para cada instância, assegurando confiabilidade estatística nos resultados.

A comparação dos algoritmos será realizada por meio de métricas como: número de chapas utilizadas, tempo médio de execução, desvio padrão e qualidade média das soluções. Os resultados serão apresentados em tabelas e gráficos comparativos, e interpretados com base nos objetivos definidos neste estudo.

4.3.1 *GRASP-Learning com busca local simples*

O processo de execução do algoritmo híbrido utilizando busca local proposto nesta dissertação tem como meta reproduzir o experimento proposto por Rebouças (2021), porém com a nova função de recompensa tardia. A Figura 8 ilustra todo o fluxo computacional realizado, desde a leitura da instância até a obtenção da solução final.

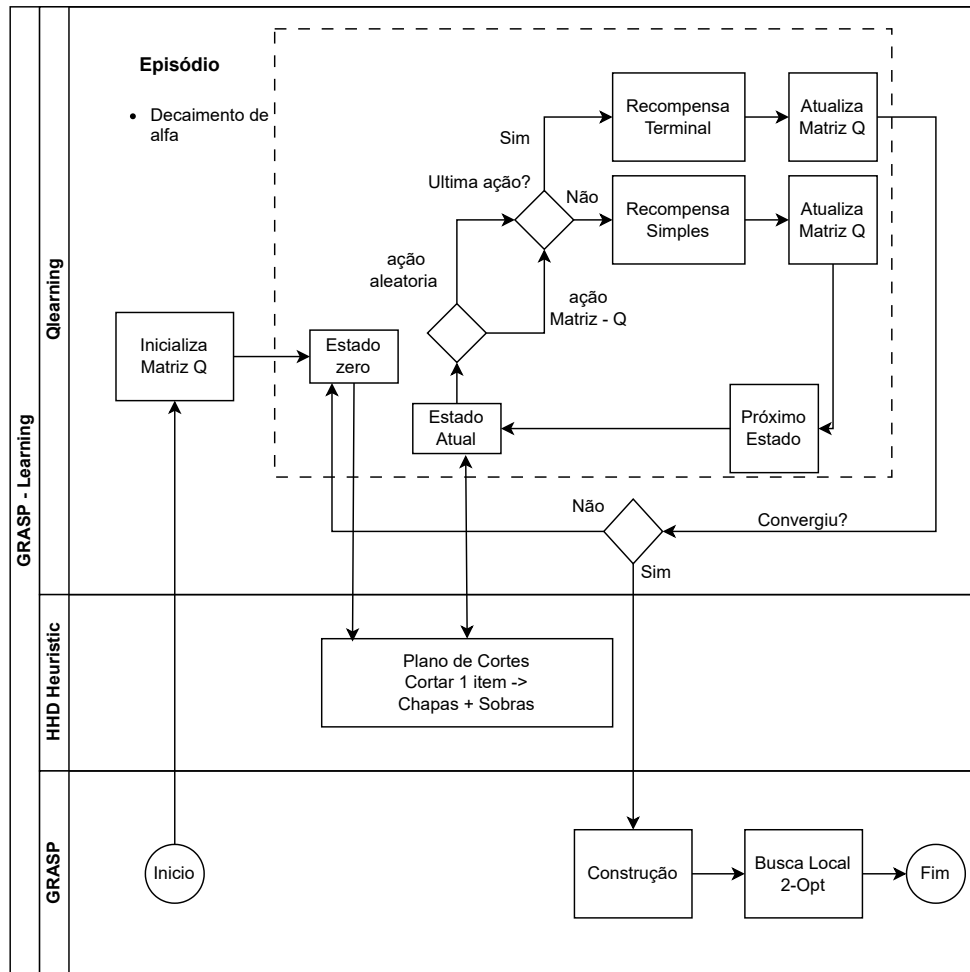
4.3.1.1 *Leitura da Instância e Inicialização do Problema*

O algoritmo recebe como entrada dados da instancia que contém os parâmetros (dimensões da chapa, lista de itens e rótulo da instância) e retornando:

- Uma lista de itens a serem cortados, cada um representado por um par de dimensões (largura, altura);
- As dimensões da chapa padrão disponível para corte (largura, altura e quantidade).

Esses dados são então utilizados para instanciar o problema para a execução do Q-Learning durante a fase de construção do GRASP.

Figura 8 – Fluxograma GRASP-Learning com Busca local



Fonte: Autoria própria (2025).

4.3.1.2 Treinamento com Q-Learning

Nesta fase, o Q-Learning é executado com um limite máximo de episódios pré-definidos, porém a execução pode ser abreviado, caso haja convergência na matriz Q . Em cada episódio, é construída uma sequência de itens a serem cortados, baseada na política ϵ -gulosa sobre uma matriz Q de valores de estado-ação. A estrutura geral da execução do Q-Learning segue as seguintes etapas:

- A matriz Q é inicializada com zeros para todos os pares (s, a) possíveis.
- Um estado inicial s_0 é definido para o primeiro item da lista $\sigma(D)$ é selecionado para início do episódio
- Para cada etapa do episódio:
 - Com probabilidade ϵ , uma ação aleatória é escolhida (exploração);

- Caso contrário, a ação com maior valor $Q(s, a)$ é selecionada (exploração).
- A recompensa simples para a transição (s, a) é calculada com base na área do item cortado em relação à área total da chapa.
- A matriz Q é atualizada usando a equação clássica do Q-Learning:

$$Q(s, a) \leftarrow Q(s, a) + \alpha \left[r + \gamma \max_{a'} Q(s', a') - Q(s, a) \right]$$

- O processo continua até que todos os itens tenham sido utilizados na construção da sequência.
- Na última ação, a recompensa aplicada é a terminal, mais punitiva caso a diferença de chapas do episódio seja maior.

4.3.1.3 Convergência

Para exatidão no treinamento executado pelo Q-learning, uma das condições de paradas do treinamento foi definida ao analisar se existe convergência da matriz Q . Nesse caso durante a criação dos episódios, o teste é realizado para garantir que se em 100 vezes a diferença entre todos os elementos da matriz Q do estado atual e a matriz Q atualiza é menor ou igual a 0,0001.

4.3.1.4 Fase Construtiva

No *GRASP-Learning*, a fase construtiva é totalmente substituída pelo algoritmo de Q-learning, que atua como mecanismo guloso-aleatório adaptativo. A cada iteração do GRASP, executam-se N_{ep} episódios de Q-learning. Originalmente o GRASP utiliza a *Lista de Candidatos* (LC) e de uma *Lista Restrita de Candidatos* (LRC), formada pelos p melhores elementos segundo uma função gulosa; um candidato é então sorteado da LRC e inserido na solução parcial até seu término. No caso do GRASP-Learning, a aleatoriedade imposta pelo α não será utilizada, visto que o próprio treinamento da matriz Q já prevê a mesma.

1. Busca Local: Após a fase construtiva (seja pelo GRASP clássico ou pelo GRASP-Learning), aplica-se o operador 2-opt para explorar o entorno da solução inicial. No contexto do problema de corte guilhotinado, cada solução pode ser vista como uma sequência de cortes ou de itens montados em uma chapa. O 2-opt atua removendo duas “arestas” na ordem de montagem — isto é, dois cortes consecutivos ou não

consecutivos — e reconectando os dois segmentos produzidos de forma invertida, gerando uma nova solução vizinha. Se a nova disposição reduzir o número de chapas usadas ou diminuir o desperdício, ela é aceita; caso contrário, descarta-se e testam-se outras trocas 2-opt até que não haja mais melhorias locais. Este processo prossegue iterativamente até atingir um ótimo local, garantindo melhorias rápidas e de baixo custo computacional na solução inicial

2. ILS: Além disso, quando substitui-se, a metaheurística ILS, na busca local, eleva a busca local simples ao combinar ciclos de refinamento (via 2-opt) com perturbações controladas, permitindo escapar de ótimos locais rasos. Primeiro, obtém-se um ótimo local aplicando 2-opt sobre a solução gerada pelo GRASP-Learning. Em seguida, executa-se uma perturbação: troca-se aleatoriamente k pares de cortes ou segmentos na sequência, criando uma solução “quase ótima” distinta. Essa solução perturbada é então refinada novamente por 2-opt. Se o resultado for melhor que a incumbente, substitui-se a incumbente; caso contrário, mantém-se a melhor conhecida. Todo esse ciclo — perturbação seguida de busca local — repete-se por N_{ILS} iterações ou até convergir, balanceando exploração e exploração para alcançar soluções de qualidade superior àquelas obtidas somente com GRASP ou 2-opt isolados.

4.3.2 *Retorno da Solução Final*

Ao término das iterações da GRASP (definidas por um número máximo de repetições), o algoritmo retorna:

- A sequência de itens da melhor solução encontrada;
- O custo final da solução, medido pelo número de chapas utilizadas.

Essa saída representa o melhor plano de corte encontrado com base na combinação da construção via Q-Learning e da otimização local via GRASP.

5 AVALIAÇÃO DAS ABORDAGENS COMPUTACIONAIS

A finalidade desse capítulo é expor os resultados computacionais dos artefatos obtidos utilizando a metodologia de *Design science*, mencionados no Capítulo 1.

5.1 Experimentos Computacionais

Esta seção descreve o ambiente de desenvolvimento e de execução utilizados para obtenção dos resultados deste trabalho. Assim como discute e analisa os resultados obtidos durante esses dois processos

5.1.1 Ambiente de Desenvolvimento

Todos os algoritmos desse trabalho foram implementados na linguagem Python 3.12.1 com utilização da IDE do Visual Code Studio 1.86.2. Além disso foi criado um repositório privado no Github para o projeto, uma plataforma de hospedagem de código-fonte e controle de versão. Para execução dos algoritmos foi utilizado um computador pessoal, com o sistema operacional Windows 11 64 bits, com 16 GB de memória RAM (1333 Mhz) e um processador AMD 5600G de Clock variando de 3.5 GHz e máximo 4.4 GHz (*Boost Mode*) de 6 núcleos e até 12 *Threads*.

5.1.2 Análise dos Resultados

Nessa sessão primeiramente será discorrido os testes realizados utilizando a função de recompensa de Rebouças e a implementação do algoritmo Q-learning, deste trabalho, utilizando os parâmetros originais, assim como variações do mesmo. A fim de validar a hipótese sobre o comportamento escalar do trabalho anterior.

No trabalho anterior foram utilizados os parâmetros para o treinamento do Q-learning descritos na Tabela 2:

Tabela 2 – Parâmetros utilizados por Rebouças (2021)

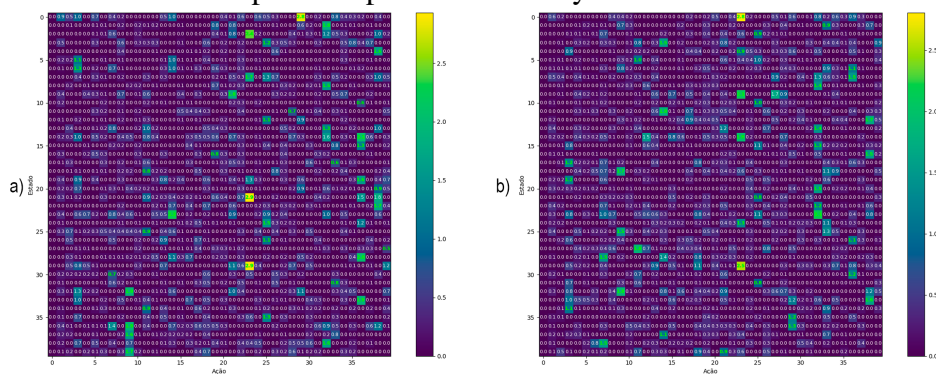
Parâmetro	Função	Valor
NEp	Número de episódios	200
α	Taxa de aprendizagem	0.9
γ	Fator de desconto	0.9
ϵ	Controle entre a gula e aleatoriedade	0.9
τ	Taxa de decaimento do ϵ	0.97

Fonte: Rebouças (2021).

Para o teste foi escolhida a instância número 464, a qual pertence a 10ª classe dos problemas, considerados mais difíceis de resolver. Ela possui 40 itens a serem recortados, utilizando chapas de 100 de altura por 100 de largura.

Entre os 40 itens, a instancia possui uma certa variedade de itens: alguns itens com grande diferença entre as dimensões como 8(altura) x 99(largura); itens que se aproximam do tamanho da chapa 100 x 100, como é o caso do item 98 x 89; e itens de dimensões menores como 19 x 2. Utilizando a instância em questão, a matriz Q final, gerada pelo treinamento foi:

Figura 9 – Matriz-Q final de dois treinamentos com 200 episódios para o método xyz



Fonte: Autoria própria (2025).

Visualmente é perceptível, que as matrizes diferem bastante, o que não deveria ser o comportamento do treinamento, visto que nos episódios finais a matriz Q deve convergir uma vez que já explorou o caráter randômico proporcionado pela política ϵ -greedy, mesmo o fator de decaimento de ϵ seja considerado alto.

A partir desse teste, realizou-se um ajuste nos parâmetros, utilizando a mesma função de recompensa, para tentar criar uma matriz Q mais definida. Os parâmetros foram:

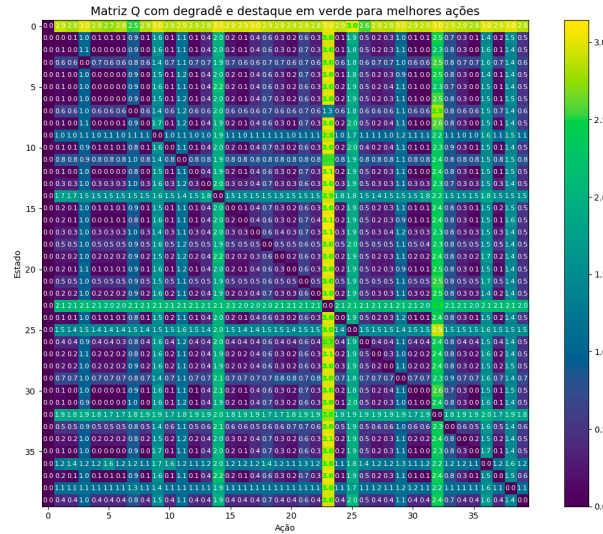
Tabela 3 – Parâmetros utilizados tentativa de convergência

Parâmetro	Função	Valor
MEp	Número máximo de episódios	15000
α	Taxa de aprendizagem	0.9
γ	Fator de desconto	0.9
ϵ	Controle entre a gula e aleatoriedade	0.9
τ	Taxa de decaimento do ϵ	0.9995

Fonte: Adaptado de Rebouças (2021).

Em seguida, como resultado, uma Matriz- Q foi obtida convergindo na faixa dos 9000 episódios e realizada a análise para averiguar a escolha das ações a partir da convergência.

Figura 10 – Matriz Q resultante da convergência com x para a próxima ação e y para a ação atual

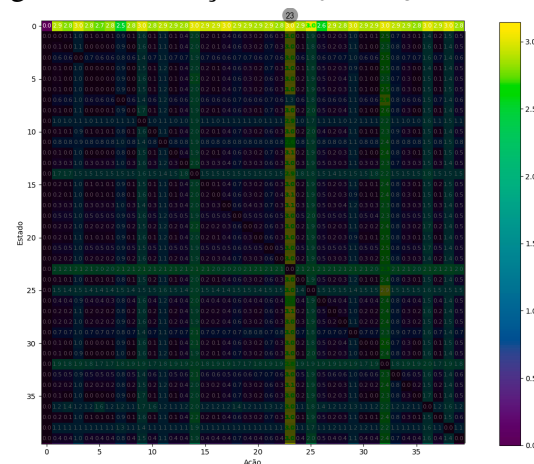


Fonte: Autoria própria (2025).

Baseado na matriz Q treinada, acima na Figura 10, quatro passos de decisão e a respectiva sequência de itens escolhidos para um suposto layout de corte podem ser selecionadas.

Na primeira iteração um item 0 (estado 0) é o primeiro do layout do corte, representado pela primeira linha da matriz Q. O próximo estado o S_{23} que representa a escolha do item 23 para ser cortado, conforme Figura 11

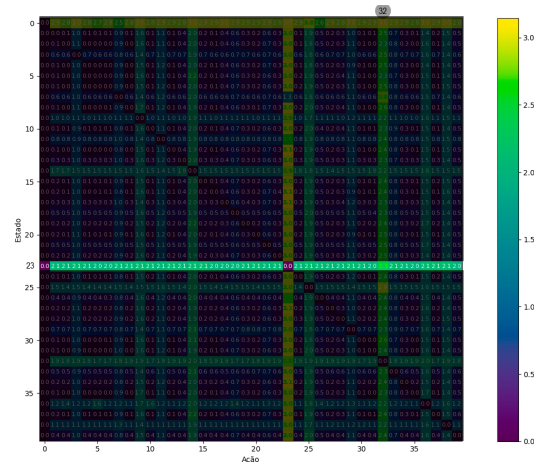
Figura 11 – Iteração 1 - $S_0 \rightarrow S_{23}$



Fonte: Autoria própria (2025).

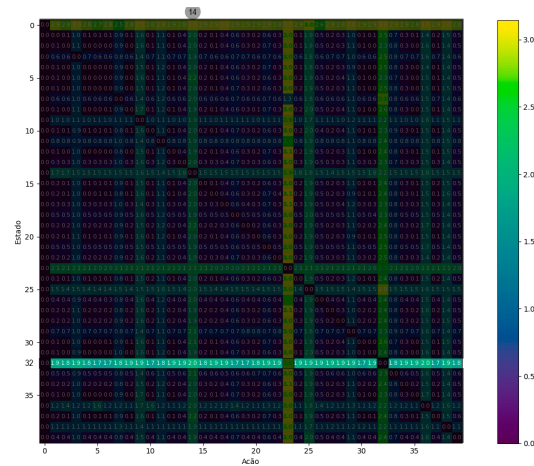
Na segunda iteração, conforme Figura 12 partindo do estado S_{23} , na 23ª linha o maior valor na linha é o da coluna 32, logo o item 32, estado 32.

Na terceira iteração, conforme Figura 13 partindo do estado S_{32} , na 32ª linha o

Figura 12 – Iteração 2 - $S_{23} \rightarrow S_{32}$ 

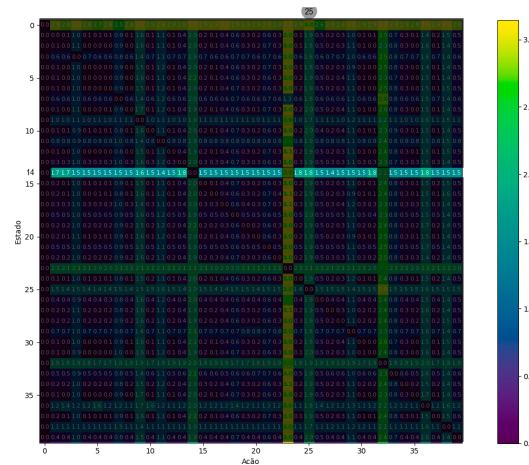
Fonte: Autoria própria (2025).

maior valor na linha é o da coluna 14, logo o item 14, lembrando que as ações anteriores em um episódio não pode ser repetida.

Figura 13 – Iteração 3 - $S_{32} \rightarrow S_{14}$ 

Fonte: Autoria própria (2025).

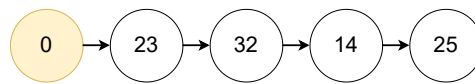
Finalmente, na quarta iteração, conforme Figura 14 partindo do estado S_{14} , na 14ª linha o maior valor na linha seria o da coluna 25, logo o item 25.

Figura 14 – Iteração 4 - $S_{14} \rightarrow S_{25}$ 

Fonte: Autoria própria (2025).

Até então, o layout parcial de corte iniciando no item 0 utiliza os seguintes itens, conforme Figura 15

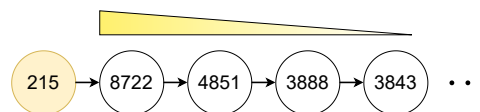
Figura 15 – Layout de corte parcial



Fonte: Autoria própria (2025).

Ao substituir o identificador dos itens por sua respectiva área. Se retido o item 0, inicial. É perceptível que o treinamento gerou um comportamento escalar, de maneira gulosa, buscando sempre utilizar os maiores itens, conforme Figura 16

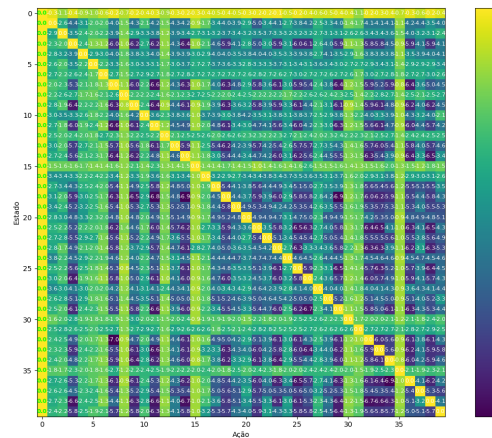
Figura 16 – Áreas de itens do layout parcial de corte



Fonte: Autoria própria (2025).

Ao utilizar a função de recompensa tardia, na mesma instancia, obteve-se uma matriz Q resultante, não escalar, visto que as ações escolhidas inicialmente possuem pesos semelhantes, conforme Figura 17, abaixo.

Figura 17 – Matriz Q resultante com função de recompensa tardia



Fonte: Autoria própria (2025).

Segue a análise comparativa de todos os algoritmos, resumida na Tabela 4, e posteriormente complementada com os resultados do estado da arte apresentados na apenas utilizando as instancias de classe 10 para o teste Tabela 5

Tabela 4 – Comparativo entre os resultados de todos os algoritmos

Metaheurística	Melhor	Pior	Média	Desvio Padrão	Tempo(s)
GRASP	7051	7112	7058,00	0,033	16,89
GRASP-Learning	7045	7068	7057,65	0,011	33,79
GRASP-Learning-Delayed (BL)	7047	7076	7076,5	0,014	48843,821
GRASP-Learning-Delayed (BLA)	7045	7069	7057	0,012	52831,562
GRASP-Learning-Delayed (ILS)	7047	7076	7060,5	0,014	78562,82

Fonte: Autoria própria (2025).

Tabela 5 – Comparativo entre os resultados da Classe 10 para os algoritmos propostos e resultados da literatura

Métodos	Quantidade de itens					Ano
	20	40	60	80	100	
FBS2	42	75	103	131	163	1999
FC2	41	74	101	130	163	1999
BuscaTabu-Lodi	43	75	104	130	166	2002
GLS	42	74	102	130	162	2003
HBP	42	74	102	130	160	2003
WA	43	74	102	129	159	2009
GRASP-Velasco	42	74	100	129	159	2010
MS-LGFi	42	74	101	129	160	2012
EA-LGFi	42	74	101	128	160	2012
FFIH	42	73	101	130	161	2013
BFIH	42	73	101	130	160	2013
CFIH	43	73	101	129	159	2013
FFIH+J4	42	73	101	130	161	2013
BFIH+J4	42	73	101	129	160	2013
CFIH+J4	42	73	101	129	159	2013
Time assíncrono - Mariano (2014)	41	73	100	127	159	2014
Time assíncrono - Silva	41	73	99	126	158	2017
BRKGA-Learning	41	73	100	130	162	2021
GRASP-Learning	41	73	99	126	158	2021
GRASP-Learning-Delayed (BL)	41	75	102	130	165	2025
GRASP-Learning-Delayed (BLA)	41	74	100	130	163	2025
GRASP-Learning-Delayed (ILS)	41	73	101	130	164	2025
Limite inferior-CJH	39	70	95	122	153	2007

Fonte: Adaptado de diversas fontes da literatura e resultados obtidos (2025).

6 CONCLUSÕES

Analisando as questões de pesquisa deste trabalho, foi possível contribuir com o estado da arte no que se refere a abordagem de treinamento utilizado no Q-Learning, especialmente na elaboração de uma nova função de recompensa e parametrização para o treinamento, assim como o desenvolvimento de três algoritmos que agora poderão servir como referência ou *benchmark* para trabalhos futuros.

Ainda referente as questões técnicas, de fato existem parametrizações mais efetivas para o problema, como foi constatado. Ficou evidente no problema que o desenho de uma função de recompensa pode levar a um comportamento escalar, o que acabou tornando o treinamento de máquina no mínimo irrelevante, visto que uma simples ordenação escalar, poderia resultar em uma mesma resposta.

Durante o desenvolvimento deste trabalho, foram identificados alguns desafios relevantes que impactaram tanto o andamento da pesquisa quanto a obtenção de resultados mais robustos. O primeiro deles foi a ausência de um repositório original contendo os códigos do trabalho anterior utilizado como referência. Essa limitação exigiu a reconstrução de boa parte da implementação a partir da descrição teórica e de informações fragmentadas, o que demandou tempo adicional e aumentou a possibilidade de variações em relação à metodologia original. Outro fator crítico foi o tempo de treinamento do algoritmo, considerando que o processo de convergência requer, no mínimo, cerca de 8.000 episódios para apresentar resultados satisfatórios. Essa característica impôs restrições ao número de experimentos que puderam ser executados, limitando a exploração de diferentes parametrizações e ajustes finos do modelo.

6.1 Trabalhos futuros

Como continuação deste estudo, diversos aprimoramentos podem ser implementados a fim de superar as limitações identificadas e ampliar o potencial do método proposto. Uma possibilidade consiste na construção de um algoritmo Q-Learning capaz de iniciar o processo de aprendizado a partir do estado inicial de qualquer item pertencente à lista $\sigma(D)$, permitindo maior flexibilidade e abrangência na exploração dos estados possíveis. Além disso, recomenda-se aprimorar a função de recompensa, de forma a torná-la mais sensível a variações intermediárias de desempenho, incentivando a convergência para soluções de

melhor qualidade. Outra proposta envolve a utilização de técnicas de paralelismo, em que cada iteração do treinamento seja inicializada com um estado distinto, possibilitando posteriormente a unificação das matrizes Q resultantes. Essa abordagem pode acelerar o processo de aprendizado e aumentar a diversidade de experiências utilizadas na atualização dos valores Q . Por fim, sugere-se a realização de testes de aceleração utilizando CPU e GPU, especialmente para os casos em que o paralelismo for aplicado, visando reduzir significativamente o tempo de processamento e permitir experimentos com construções de matriz Q com seções de episódios que comecem de outros estados.

REFERÊNCIAS

- BENGIO, Y.; LODI, A.; PROUVOST, A. Machine learning for combinatorial optimization: a methodological tour d'horizon. **European Journal of Operational Research**, Elsevier, v. 290, n. 2, p. 405–421, 2021.
- BERKEY, J. O.; WANG, P. Y. Two dimensional finite bin-packing algorithms. **Journal of the Operational Research Society**, Palgrave Macmillan, v. 38, n. 5, p. 423–429, 1987.
- BLUM, C.; ROLI, A. Metaheuristics in combinatorial optimization: Overview and conceptual comparison. **ACM Computing Surveys**, ACM, v. 35, n. 3, p. 268–308, 2003.
- CLAUTIAUX, F.; HAYEK, A. E.; JOUGLET, A. New lower bounds for bin-packing and strip-packing problems. In: **Proceedings of the International Network Optimization Conference (INOC)**. Spa, Belgium: [s.n.], 2007.
- DELL'AMICO, M.; MARTELLO, S.; VIGO, D. A lower bound for the two-dimensional bin packing problem. **Discrete Applied Mathematics**, v. 118, n. 1–2, p. 13–24, 2002.
- DYCKHOFF, R. A typology of cutting and packing problems. **European Journal of Operational Research**, v. 44, n. 2, p. 145–159, 1990.
- EVEN-DAR, E.; MANNOR, S. Convergence of optimistic and incremental q-learning. In: **Advances in Neural Information Processing Systems**. [S.l.: s.n.], 2004. v. 17.
- GILMORE, P. C.; GOMORY, R. E. Multistage cutting stock problems of two and more dimensions. **Operations Research**, v. 13, n. 1, p. 94–120, fev. 1965.
- GONCALVES, J. R.; RESENDE, M. G. C. Biased random-key genetic algorithms for combinatorial optimization. **Journal of Heuristics**, v. 17, n. 4, p. 487–525, 2011.
- JIANG, Y.; CAO, Z.; ZHANG, J. Solving 3d bin packing problem via multimodal deep reinforcement learning. In: **Proc. AAMAS 2021**. [S.l.: s.n.], 2021.
- LEE, B.; NAM, C. A hierarchical bin packing framework with dual manipulators via heuristic search and deep reinforcement learning. **arXiv preprint**, 2025.
- LODI, A.; MARTELLO, S.; VIGO, D. Heuristic and metaheuristic approaches for a class of two-dimensional bin packing problems. **INFORMS Journal on Computing**, v. 11, n. 4, p. 345–357, 1999.
- LODI, A.; MARTELLO, S.; VIGO, D. Two-dimensional packing problems: A survey. **European Journal of Operational Research**, Elsevier, v. 141, n. 2, p. 241–252, 2002.
- MARIANO, M. F. D. S. **Uma metaheurística A-teams aplicada ao problema do corte bidimensional guilhotinado**. Dissertação (Dissertação (Mestrado)) — Universidade do Estado do Rio Grande do Norte, 2014.
- MARTELLO, S.; VIGO, D. Exact solution of the two-dimensional finite bin packing problem. **Management Science**, v. 44, n. 3, p. 388–399, 1998.
- MURDIVIEN, S. A.; UM, J. Boxstacker: Deep reinforcement learning for 3d bin packing problem in virtual environment of logistics systems. **Sensors**, v. 23, n. 15, p. 6928, 2023.

NASCIMENTO, H. A. D. d.; LONGO, H. J.; ALOISE, D. J. Uma heurística *o(mn)* para o corte bidimensional guilhotinado. In: **Anais do XXXI Simpósio Brasileiro de Pesquisa Operacional (SBPO)**. Juiz de Fora, MG, Brasil: SOBRAPO, 1999. p. 1197–1206.

OLIVEIRA, T. H. F. de. **Algoritmos de aprendizagem por reforço para problemas de otimização multiobjetivo**. Tese (Doutorado) — Universidade Federal do Rio Grande do Norte, Natal, Brazil, 2021. Tese (Doutorado em Engenharia Elétrica e de Computação).

OSMAN, I. H.; KELLY, J. P. Meta-heuristics: A literature review. **Annals of Operations Research**, Springer, v. 63, p. 511–548, 1996.

PAN, Y.; CHEN, Y.; LIN, F. Adjustable robust reinforcement learning for online 3d bin packing. In: **NeurIPS 2023**. [S.l.: s.n.], 2023.

PITOMBEIRA, R. V.; JÚNIOR, F. C. L.; OLIVEIRA, T. H. F.; ALOISE, D. J. A reinforcement learning approach to the two-dimensional guillotine cutting stock problem. In: IEEE. **Proceedings of the 2021 IEEE Congress on Evolutionary Computation (CEC)**. [S.l.], 2021. p. 1895–1902.

PUGLIESE, V. U.; FERREIRA, O. F. d. A.; FARIA, F. A. Optimizing 2d+1 packing in constrained environments using deep reinforcement learning. **ScitePress – Transactions on Engineering and Technology**, 2025. Preprint.

REBOUCAS, A. D. C. **O Problema de Corte Bidimensional Guilhotinado Não-Estagiado: Uma Abordagem via Metaheurísticas GRASP e BRKGA com Aprendizagem por Reforço**. Dissertação (Dissertação de Mestrado) — Universidade Federal Rural do Semi-Árido (UFERSA), Mossoró, RN, Brasil, 2021.

REBOUÇAS, A. D. C.; SILVA, J. L.; OLIVEIRA, T. H. F.; ALOISE, D. J.; NASCIMENTO, H. A. D.; LONGO, H. J. A hybrid grasp with q-learning for the guillotine cutting stock problem. In: **Anais do XLV Simpósio Brasileiro de Pesquisa Operacional (SBPO)**. Caxambu, MG, Brasil: SOBRAPO, 2023.

RESENDE, M. G. C.; RIBEIRO, C. C. Grasp: Greedy randomized adaptive search procedures. **Handbook of Metaheuristics**, Springer, v. 57, p. 219–249, 2005.

SILVA, C. T. L. d. **Otimização de processos acoplados: programação da produção e corte de estoque**. Tese (Tese (Doutorado)) — Universidade de São Paulo, 2009.

TOKIC, M.; PALM, G. Value-difference based exploration: Adaptive control between epsilon-greedy and softmax. In: **Proceedings of the Annual Conference on Artificial Intelligence (KI 2011)**. [S.l.]: Springer, 2011. (Lecture Notes in Computer Science, v. 7006), p. 335–346.

WÄSCHER, G.; HAUSSNER, H.; SCHUMANN, H. An improved typology of cutting and packing problems. **European Journal of Operational Research**, v. 183, n. 3, p. 1109–1130, 2007.

WATKINS, C. J. C. H. **Learning from Delayed Rewards**. Tese (Doutorado) — King's College, Cambridge, 1989.

WIERINGA, R. **Design Science Methodology for Information Systems and Software Engineering**. Springer, 2009. ISBN 978-3-540-92966-4. Disponível em: <<https://doi.org/10.1007/978-3-540-92966-4>>.

XIONG, H.; GUO, C.; PENG, J.; DING, K.; CHEN, W.; QIU, X.; BAI, L.; XU, J. Gopt: Generalizable online 3d bin packing via transformer-based deep reinforcement learning. **arXiv preprint**, 2024.

ZHANG, J.; ZI, B.; GE, X. Attend2pack: Bin packing through deep reinforcement learning with attention. **arXiv preprint**, 2021.

ZHAO, H.; YU, Y.; XU, K. Learning efficient online 3d bin packing on packing configuration trees. In: **ICLR 2022**. [S.l.: s.n.], 2022.

APÊNDICE A – RESULTADOS - GRASP-LEARNING-DELAYED

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
1_20_1_451	6	6	6	0,00	110,53
1_20_2_452	3	3	3	0,00	106,75
1_20_3_453	4	4	4	0,00	110,91
1_20_4_454	4	4	4	0,00	111,77
1_20_5_455	4	4	4	0,00	109,99
1_20_6_456	4	4	4	0,00	112,38
1_20_7_457	5	5	5	0,00	112,05
1_20_8_458	3	3	3	0,00	110,81
1_20_9_459	5	5	5	0,00	112,79
1_20_10_460	3	3	3	0,00	111,91
1_40_1_461	8	8	8	0,00	453,54
1_40_2_462	8	8	8	0,00	465,04
1_40_3_463	9	9	9	0,00	512,04
1_40_4_464	6	6	6	0,00	553,24
1_40_5_465	6	6	6	0,00	464,83
1_40_6_466	6	6	6	0,00	455,63
1_40_7_467	7	7	7	0,00	470,29
1_40_8_468	7	7	7	0,00	473,14
1_40_9_469	8	8	8	0,00	466,96
1_40_10_470	8	8	8	0,00	537,01
1_60_1_471	12	12	12	0,00	1427,42
1_60_2_472	12	12	12	0,00	1842,38
1_60_3_473	12	12	12	0,00	1428,73
1_60_4_474	8	8	8	0,00	1351,81
1_60_5_475	8	8	8	0,00	1393,77
1_60_6_476	13	13	13	0,00	1417,05
1_60_7_477	10	10	10	0,00	1424,76
1_60_8_478	10	10	10	0,00	1419,63

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
1_60_9_479	9	9	9	0,00	1411,16
1_60_10_480	8	8	8	0,00	1428,81
1_80_1_481	13	13	13	0,00	3477,59
1_80_2_482	11	11	11	0,00	3387,82
1_80_3_483	11	11	11	0,00	3492,17
1_80_4_484	14	14	14	0,00	3415,02
1_80_5_485	14	14	14	0,00	4151,20
1_80_6_486	14	14	14	0,00	3360,11
1_80_7_487	15	15	15	0,00	3411,93
1_80_8_488	10	10	10	0,00	3331,79
1_80_9_489	13	13	13	0,00	3483,10
1_80_10_490	15	15	15	0,00	3494,62
1_100_1_491	15	15	15	0,00	6975,98
1_100_2_492	16	16	16	0,00	7501,59
1_100_3_493	16	17	16,50	0,71	7554,53
1_100_4_494	18	18	18	0,00	6960,83
1_100_5_495	19	19	19	0,00	6808,27
1_100_6_496	14	14	14	0,00	6676,91
1_100_7_497	14	14	14	0,00	6952,45
1_100_8_498	19	19	19	0,00	7355,08
1_100_9_499	17	17	17	0,00	7338,27
1_100_10_500	16	16	16	0,00	7178,50
2_20_1_51	1	1	1	0,00	2,27
2_20_2_52	1	1	1	0,00	2,30
2_20_3_53	1	1	1	0,00	2,29
2_20_4_54	1	1	1	0,00	2,38
2_20_5_55	1	1	1	0,00	2,34
2_20_6_56	1	1	1	0,00	2,32
2_20_7_57	1	1	1	0,00	2,28
2_20_8_58	1	1	1	0,00	2,29

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
2_20_9_59	1	1	1	0,00	2,30
2_20_10_60	1	1	1	0,00	2,27
2_40_1_61	2	2	2	0,00	10,89
2_40_2_62	2	2	2	0,00	11,03
2_40_3_63	2	2	2	0,00	11,40
2_40_4_64	2	2	2	0,00	11,31
2_40_5_65	2	2	2	0,00	11,29
2_40_6_66	2	2	2	0,00	11,02
2_40_7_67	2	2	2	0,00	11,09
2_40_8_68	2	2	2	0,00	11,44
2_40_9_69	2	2	2	0,00	11,26
2_40_10_70	2	2	2	0,00	11,29
2_60_1_71	3	3	3	0,00	35,51
2_60_2_72	3	3	3	0,00	36,13
2_60_3_73	3	3	3	0,00	35,70
2_60_4_74	3	3	3	0,00	35,21
2_60_5_75	2	2	2	0,00	34,28
2_60_6_76	2	2	2	0,00	35,47
2_60_7_77	2	2	2	0,00	33,77
2_60_8_78	3	3	3	0,00	35,31
2_60_9_79	2	2	2	0,00	41,88
2_60_10_80	3	3	3	0,00	36,45
2_80_1_81	3	3	3	0,00	85,89
2_80_2_82	3	3	3	0,00	86,18
2_80_3_83	3	3	3	0,00	84,58
2_80_4_84	3	3	3	0,00	84,10
2_80_5_85	3	3	3	0,00	86,11
2_80_6_86	3	3	3	0,00	86,15
2_80_7_87	3	4	3,97	0,18	90,08
2_80_8_88	4	4	4	0,00	86,73

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
2_80_9_89	4	4	4	0,00	87,89
2_80_10_90	3	3	3	0,00	83,80
2_100_1_91	4	4	4	0,00	173,32
2_100_2_92	4	4	4	0,00	177,72
2_100_3_93	4	4	4	0,00	172,05
2_100_4_94	4	4	4	0,00	176,04
2_100_5_95	4	4	4	0,00	179,71
2_100_6_96	4	4	4	0,00	180,56
2_100_7_97	4	4	4	0,00	172,02
2_100_8_98	4	4	4	0,00	175,26
2_100_9_99	4	4	4	0,00	172,70
2_100_10_100	4	4	4	0,00	179,06
3_20_1_101	5	5	5	0,00	2,66
3_20_2_102	3	3	3	0,00	2,56
3_20_3_103	5	6	5,03	0,18	3,12
3_20_4_104	4	4	4	0,00	2,47
3_20_5_105	4	4	4	0,00	2,50
3_20_6_106	6	6	6	0,00	2,51
3_20_7_107	4	4	4	0,00	2,37
3_20_8_108	4	4	4	0,00	2,41
3_20_9_109	5	5	5	0,00	2,46
3_20_10_110	6	6	6	0,00	2,49
3_40_1_111	7	7	7	0,00	12,84
3_40_2_112	8	9	8,07	0,25	16,95
3_40_3_113	11	11	11	0,00	14,57
3_40_4_114	9	10	9,87	0,35	14,04
3_40_5_115	11	11	11	0,00	14,50
3_40_6_116	10	10	10	0,00	13,25
3_40_7_117	8	9	8,77	0,43	13,93
3_40_8_118	13	13	13	0,00	14,07

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
3_40_9_119	8	8	8	0,00	13,74
3_40_10_120	8	8	8	0,00	13,45
3_60_1_121	16	16	16	0,00	56,75
3_60_2_122	13	14	13,07	0,25	58,15
3_60_3_123	14	15	14,10	0,31	56,69
3_60_4_124	15	15	15	0,00	56,32
3_60_5_125	13	13	13	0,00	45,85
3_60_6_126	13	13	13	0,00	52,46
3_60_7_127	11	11	11	0,00	49,07
3_60_8_128	15	16	15,03	0,18	59,91
3_60_9_129	13	14	13,07	0,25	59,82
3_60_10_130	17	17	17	0,00	53,79
3_80_1_131	17	18	17,80	0,41	147,06
3_80_2_132	18	18	18	0,00	141,46
3_80_3_133	18	19	18,47	0,51	159,37
3_80_4_134	19	19	19	0,00	153,53
3_80_5_135	18	18	18	0,00	135,84
3_80_6_136	20	21	20,20	0,41	161,56
3_80_7_137	21	21	21	0,00	156,08
3_80_8_138	21	22	21,93	0,25	145,97
3_80_9_139	22	22	22	0,00	150,14
3_80_10_140	18	19	18,90	0,31	137,89
3_100_1_141	20	20	20	0,00	311,04
3_100_2_142	23	24	23,57	0,50	326,96
3_100_3_143	20	21	20,10	0,31	333,71
3_100_4_144	21	22	21,93	0,25	303,50
3_100_5_145	23	24	23,83	0,38	327,21
3_100_6_146	26	27	26,97	0,18	311,22
3_100_7_147	21	21	21	0,00	291,30
3_100_8_148	24	24	24	0,00	312,07

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
3_100_9_149	22	23	22,50	0,51	335,77
3_100_10_150	28	28	28	0,00	317,43
4_20_1_151	1	1	1	0,00	2,38
4_20_2_152	1	1	1	0,00	2,37
4_20_3_153	1	1	1	0,00	2,38
4_20_4_154	1	1	1	0,00	2,35
4_20_5_155	1	1	1	0,00	2,40
4_20_6_156	1	1	1	0,00	2,38
4_20_7_157	1	1	1	0,00	2,35
4_20_8_158	1	1	1	0,00	2,37
4_20_9_159	1	1	1	0,00	2,39
4_20_10_160	1	1	1	0,00	2,41
4_40_1_161	1	2	1,20	0,41	14,81
4_40_2_162	2	2	2	0,00	12,38
4_40_3_163	2	2	2	0,00	12,72
4_40_4_164	2	2	2	0,00	12,53
4_40_5_165	2	2	2	0,00	12,59
4_40_6_166	2	2	2	0,00	12,33
4_40_7_167	2	2	2	0,00	12,36
4_40_8_168	2	2	2	0,00	12,60
4_40_9_169	2	2	2	0,00	12,51
4_40_10_170	2	2	2	0,00	12,48
4_60_1_171	3	3	3	0,00	42,86
4_60_2_172	2	2	2	0,00	55,65
4_60_3_173	3	3	3	0,00	42,60
4_60_4_174	3	3	3	0,00	42,19
4_60_5_175	2	2	2	0,00	41,83
4_60_6_176	2	2	2	0,00	46,36
4_60_7_177	2	2	2	0,00	40,86
4_60_8_178	3	3	3	0,00	42,64

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
4_60_9_179	2	3	2,10	0,31	51,93
4_60_10_180	3	3	3	0,00	43,27
4_80_1_181	3	3	3	0,00	109,19
4_80_2_182	3	3	3	0,00	108,83
4_80_3_183	3	3	3	0,00	107,72
4_80_4_184	3	3	3	0,00	107,70
4_80_5_185	3	3	3	0,00	109,56
4_80_6_186	3	3	3	0,00	142,33
4_80_7_187	4	4	4	0,00	112,34
4_80_8_188	4	4	4	0,00	110,34
4_80_9_189	4	4	4	0,00	111,13
4_80_10_190	3	3	3	0,00	107,35
4_100_1_191	4	4	4	0,00	235,12
4_100_2_192	4	4	4	0,00	238,74
4_100_3_193	4	4	4	0,00	230,53
4_100_4_194	4	4	4	0,00	236,53
4_100_5_195	4	4	4	0,00	238,28
4_100_6_196	4	4	4	0,00	282,09
4_100_7_197	4	4	4	0,00	229,25
4_100_8_198	4	4	4	0,00	234,18
4_100_9_199	4	4	4	0,00	232,95
4_100_10_200	4	4	4	0,00	281,24
5_20_1_201	6	6	6	0,00	2,62
5_20_2_202	4	4	4	0,00	2,43
5_20_3_203	7	7	7	0,00	2,47
5_20_4_204	5	5	5	0,00	2,45
5_20_5_205	5	5	5	0,00	2,51
5_20_6_206	8	8	8	0,00	2,49
5_20_7_207	5	5	5	0,00	2,41
5_20_8_208	5	5	5	0,00	2,44

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
5_20_9_209	6	6	6	0,00	3,39
5_20_10_210	6	6	6	0,00	2,77
5_40_1_211	8	8	8	0,00	14,56
5_40_2_212	10	10	10	0,00	15,37
5_40_3_213	14	14	14	0,00	14,46
5_40_4_214	12	12	12	0,00	15,41
5_40_5_215	13	13	13	0,00	14,95
5_40_6_216	12	12	12	0,00	13,68
5_40_7_217	10	10	10	0,00	14,63
5_40_8_218	16	16	16	0,00	14,48
5_40_9_219	10	10	10	0,00	14,09
5_40_10_220	10	10	10	0,00	14,83
5_60_1_221	20	20	20	0,00	61,56
5_60_2_222	17	17	17	0,00	52,95
5_60_3_223	18	19	18,37	0,49	67,01
5_60_4_224	19	20	19,03	0,18	63,59
5_60_5_225	16	16	16	0,00	55,19
5_60_6_226	16	16	16	0,00	58,70
5_60_7_227	14	14	14	0,00	54,85
5_60_8_228	19	19	19	0,00	58,19
5_60_9_229	16	16	16	0,00	59,14
5_60_10_230	22	22	22	0,00	57,75
5_80_1_231	22	23	22,83	0,38	162,20
5_80_2_232	23	23	23	0,00	160,63
5_80_3_233	24	24	24	0,00	152,63
5_80_4_234	24	24	24	0,00	160,03
5_80_5_235	22	23	22,93	0,25	157,69
5_80_6_236	25	26	25,77	0,43	170,36
5_80_7_237	26	27	26,27	0,45	175,35
5_80_8_238	27	27	27	0,00	166,96

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
5_80_9_239	27	28	27,10	0,31	193,64
5_80_10_240	23	23	23	0,00	162,75
5_100_1_241	25	26	25,10	0,31	400,03
5_100_2_242	29	29	29	0,00	387,12
5_100_3_243	25	26	25,03	0,18	409,49
5_100_4_244	27	28	27,60	0,50	390,56
5_100_5_245	29	30	29,70	0,47	390,49
5_100_6_246	33	34	33,10	0,31	450,33
5_100_7_247	26	26	26	0,00	331,68
5_100_8_248	30	31	30,07	0,25	421,96
5_100_9_249	28	28	28	0,00	343,58
5_100_10_250	34	34	34	0,00	388,56
6_20_1_251	1	1	1	0,00	2,32
6_20_2_252	1	1	1	0,00	2,35
6_20_3_253	1	1	1	0,00	2,34
6_20_4_254	1	1	1	0,00	2,31
6_20_5_255	1	1	1	0,00	2,35
6_20_6_256	1	1	1	0,00	2,35
6_20_7_257	1	1	1	0,00	2,33
6_20_8_258	1	1	1	0,00	2,38
6_20_9_259	1	1	1	0,00	2,33
6_20_10_260	1	1	1	0,00	2,32
6_40_1_261	1	1	1	0,00	12,35
6_40_2_262	2	2	2	0,00	12,98
6_40_3_263	2	2	2	0,00	12,78
6_40_4_264	2	2	2	0,00	12,73
6_40_5_265	2	2	2	0,00	12,73
6_40_6_266	2	2	2	0,00	12,52
6_40_7_267	2	2	2	0,00	12,67
6_40_8_268	2	2	2	0,00	12,71

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
6_40_9_269	2	2	2	0,00	12,71
6_40_10_270	2	2	2	0,00	12,72
6_60_1_271	3	3	3	0,00	44,67
6_60_2_272	2	2	2	0,00	44,23
6_60_3_273	2	2	2	0,00	45,60
6_60_4_274	2	2	2	0,00	46,28
6_60_5_275	2	2	2	0,00	43,96
6_60_6_276	2	2	2	0,00	44,30
6_60_7_277	2	2	2	0,00	43,41
6_60_8_278	2	2	2	0,00	53,55
6_60_9_279	2	2	2	0,00	43,99
6_60_10_280	3	3	3	0,00	45,56
6_80_1_281	3	3	3	0,00	118,32
6_80_2_282	3	3	3	0,00	118,61
6_80_3_283	3	3	3	0,00	115,80
6_80_4_284	3	3	3	0,00	116,70
6_80_5_285	3	3	3	0,00	117,30
6_80_6_286	3	3	3	0,00	117,88
6_80_7_287	3	3	3	0,00	118,80
6_80_8_288	3	3	3	0,00	117,82
6_80_9_289	3	3	3	0,00	120,76
6_80_10_290	3	3	3	0,00	115,78
6_100_1_291	3	3	3	0,00	258,50
6_100_2_292	4	4	4	0,00	260,77
6_100_3_293	3	3	3	0,00	255,17
6_100_4_294	3	3	3	0,00	278,53
6_100_5_295	4	4	4	0,00	261,01
6_100_6_296	4	4	4	0,00	264,12
6_100_7_297	3	3	3	0,00	253,23
6_100_8_298	4	4	4	0,00	257,70

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
6_100_9_299	3	4	3,90	0,32	258,07
6_100_10_300	4	4	4	0,00	259,34
7_20_1_301	5	5	5	0,00	2,38
7_20_2_302	5	5	5	0,00	2,38
7_20_3_303	5	5	5	0,00	2,39
7_20_4_304	6	6	6	0,00	2,62
7_20_5_305	5	5	5	0,00	2,48
7_20_6_306	5	5	5	0,00	2,44
7_20_7_307	4	4	4	0,00	2,39
7_20_8_308	5	5	5	0,00	2,70
7_20_9_309	5	6	5,80	0,42	2,70
7_20_10_310	4	4	4	0,00	2,47
7_40_1_311	10	10	10	0,00	13,60
7_40_2_312	12	12	12	0,00	15,40
7_40_3_313	10	10	10	0,00	13,75
7_40_4_314	13	13	13	0,00	15,06
7_40_5_315	9	10	9,80	0,42	14,13
7_40_6_316	11	11	11	0,00	18,63
7_40_7_317	11	11	11	0,00	14,80
7_40_8_318	11	11	11	0,00	15,03
7_40_9_319	8	8	8	0,00	13,44
7_40_10_320	12	12	12	0,00	15,17
7_60_1_321	17	17	17	0,00	56,41
7_60_2_322	14	14	14	0,00	48,40
7_60_3_323	16	16	16	0,00	62,05
7_60_4_324	15	15	15	0,00	55,39
7_60_5_325	14	14	14	0,00	52,02
7_60_6_326	14	14	14	0,00	53,84
7_60_7_327	14	15	14,80	0,42	58,55
7_60_8_328	17	17	17	0,00	56,34

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
7_60_9_329	14	14	14	0,00	63,77
7_60_10_330	19	19	19	0,00	59,75
7_80_1_331	21	21	21	0,00	150,71
7_80_2_332	24	24	24	0,00	181,97
7_80_3_333	20	21	20,30	0,48	172,46
7_80_4_334	21	22	21,20	0,42	170,21
7_80_5_335	22	22	22	0,00	199,38
7_80_6_336	22	23	22,40	0,52	164,21
7_80_7_337	24	24	24	0,00	146,06
7_80_8_338	21	22	21,90	0,32	140,98
7_80_9_339	22	23	22,50	0,53	183,15
7_80_10_340	23	24	23,10	0,32	183,34
7_100_1_341	27	27	27	0,00	351,05
7_100_2_342	27	27	27	0,00	381,54
7_100_3_343	24	24	24	0,00	317,81
7_100_4_344	25	26	25,90	0,32	344,33
7_100_5_345	24	24	24	0,00	335,34
7_100_6_346	29	29	29	0,00	417,74
7_100_7_347	26	26	26	0,00	419,19
7_100_8_348	28	28	28	0,00	356,86
7_100_9_349	25	26	25,80	0,42	344,06
7_100_10_350	32	32	32	0,00	397,22
8_20_1_351	5	5	5	0,00	2,96
8_20_2_352	6	6	6	0,00	2,43
8_20_3_353	5	5	5	0,00	2,57
8_20_4_354	6	6	6	0,00	2,47
8_20_5_355	4	5	4,90	0,32	2,50
8_20_6_356	6	6	6	0,00	2,41
8_20_7_357	4	4	4	0,00	2,40
8_20_8_358	5	5	5	0,00	2,40

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
8_20_9_359	6	6	6	0,00	2,41
8_20_10_360	4	4	4	0,00	2,38
8_40_1_361	11	11	11	0,00	15,29
8_40_2_362	12	13	12,90	0,32	15,03
8_40_3_363	10	11	10,10	0,32	18,61
8_40_4_364	11	12	11,50	0,53	16,10
8_40_5_365	8	9	8,50	0,53	15,93
8_40_6_366	11	12	11,10	0,32	17,34
8_40_7_367	10	11	10,60	0,52	15,71
8_40_8_368	11	11	11	0,00	13,64
8_40_9_369	9	9	9	0,00	14,26
8_40_10_370	12	13	12,80	0,42	14,85
8_60_1_371	17	17	17	0,00	56,38
8_60_2_372	17	17	17	0,00	50,26
8_60_3_373	16	17	16,10	0,32	62,62
8_60_4_374	15	15	15	0,00	50,18
8_60_5_375	14	14	14	0,00	54,54
8_60_6_376	15	15	15	0,00	56,87
8_60_7_377	14	14	14	0,00	56,41
8_60_8_378	17	18	17,60	0,52	60,76
8_60_9_379	16	16	16	0,00	50,94
8_60_10_380	17	17	17	0,00	58,49
8_80_1_381	21	22	21,50	0,53	172,78
8_80_2_382	22	23	22,80	0,42	149,92
8_80_3_383	20	20	20	0,00	141,81
8_80_4_384	19	20	19,20	0,42	164,20
8_80_5_385	26	26	26	0,00	179,09
8_80_6_386	22	22	22	0,00	143,95
8_80_7_387	21	21	21	0,00	149,64
8_80_8_388	22	22	22	0,00	165,95

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
8_80_9_389	21	22	21,10	0,32	188,83
8_80_10_390	24	24	24	0,00	170,01
8_100_1_391	25	26	25,40	0,52	346,85
8_100_2_392	26	27	26,80	0,42	334,92
8_100_3_393	23	24	23,70	0,48	346,30
8_100_4_394	30	31	30,20	0,42	448,94
8_100_5_395	27	28	27,90	0,32	367,13
8_100_6_396	26	26	26	0,00	359,70
8_100_7_397	26	26	26	0,00	369,16
8_100_8_398	28	29	28,10	0,32	424,05
8_100_9_399	27	27	27	0,00	367,20
8_100_10_400	32	32	32	0,00	375,35
8_100_1_391	4	4	4	0,00	235,12
8_100_2_392	4	4	4	0,00	238,74
8_100_3_393	4	4	4	0,00	230,53
8_100_4_394	4	4	4	0,00	236,53
8_100_5_395	4	4	4	0,00	238,28
8_100_6_396	4	4	4	0,00	282,09
8_100_7_397	4	4	4	0,00	229,25
8_100_8_398	4	4	4	0,00	234,18
8_100_9_399	4	4	4	0,00	232,95
8_100_10_400	4	4	4	0,00	281,24
9_20_1_401	18	18	18	0,00	2,87
9_20_2_402	13	13	13	0,00	2,61
9_20_3_403	13	13	13	0,00	2,70
9_20_4_404	15	15	15	0,00	2,73
9_20_5_405	15	15	15	0,00	2,74
9_20_6_406	14	14	14	0,00	2,74
9_20_7_407	8	8	8	0,00	2,57
9_20_8_408	13	13	13	0,00	2,64

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
9_20_9_409	14	14	14	0,00	2,72
9_20_10_410	13	13	13	0,00	2,63
9_40_1_411	24	24	24	0,00	15,58
9_40_2_412	31	31	31	0,00	17,35
9_40_3_413	27	27	27	0,00	16,45
9_40_4_414	30	30	30	0,00	17,26
9_40_5_415	27	27	27	0,00	16,03
9_40_6_416	29	29	29	0,00	16,68
9_40_7_417	23	23	23	0,00	15,49
9_40_8_418	25	25	25	0,00	16,30
9_40_9_419	21	21	21	0,00	15,37
9_40_10_420	32	32	32	0,00	17,57
9_60_1_421	45	45	45	0,00	66,81
9_60_2_422	43	43	43	0,00	64,22
9_60_3_423	45	45	45	0,00	66,89
9_60_4_424	44	44	44	0,00	65,19
9_60_5_425	41	41	41	0,00	62,95
9_60_6_426	37	37	37	0,00	61,16
9_60_7_427	39	39	39	0,00	62,43
9_60_8_428	47	47	47	0,00	67,21
9_60_9_429	44	44	44	0,00	64,23
9_60_10_430	44	44	44	0,00	65,25
9_80_1_431	57	57	57	0,00	181,80
9_80_2_432	56	56	56	0,00	182,04
9_80_3_433	57	57	57	0,00	182,87
9_80_4_434	52	52	52	0,00	171,99
9_80_5_435	60	60	60	0,00	187,41
9_80_6_436	62	62	62	0,00	188,49
9_80_7_437	58	58	58	0,00	180,78
9_80_8_438	58	58	58	0,00	179,38

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
9_80_9_439	48	48	48	0,00	165,55
9_80_10_440	59	59	59	0,00	183,92
9_100_1_441	70	70	70	0,00	411,52
9_100_2_442	64	64	64	0,00	387,52
9_100_3_443	68	68	68	0,00	401,73
9_100_4_444	77	77	77	0,00	434,95
9_100_5_445	64	64	64	0,00	389,07
9_100_6_446	70	70	70	0,00	408,75
9_100_7_447	65	65	65	0,00	391,09
9_100_8_448	72	72	72	0,00	417,01
9_100_9_449	65	65	65	0,00	395,54
9_100_10_450	71	71	71	0,00	412,81
10_20_1_451	6	6	6	0,00	2,51
10_20_2_452	3	3	3	0,00	2,46
10_20_3_453	4	4	4	0,00	2,79
10_20_4_454	4	4	4	0,00	2,50
10_20_5_455	4	4	4	0,00	2,43
10_20_6_456	4	4	4	0,00	2,49
10_20_7_457	5	5	5	0,00	2,44
10_20_8_458	3	3	3	0,00	2,40
10_20_9_459	5	5	5	0,00	2,63
10_20_10_460	3	3	3	0,00	2,44
10_40_1_461	8	8	8	0,00	13,34
10_40_2_462	8	8	8	0,00	13,48
10_40_3_463	9	9	9	0,00	14,22
10_40_4_464	7	7	7	0,00	13,47
10_40_5_465	6	6	6	0,00	13,38
10_40_6_466	6	6	6	0,00	13,01
10_40_7_467	7	7	7	0,00	13,24
10_40_8_468	7	7	7	0,00	13,66

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
10_40_9_469	8	8	8	0,00	13,49
10_40_10_470	9	9	9	0,00	13,49
10_60_1_471	12	12	12	0,00	47,88
10_60_2_472	12	13	12,40	0,55	61,49
10_60_3_473	12	12	12	0,00	49,07
10_60_4_474	8	8	8	0,00	44,99
10_60_5_475	8	8	8	0,00	50,20
10_60_6_476	13	13	13	0,00	48,24
10_60_7_477	10	10	10	0,00	49,04
10_60_8_478	10	11	10,40	0,55	50,44
10_60_9_479	9	9	9	0,00	47,30
10_60_10_480	8	9	8,40	0,55	58,06
10_80_1_481	13	14	13,60	0,55	143,09
10_80_2_482	11	11	11	0,00	132,87
10_80_3_483	11	12	11,20	0,45	143,29
10_80_4_484	14	14	14	0,00	139,27
10_80_5_485	14	15	14,40	0,55	174,92
10_80_6_486	14	14	14	0,00	132,15
10_80_7_487	15	15	15	0,00	131,09
10_80_8_488	10	10	10	0,00	169,28
10_80_9_489	13	13	13	0,00	139,21
10_80_10_490	15	15	15	0,00	145,78
10_100_1_491	15	15	15	0,00	296,60
10_100_2_492	16	16	16	0,00	348,02
10_100_3_493	17	17	17	0,00	273,00
10_100_4_494	18	18	18	0,00	336,39
10_100_5_495	19	19	19	0,00	279,38
10_100_6_496	14	14	14	0,00	272,78
10_100_7_497	14	15	14,40	0,55	281,16
10_100_8_498	19	19	19	0,00	326,55

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
10_100_9_499	17	17	17	0,00	322,39
10_100_10_500	16	16	16	0,00	294,79

Tabela 6 – Resultados experimentais para diferentes instâncias

APÊNDICE B – RESULTADOS - GRASP-LEARNING-DELAYED (BLA)

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
1_20_1_451	6	6	6	0,00	110,53
1_20_2_452	3	3	3	0,00	106,75
1_20_3_453	4	4	4	0,00	110,91
1_20_4_454	4	4	4	0,00	111,77
1_20_5_455	4	4	4	0,00	109,99
1_20_6_456	4	4	4	0,00	112,38
1_20_7_457	5	5	5	0,00	112,05
1_20_8_458	3	3	3	0,00	110,81
1_20_9_459	5	5	5	0,00	112,79
1_20_10_460	3	3	3	0,00	111,91
1_40_1_461	8	8	8	0,00	453,54
1_40_2_462	8	8	8	0,00	465,04
1_40_3_463	9	9	9	0,00	512,04
1_40_4_464	6	6	6	0,00	553,24
1_40_5_465	6	6	6	0,00	464,83
1_40_6_466	6	6	6	0,00	455,63
1_40_7_467	7	7	7	0,00	470,29
1_40_8_468	7	7	7	0,00	473,14
1_40_9_469	8	8	8	0,00	466,96
1_40_10_470	8	8	8	0,00	537,01
1_60_1_471	12	12	12	0,00	1427,42
1_60_2_472	12	12	12	0,00	1842,38
1_60_3_473	12	12	12	0,00	1428,73
1_60_4_474	8	8	8	0,00	1351,81
1_60_5_475	8	8	8	0,00	1393,77
1_60_6_476	13	13	13	0,00	1417,05
1_60_7_477	10	10	10	0,00	1424,76
1_60_8_478	10	10	10	0,00	1419,63

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
1_60_9_479	9	9	9	0,00	1411,16
1_60_10_480	8	8	8	0,00	1428,81
1_80_1_481	13	13	13	0,00	3477,59
1_80_2_482	11	11	11	0,00	3387,82
1_80_3_483	11	11	11	0,00	3492,17
1_80_4_484	14	14	14	0,00	3415,02
1_80_5_485	14	14	14	0,00	4151,20
1_80_6_486	14	14	14	0,00	3360,11
1_80_7_487	15	15	15	0,00	3411,93
1_80_8_488	10	10	10	0,00	3331,79
1_80_9_489	13	13	13	0,00	3483,10
1_80_10_490	15	15	15	0,00	3494,62
1_100_1_491	15	15	15	0,00	6975,98
1_100_2_492	16	16	16	0,00	7501,59
1_100_3_493	16	17	16,50	0,71	7554,53
1_100_4_494	18	18	18	0,00	6960,83
1_100_5_495	19	19	19	0,00	6808,27
1_100_6_496	14	14	14	0,00	6676,91
1_100_7_497	14	14	14	0,00	6952,45
1_100_8_498	19	19	19	0,00	7355,08
1_100_9_499	17	17	17	0,00	7338,27
1_100_10_500	16	16	16	0,00	7178,50
2_20_1_51	1	1	1	0,00	2,27
2_20_2_52	1	1	1	0,00	2,30
2_20_3_53	1	1	1	0,00	2,29
2_20_4_54	1	1	1	0,00	2,38
2_20_5_55	1	1	1	0,00	2,34
2_20_6_56	1	1	1	0,00	2,32
2_20_7_57	1	1	1	0,00	2,28
2_20_8_58	1	1	1	0,00	2,29

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
2_20_9_59	1	1	1	0,00	2,30
2_20_10_60	1	1	1	0,00	2,27
2_40_1_61	2	2	2	0,00	10,89
2_40_2_62	2	2	2	0,00	11,03
2_40_3_63	2	2	2	0,00	11,40
2_40_4_64	2	2	2	0,00	11,31
2_40_5_65	2	2	2	0,00	11,29
2_40_6_66	2	2	2	0,00	11,02
2_40_7_67	2	2	2	0,00	11,09
2_40_8_68	2	2	2	0,00	11,44
2_40_9_69	2	2	2	0,00	11,26
2_40_10_70	2	2	2	0,00	11,29
2_60_1_71	3	3	3	0,00	35,51
2_60_2_72	3	3	3	0,00	36,13
2_60_3_73	3	3	3	0,00	35,70
2_60_4_74	3	3	3	0,00	35,21
2_60_5_75	2	2	2	0,00	34,28
2_60_6_76	2	2	2	0,00	35,47
2_60_7_77	2	2	2	0,00	33,77
2_60_8_78	3	3	3	0,00	35,31
2_60_9_79	2	2	2	0,00	41,88
2_60_10_80	3	3	3	0,00	36,45
2_80_1_81	3	3	3	0,00	85,89
2_80_2_82	3	3	3	0,00	86,18
2_80_3_83	3	3	3	0,00	84,58
2_80_4_84	3	3	3	0,00	84,10
2_80_5_85	3	3	3	0,00	86,11
2_80_6_86	3	3	3	0,00	86,15
2_80_7_87	3	4	3,97	0,18	90,08
2_80_8_88	4	4	4	0,00	86,73

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
2_80_9_89	4	4	4	0,00	87,89
2_80_10_90	3	3	3	0,00	83,80
2_100_1_91	4	4	4	0,00	173,32
2_100_2_92	4	4	4	0,00	177,72
2_100_3_93	4	4	4	0,00	172,05
2_100_4_94	4	4	4	0,00	176,04
2_100_5_95	4	4	4	0,00	179,71
2_100_6_96	4	4	4	0,00	180,56
2_100_7_97	4	4	4	0,00	172,02
2_100_8_98	4	4	4	0,00	175,26
2_100_9_99	4	4	4	0,00	172,70
2_100_10_100	4	4	4	0,00	179,06
3_20_1_101	5	5	5	0,00	2,66
3_20_2_102	3	3	3	0,00	2,56
3_20_3_103	5	6	5,03	0,18	3,12
3_20_4_104	4	4	4	0,00	2,47
3_20_5_105	4	4	4	0,00	2,50
3_20_6_106	6	6	6	0,00	2,51
3_20_7_107	4	4	4	0,00	2,37
3_20_8_108	4	4	4	0,00	2,41
3_20_9_109	5	5	5	0,00	2,46
3_20_10_110	6	6	6	0,00	2,49
3_40_1_111	7	7	7	0,00	12,84
3_40_2_112	8	9	8,07	0,25	16,95
3_40_3_113	11	11	11	0,00	14,57
3_40_4_114	9	10	9,87	0,35	14,04
3_40_5_115	11	11	11	0,00	14,50
3_40_6_116	10	10	10	0,00	13,25
3_40_7_117	8	9	8,77	0,43	13,93
3_40_8_118	13	13	13	0,00	14,07

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
3_40_9_119	8	8	8	0,00	13,74
3_40_10_120	8	8	8	0,00	13,45
3_60_1_121	16	16	16	0,00	56,75
3_60_2_122	13	14	13,07	0,25	58,15
3_60_3_123	14	15	14,10	0,31	56,69
3_60_4_124	15	15	15	0,00	56,32
3_60_5_125	13	13	13	0,00	45,85
3_60_6_126	13	13	13	0,00	52,46
3_60_7_127	11	11	11	0,00	49,07
3_60_8_128	15	16	15,03	0,18	59,91
3_60_9_129	13	14	13,07	0,25	59,82
3_60_10_130	17	17	17	0,00	53,79
3_80_1_131	17	18	17,80	0,41	147,06
3_80_2_132	18	18	18	0,00	141,46
3_80_3_133	18	19	18,47	0,51	159,37
3_80_4_134	19	19	19	0,00	153,53
3_80_5_135	18	18	18	0,00	135,84
3_80_6_136	20	21	20,20	0,41	161,56
3_80_7_137	21	21	21	0,00	156,08
3_80_8_138	21	22	21,93	0,25	145,97
3_80_9_139	22	22	22	0,00	150,14
3_80_10_140	18	19	18,90	0,31	137,89
3_100_1_141	20	20	20	0,00	311,04
3_100_2_142	23	24	23,57	0,50	326,96
3_100_3_143	20	21	20,10	0,31	333,71
3_100_4_144	21	22	21,93	0,25	303,50
3_100_5_145	23	24	23,83	0,38	327,21
3_100_6_146	26	27	26,97	0,18	311,22
3_100_7_147	21	21	21	0,00	291,30
3_100_8_148	24	24	24	0,00	312,07

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
3_100_9_149	22	23	22,50	0,51	335,77
3_100_10_150	28	28	28	0,00	317,43
4_20_1_151	1	1	1	0,00	2,38
4_20_2_152	1	1	1	0,00	2,37
4_20_3_153	1	1	1	0,00	2,38
4_20_4_154	1	1	1	0,00	2,35
4_20_5_155	1	1	1	0,00	2,40
4_20_6_156	1	1	1	0,00	2,38
4_20_7_157	1	1	1	0,00	2,35
4_20_8_158	1	1	1	0,00	2,37
4_20_9_159	1	1	1	0,00	2,39
4_20_10_160	1	1	1	0,00	2,41
4_40_1_161	1	2	1,20	0,41	14,81
4_40_2_162	2	2	2	0,00	12,38
4_40_3_163	2	2	2	0,00	12,72
4_40_4_164	2	2	2	0,00	12,53
4_40_5_165	2	2	2	0,00	12,59
4_40_6_166	2	2	2	0,00	12,33
4_40_7_167	2	2	2	0,00	12,36
4_40_8_168	2	2	2	0,00	12,60
4_40_9_169	2	2	2	0,00	12,51
4_40_10_170	2	2	2	0,00	12,48
4_60_1_171	3	3	3	0,00	42,86
4_60_2_172	2	2	2	0,00	55,65
4_60_3_173	3	3	3	0,00	42,60
4_60_4_174	3	3	3	0,00	42,19
4_60_5_175	2	2	2	0,00	41,83
4_60_6_176	2	2	2	0,00	46,36
4_60_7_177	2	2	2	0,00	40,86
4_60_8_178	3	3	3	0,00	42,64

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
4_60_9_179	2	3	2,10	0,31	51,93
4_60_10_180	3	3	3	0,00	43,27
4_80_1_181	3	3	3	0,00	109,19
4_80_2_182	3	3	3	0,00	108,83
4_80_3_183	3	3	3	0,00	107,72
4_80_4_184	3	3	3	0,00	107,70
4_80_5_185	3	3	3	0,00	109,56
4_80_6_186	3	3	3	0,00	142,33
4_80_7_187	4	4	4	0,00	112,34
4_80_8_188	4	4	4	0,00	110,34
4_80_9_189	4	4	4	0,00	111,13
4_80_10_190	3	3	3	0,00	107,35
4_100_1_191	4	4	4	0,00	235,12
4_100_2_192	4	4	4	0,00	238,74
4_100_3_193	4	4	4	0,00	230,53
4_100_4_194	4	4	4	0,00	236,53
4_100_5_195	4	4	4	0,00	238,28
4_100_6_196	4	4	4	0,00	282,09
4_100_7_197	4	4	4	0,00	229,25
4_100_8_198	4	4	4	0,00	234,18
4_100_9_199	4	4	4	0,00	232,95
4_100_10_200	4	4	4	0,00	281,24
5_20_1_201	6	6	6	0,00	2,62
5_20_2_202	4	4	4	0,00	2,43
5_20_3_203	7	7	7	0,00	2,47
5_20_4_204	5	5	5	0,00	2,45
5_20_5_205	5	5	5	0,00	2,51
5_20_6_206	8	8	8	0,00	2,49
5_20_7_207	5	5	5	0,00	2,41
5_20_8_208	5	5	5	0,00	2,44

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
5_20_9_209	6	6	6	0,00	3,39
5_20_10_210	6	6	6	0,00	2,77
5_40_1_211	8	8	8	0,00	14,56
5_40_2_212	10	10	10	0,00	15,37
5_40_3_213	14	14	14	0,00	14,46
5_40_4_214	12	12	12	0,00	15,41
5_40_5_215	13	13	13	0,00	14,95
5_40_6_216	12	12	12	0,00	13,68
5_40_7_217	10	10	10	0,00	14,63
5_40_8_218	16	16	16	0,00	14,48
5_40_9_219	10	10	10	0,00	14,09
5_40_10_220	10	10	10	0,00	14,83
5_60_1_221	20	20	20	0,00	61,56
5_60_2_222	17	17	17	0,00	52,95
5_60_3_223	18	19	18,37	0,49	67,01
5_60_4_224	19	20	19,03	0,18	63,59
5_60_5_225	16	16	16	0,00	55,19
5_60_6_226	16	16	16	0,00	58,70
5_60_7_227	14	14	14	0,00	54,85
5_60_8_228	19	19	19	0,00	58,19
5_60_9_229	16	16	16	0,00	59,14
5_60_10_230	22	22	22	0,00	57,75
5_80_1_231	22	23	22,83	0,38	162,20
5_80_2_232	23	23	23	0,00	160,63
5_80_3_233	24	24	24	0,00	152,63
5_80_4_234	24	24	24	0,00	160,03
5_80_5_235	22	23	22,93	0,25	157,69
5_80_6_236	25	26	25,77	0,43	170,36
5_80_7_237	26	27	26,27	0,45	175,35
5_80_8_238	27	27	27	0,00	166,96

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
5_80_9_239	27	28	27,10	0,31	193,64
5_80_10_240	23	23	23	0,00	162,75
5_100_1_241	25	26	25,10	0,31	400,03
5_100_2_242	29	29	29	0,00	387,12
5_100_3_243	25	26	25,03	0,18	409,49
5_100_4_244	27	28	27,60	0,50	390,56
5_100_5_245	29	30	29,70	0,47	390,49
5_100_6_246	33	34	33,10	0,31	450,33
5_100_7_247	26	26	26	0,00	331,68
5_100_8_248	30	31	30,07	0,25	421,96
5_100_9_249	28	28	28	0,00	343,58
5_100_10_250	34	34	34	0,00	388,56
6_20_1_251	1	1	1	0,00	2,32
6_20_2_252	1	1	1	0,00	2,35
6_20_3_253	1	1	1	0,00	2,34
6_20_4_254	1	1	1	0,00	2,31
6_20_5_255	1	1	1	0,00	2,35
6_20_6_256	1	1	1	0,00	2,35
6_20_7_257	1	1	1	0,00	2,33
6_20_8_258	1	1	1	0,00	2,38
6_20_9_259	1	1	1	0,00	2,33
6_20_10_260	1	1	1	0,00	2,32
6_40_1_261	1	1	1	0,00	12,35
6_40_2_262	2	2	2	0,00	12,98
6_40_3_263	2	2	2	0,00	12,78
6_40_4_264	2	2	2	0,00	12,73
6_40_5_265	2	2	2	0,00	12,73
6_40_6_266	2	2	2	0,00	12,52
6_40_7_267	2	2	2	0,00	12,67
6_40_8_268	2	2	2	0,00	12,71

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
6_40_9_269	2	2	2	0,00	12,71
6_40_10_270	2	2	2	0,00	12,72
6_60_1_271	3	3	3	0,00	44,67
6_60_2_272	2	2	2	0,00	44,23
6_60_3_273	2	2	2	0,00	45,60
6_60_4_274	2	2	2	0,00	46,28
6_60_5_275	2	2	2	0,00	43,96
6_60_6_276	2	2	2	0,00	44,30
6_60_7_277	2	2	2	0,00	43,41
6_60_8_278	2	2	2	0,00	53,55
6_60_9_279	2	2	2	0,00	43,99
6_60_10_280	3	3	3	0,00	45,56
6_80_1_281	3	3	3	0,00	118,32
6_80_2_282	3	3	3	0,00	118,61
6_80_3_283	3	3	3	0,00	115,80
6_80_4_284	3	3	3	0,00	116,70
6_80_5_285	3	3	3	0,00	117,30
6_80_6_286	3	3	3	0,00	117,88
6_80_7_287	3	3	3	0,00	118,80
6_80_8_288	3	3	3	0,00	117,82
6_80_9_289	3	3	3	0,00	120,76
6_80_10_290	3	3	3	0,00	115,78
6_100_1_291	3	3	3	0,00	258,50
6_100_2_292	4	4	4	0,00	260,77
6_100_3_293	3	3	3	0,00	255,17
6_100_4_294	3	3	3	0,00	278,53
6_100_5_295	4	4	4	0,00	261,01
6_100_6_296	4	4	4	0,00	264,12
6_100_7_297	3	3	3	0,00	253,23
6_100_8_298	4	4	4	0,00	257,70

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
6_100_9_299	3	4	3,90	0,32	258,07
6_100_10_300	4	4	4	0,00	259,34
7_20_1_301	5	5	5	0,00	2,38
7_20_2_302	5	5	5	0,00	2,38
7_20_3_303	5	5	5	0,00	2,39
7_20_4_304	6	6	6	0,00	2,62
7_20_5_305	5	5	5	0,00	2,48
7_20_6_306	5	5	5	0,00	2,44
7_20_7_307	4	4	4	0,00	2,39
7_20_8_308	5	5	5	0,00	2,70
7_20_9_309	5	6	5,80	0,42	2,70
7_20_10_310	4	4	4	0,00	2,47
7_40_1_311	10	10	10	0,00	13,60
7_40_2_312	12	12	12	0,00	15,40
7_40_3_313	10	10	10	0,00	13,75
7_40_4_314	13	13	13	0,00	15,06
7_40_5_315	9	10	9,80	0,42	14,13
7_40_6_316	11	11	11	0,00	18,63
7_40_7_317	11	11	11	0,00	14,80
7_40_8_318	11	11	11	0,00	15,03
7_40_9_319	8	8	8	0,00	13,44
7_40_10_320	12	12	12	0,00	15,17
7_60_1_321	17	17	17	0,00	56,41
7_60_2_322	14	14	14	0,00	48,40
7_60_3_323	16	16	16	0,00	62,05
7_60_4_324	15	15	15	0,00	55,39
7_60_5_325	14	14	14	0,00	52,02
7_60_6_326	14	14	14	0,00	53,84
7_60_7_327	14	15	14,80	0,42	58,55
7_60_8_328	17	17	17	0,00	56,34

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
7_60_9_329	14	14	14	0,00	63,77
7_60_10_330	19	19	19	0,00	59,75
7_80_1_331	21	21	21	0,00	150,71
7_80_2_332	24	24	24	0,00	181,97
7_80_3_333	20	21	20,30	0,48	172,46
7_80_4_334	21	22	21,20	0,42	170,21
7_80_5_335	22	22	22	0,00	199,38
7_80_6_336	22	23	22,40	0,52	164,21
7_80_7_337	24	24	24	0,00	146,06
7_80_8_338	21	22	21,90	0,32	140,98
7_80_9_339	22	23	22,50	0,53	183,15
7_80_10_340	23	24	23,10	0,32	183,34
7_100_1_341	27	27	27	0,00	351,05
7_100_2_342	27	27	27	0,00	381,54
7_100_3_343	24	24	24	0,00	317,81
7_100_4_344	25	26	25,90	0,32	344,33
7_100_5_345	24	24	24	0,00	335,34
7_100_6_346	29	29	29	0,00	417,74
7_100_7_347	26	26	26	0,00	419,19
7_100_8_348	28	28	28	0,00	356,86
7_100_9_349	25	26	25,80	0,42	344,06
7_100_10_350	32	32	32	0,00	397,22
8_20_1_351	5	5	5	0,00	2,96
8_20_2_352	6	6	6	0,00	2,43
8_20_3_353	5	5	5	0,00	2,57
8_20_4_354	6	6	6	0,00	2,47
8_20_5_355	4	5	4,90	0,32	2,50
8_20_6_356	6	6	6	0,00	2,41
8_20_7_357	4	4	4	0,00	2,40
8_20_8_358	5	5	5	0,00	2,40

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
8_20_9_359	6	6	6	0,00	2,41
8_20_10_360	4	4	4	0,00	2,38
8_40_1_361	11	11	11	0,00	15,29
8_40_2_362	12	13	12,90	0,32	15,03
8_40_3_363	10	11	10,10	0,32	18,61
8_40_4_364	11	12	11,50	0,53	16,10
8_40_5_365	8	9	8,50	0,53	15,93
8_40_6_366	11	12	11,10	0,32	17,34
8_40_7_367	10	11	10,60	0,52	15,71
8_40_8_368	11	11	11	0,00	13,64
8_40_9_369	9	9	9	0,00	14,26
8_40_10_370	12	13	12,80	0,42	14,85
8_60_1_371	17	17	17	0,00	56,38
8_60_2_372	17	17	17	0,00	50,26
8_60_3_373	16	17	16,10	0,32	62,62
8_60_4_374	15	15	15	0,00	50,18
8_60_5_375	14	14	14	0,00	54,54
8_60_6_376	15	15	15	0,00	56,87
8_60_7_377	14	14	14	0,00	56,41
8_60_8_378	17	18	17,60	0,52	60,76
8_60_9_379	16	16	16	0,00	50,94
8_60_10_380	17	17	17	0,00	58,49
8_80_1_381	21	22	21,50	0,53	172,78
8_80_2_382	22	23	22,80	0,42	149,92
8_80_3_383	20	20	20	0,00	141,81
8_80_4_384	19	20	19,20	0,42	164,20
8_80_5_385	26	26	26	0,00	179,09
8_80_6_386	22	22	22	0,00	143,95
8_80_7_387	21	21	21	0,00	149,64
8_80_8_388	22	22	22	0,00	165,95

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
8_80_9_389	21	22	21,10	0,32	188,83
8_80_10_390	24	24	24	0,00	170,01
8_100_1_391	25	26	25,40	0,52	346,85
8_100_2_392	26	27	26,80	0,42	334,92
8_100_3_393	23	24	23,70	0,48	346,30
8_100_4_394	30	31	30,20	0,42	448,94
8_100_5_395	27	28	27,90	0,32	367,13
8_100_6_396	26	26	26	0,00	359,70
8_100_7_397	26	26	26	0,00	369,16
8_100_8_398	28	29	28,10	0,32	424,05
8_100_9_399	27	27	27	0,00	367,20
8_100_10_400	32	32	32	0,00	375,35
8_100_1_391	4	4	4	0,00	235,12
8_100_2_392	4	4	4	0,00	238,74
8_100_3_393	4	4	4	0,00	230,53
8_100_4_394	4	4	4	0,00	236,53
8_100_5_395	4	4	4	0,00	238,28
8_100_6_396	4	4	4	0,00	282,09
8_100_7_397	4	4	4	0,00	229,25
8_100_8_398	4	4	4	0,00	234,18
8_100_9_399	4	4	4	0,00	232,95
8_100_10_400	4	4	4	0,00	281,24
9_20_1_401	18	18	18	0,00	2,87
9_20_2_402	13	13	13	0,00	2,61
9_20_3_403	13	13	13	0,00	2,70
9_20_4_404	15	15	15	0,00	2,73
9_20_5_405	15	15	15	0,00	2,74
9_20_6_406	14	14	14	0,00	2,74
9_20_7_407	8	8	8	0,00	2,57
9_20_8_408	13	13	13	0,00	2,64

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
9_20_9_409	14	14	14	0,00	2,72
9_20_10_410	13	13	13	0,00	2,63
9_40_1_411	24	24	24	0,00	15,58
9_40_2_412	31	31	31	0,00	17,35
9_40_3_413	27	27	27	0,00	16,45
9_40_4_414	30	30	30	0,00	17,26
9_40_5_415	27	27	27	0,00	16,03
9_40_6_416	29	29	29	0,00	16,68
9_40_7_417	23	23	23	0,00	15,49
9_40_8_418	25	25	25	0,00	16,30
9_40_9_419	21	21	21	0,00	15,37
9_40_10_420	32	32	32	0,00	17,57
9_60_1_421	45	45	45	0,00	66,81
9_60_2_422	43	43	43	0,00	64,22
9_60_3_423	45	45	45	0,00	66,89
9_60_4_424	44	44	44	0,00	65,19
9_60_5_425	41	41	41	0,00	62,95
9_60_6_426	37	37	37	0,00	61,16
9_60_7_427	39	39	39	0,00	62,43
9_60_8_428	47	47	47	0,00	67,21
9_60_9_429	44	44	44	0,00	64,23
9_60_10_430	44	44	44	0,00	65,25
9_80_1_431	57	57	57	0,00	181,80
9_80_2_432	56	56	56	0,00	182,04
9_80_3_433	57	57	57	0,00	182,87
9_80_4_434	52	52	52	0,00	171,99
9_80_5_435	60	60	60	0,00	187,41
9_80_6_436	62	62	62	0,00	188,49
9_80_7_437	58	58	58	0,00	180,78
9_80_8_438	58	58	58	0,00	179,38

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
9_80_9_439	48	48	48	0,00	165,55
9_80_10_440	59	59	59	0,00	183,92
9_100_1_441	70	70	70	0,00	411,52
9_100_2_442	64	64	64	0,00	387,52
9_100_3_443	68	68	68	0,00	401,73
9_100_4_444	77	77	77	0,00	434,95
9_100_5_445	64	64	64	0,00	389,07
9_100_6_446	70	70	70	0,00	408,75
9_100_7_447	65	65	65	0,00	391,09
9_100_8_448	72	72	72	0,00	417,01
9_100_9_449	65	65	65	0,00	395,54
9_100_10_450	71	71	71	0,00	412,81
10_20_1_451	6	6	6	0,00	2,51
10_20_2_452	3	3	3	0,00	2,46
10_20_3_453	4	4	4	0,00	2,79
10_20_4_454	4	4	4	0,00	2,50
10_20_5_455	4	4	4	0,00	2,43
10_20_6_456	4	4	4	0,00	2,49
10_20_7_457	5	5	5	0,00	2,44
10_20_8_458	3	3	3	0,00	2,40
10_20_9_459	5	5	5	0,00	2,63
10_20_10_460	3	3	3	0,00	2,44
10_40_1_461	8	8	8	0,00	13,34
10_40_2_462	8	8	8	0,00	13,48
10_40_3_463	9	9	9	0,00	14,22
10_40_4_464	7	7	7	0,00	13,47
10_40_5_465	6	6	6	0,00	13,38
10_40_6_466	6	6	6	0,00	13,01
10_40_7_467	7	7	7	0,00	13,24
10_40_8_468	7	7	7	0,00	13,66

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
10_40_9_469	8	8	8	0,00	13,49
10_40_10_470	9	9	9	0,00	13,49
10_60_1_471	12	12	12	0,00	47,88
10_60_2_472	12	13	12,40	0,55	61,49
10_60_3_473	12	12	12	0,00	49,07
10_60_4_474	8	8	8	0,00	44,99
10_60_5_475	8	8	8	0,00	50,20
10_60_6_476	13	13	13	0,00	48,24
10_60_7_477	10	10	10	0,00	49,04
10_60_8_478	10	11	10,40	0,55	50,44
10_60_9_479	9	9	9	0,00	47,30
10_60_10_480	8	9	8,40	0,55	58,06
10_80_1_481	13	14	13,60	0,55	143,09
10_80_2_482	11	11	11	0,00	132,87
10_80_3_483	11	12	11,20	0,45	143,29
10_80_4_484	14	14	14	0,00	139,27
10_80_5_485	14	15	14,40	0,55	174,92
10_80_6_486	14	14	14	0,00	132,15
10_80_7_487	15	15	15	0,00	131,09
10_80_8_488	10	10	10	0,00	169,28
10_80_9_489	13	13	13	0,00	139,21
10_80_10_490	15	15	15	0,00	145,78
10_100_1_491	15	15	15	0,00	296,60
10_100_2_492	16	16	16	0,00	348,02
10_100_3_493	17	17	17	0,00	273,00
10_100_4_494	18	18	18	0,00	336,39
10_100_5_495	19	19	19	0,00	279,38
10_100_6_496	14	14	14	0,00	272,78
10_100_7_497	14	15	14,40	0,55	281,16
10_100_8_498	19	19	19	0,00	326,55

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
10_100_9_499	17	17	17	0,00	322,39
10_100_10_500	16	16	16	0,00	294,79

Tabela 7 – Resultados experimentais para diferentes instâncias

APÊNDICE C – RESULTADOS - GRASP-LEARNING-DELAYED (ILSA)

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
10_20_1_451	6	6	6	0,00	110,53
10_20_2_452	3	3	3	0,00	106,75
10_20_3_453	4	4	4	0,00	110,91
10_20_4_454	4	4	4	0,00	111,77
10_20_5_455	4	4	4	0,00	109,99
10_20_6_456	4	4	4	0,00	112,38
10_20_7_457	5	5	5	0,00	112,05
10_20_8_458	3	3	3	0,00	110,81
10_20_9_459	5	5	5	0,00	112,79
10_20_10_460	3	3	3	0,00	111,91
10_40_1_461	8	8	8	0,00	453,54
10_40_2_462	8	8	8	0,00	465,04
10_40_3_463	9	9	9	0,00	512,04
10_40_4_464	6	6	6	0,00	553,24
10_40_5_465	6	6	6	0,00	464,83
10_40_6_466	6	6	6	0,00	455,63
10_40_7_467	7	7	7	0,00	470,29
10_40_8_468	7	7	7	0,00	473,14
10_40_9_469	8	8	8	0,00	466,96
10_40_10_470	8	8	8	0,00	537,01
10_60_1_471	12	12	12	0,00	1427,42
10_60_2_472	12	12	12	0,00	1842,38
10_60_3_473	12	12	12	0,00	1428,73
10_60_4_474	8	8	8	0,00	1351,81
10_60_5_475	8	8	8	0,00	1393,77
10_60_6_476	12	13	12,9	0,30	1917,05
10_60_7_477	10	10	10	0,00	1424,76
10_60_8_478	10	10	10	0,00	1419,63

Instancia	Melhor	Pior	Média	Desvio Padrão	Tempo
10_60_9_479	9	9	9	0,00	1411,16
10_60_10_480	8	8	8	0,00	1428,81
10_80_1_481	13	13	13	0,00	3477,59
10_80_2_482	11	11	11	0,00	3387,82
10_80_3_483	11	11	11	0,00	3492,17
10_80_4_484	14	14	14	0,00	3415,02
10_80_5_485	14	14	14	0,00	4151,20
10_80_6_486	14	14	14	0,00	3360,11
10_80_7_487	15	15	15	0,00	3411,93
10_80_8_488	10	10	10	0,00	3331,79
10_80_9_489	13	13	13	0,00	3483,10
10_80_10_490	15	15	15	0,00	3494,62
10_100_1_491	15	15	15	0,00	6975,98
10_100_2_492	16	16	16	0,00	7501,59
10_100_3_493	16	17	16,50	0,71	7554,53
10_100_4_494	18	18	18	0,00	6960,83
10_100_5_495	19	19	19	0,00	6808,27
10_100_6_496	14	14	14	0,00	6676,91
10_100_7_497	14	14	14	0,00	6952,45
10_100_8_498	19	19	19	0,00	7355,08
10_100_9_499	17	17	17	0,00	7338,27
10_100_10_500	16	16	16	0,00	7178,50

Tabela 8 – Resultados experimentais para diferentes instâncias