
Identificação de Viabilidade de
Leveduras Com Corante Vital Utilizando
Histogramas de Palavras Visuais em
Imagens Coloridas

Junior Silva Souza

SERVIÇO DE PÓS-GRADUAÇÃO DA FACOM-UFMS

Data de Depósito:

Assinatura:_____

Identificação de Viabilidade de Leveduras Com Corante Vital Utilizando
Histogramas de Palavras Visuais em Imagens Coloridas

Junior Silva Souza

Orientador: *Prof^o. Dr^o. Hemerson Pistori*

Co-Orientador: *Prof^o. Dr^o. Wesley Nunes Gonçalves*

Dissertação apresentada ao curso de Pós Graduação em Ciências da Computação, da Universidade Federal de Mato Grosso do Sul, como requisito parcial à obtenção do título de Mestre em Ciências da Computação.

**FACOM - Universidade Federal de Mato Grosso do Sul.
Fevereiro/2015**

A Deus em primeiro lugar,

*Aos meus pais,
José e Lúcia,*

*Ao meu orientador,
Prof^o. Dr^o. Hemerson Pistori,*

*Ao meu co-orientador,
Prof^o. Dr^o. Wesley Nunes Gonçalves.*

Agradecimentos

Primeiramente agradeço a Deus pela realização de mais um sonho. Pela força inexplicável nos momentos mais difíceis nesta jornada.

Agradeço aos meus pais, José Ademir Santana de Souza e Lúcia Tereza da Silva, pela motivação dada em toda esta jornada. Também pelos esforços que fizeram para realizar meus sonhos.

Agradeço ao meu orientador, Prof^o. Dr^o. Hemerson Pistori, pela paciência, dedicação e compromisso. Não existiu dia, feriado ou final de semana, que ele não pudesse responder a alguma pergunta ou orientação. A segurança em trabalhar com um excelente profissional, nos leva a procurar trabalhar com satisfação. Durante toda esta jornada, sempre trabalhei seguro.

Agradeço ao meu co-orientador, Prof^o. Dr^o. Wesley Nunes Gonçalves, também pela paciência e dedicação. Sempre dando conselho, orientando e dando muitas dicas que auxiliaram muito na realização deste trabalho.

Agradeço à Prof^a. Dr^a. Marney Pascoli Cereda, por sua dedicação, por fazer acreditar ainda mais no meu trabalho. Seus ensinamentos e conselhos auxiliaram muito e foram muito proveitosos.

Agradeço à Prof^a. Dr^a. Valguima Victoria Viana Aguiar Odakura, também por sua dedicação, conhecimento. Sempre se disponibilizou para auxiliar e teve muita calma.

Agradeço ao grupo INOVISÃO, meus colegas que me ajudaram muito com o conhecimento e conselhos. Vou citar alguns nomes: Ariadne Gonçalves, Gercina Gonçalves, Uéliton Freitas, Diogo Soares. Existem muito outros nomes que desde já agradeço a todos.

Agradeço aos amigos, Existem muitos amigos que conheci durante esta jornada, vou citar alguns: Guilherme Figueiredo Terenciani, Marcelo Figueiredo Terenciani, Bárbara Purkott Cezar e Alex Zanella Zaccaron. Que me fizeram sorrir nas horas difíceis e sempre estiveram disponíveis para me ajudar.

Agradeço, por fim, a instituição UFMS e o órgão CNPq. Por acreditaram no meu trabalho e me auxiliarem durante esta jornada.

Resumo

Neste trabalho é apresentado um sistema para automatizar o processo de identificação de leveduras viáveis que são importantes na produção do etanol. A produção do etanol depende das leveduras viáveis, responsáveis pelo processo de fermentação do caldo da cana. Esta fermentação transforma o caldo da cana em etanol e gás carbônico. Portanto, manter um controle da população das leveduras viáveis é uma tarefa crucial no processo de produção do etanol, e para isto, são feitas análises para identificar e contar as leveduras. Apesar da importância dessa etapa, a identificação e contagem é feita por visão humana em um microscópio, sendo uma tarefa repetitiva e suscetível a erros.

Neste trabalho, técnicas de visão computacional e aprendizagem supervisionada foram avaliadas para automatizar a identificação destas leveduras. A principal técnica de visão computacional estudada foi o histograma de palavras visuais (*Bag-of-Visual-Words*), que é aplicado em imagens em tons de cinza. Além desta técnica utilizamos variantes que adicionam a informação de cor, como: CCV (*Color Coherence Vectors*), CM (*Color Moments*), BoC (*Bag-of-Color*) e o OpC (*Opponent Color*). Os atributos extraídos através destes algoritmos e suas variantes, foram utilizados para o teste e treinamento dos classificadores obtidos de técnicas de aprendizagem supervisionada. Entre as técnicas, utilizamos o Naive Bayes, KNN, J48 e SVM que estão disponíveis no ambiente Weka. A avaliação de desempenho, por meio da porcentagem de classificação correta, foi realizada através dos testes de hipótese ANOVA e Friedman.

Foi utilizado o banco de imagens do projeto BioViC¹ que foi proposto para a identificação de leveduras. Este banco possui um conjunto de imagens de leveduras capturadas em laboratório através de microscópio. A câmera de Neubauer foi utilizada para auxiliar a contagem das leveduras, de modo que, 6 repetições com 4 quadrantes nas concentrações de Brix 3, 6 e 12 foram

¹<http://biovic.weebly.com/bancos-de-imagens.html>

definidas para análise. As imagens obtidas com Brix 03 foram recortadas separando as imagens das leveduras viáveis, inviáveis e também o fundo que corresponde toda região que não possui leveduras viáveis ou inviáveis. O total de imagens obtidas e separadas em três classes (viável, inviável e fundo da imagem) foram 2614, utilizadas para o treinamento e identificação.

Os resultados foram analisados através do software R, que na análise de variância ANOVA apresentou um valor-p igual a $2e^{-16}$ indicando uma diferença significativa entre as técnicas utilizadas, descartando a hipótese nula. A técnica OpC com o classificador SMO apresentou o maior desempenho, em torno de 95% em relação a outras técnicas analisadas. Na validação do software BioViC a técnica detecção de contornos em conjunto com a técnica SMOOpC apresentaram uma contagem de leveduras que não diferenciou da contagem manual realizada por um especialista.

Palavras-chave: histograma de palavras visuais, cor, aprendizagem supervisionada, *saccharomyces cerevisiae*.

Abstract

This work addressed a way to automate the process of identifying yeasts that are important in the production of ethanol. Ethanol production relies heavily on the viable yeasts, as they are responsible for the fermentation process of sugarcane juice. This fermentation transforms the cane juice into ethanol and carbon dioxide. So keep a control population of viable yeast is a crucial task in the ethanol production process, and for this, analyzes are conducted to identify and count the yeast. Despite the importance of this step, the identification and counting is mostly done by visual way through a microscope, with a repetitive task and susceptible to errors.

In this work, computer vision and supervised learning techniques have been evaluated to automate the identification of yeasts. The main computer vision technique studied was the Bag of Visual Words that uses grayscale images, and extensions that use this technique to add color information, such as CCV (Color Coherence Vectors), CM (Color Moments), Boc (Bag-of-Color) and the OpC (Opponent Color). The attributes extracted from these algorithms were used for testing and training of classifiers obtained by supervised learning techniques. Among the techniques we use the Naive Bayes, KNN, SVM and J48 that are available in the Weka environment. The performance evaluation, with the percentage of correct classification was performed using ANOVA and the Friedman hypothesis tests.

The image dataset BioViC was used. This dataset has a set of yeast images captured in the laboratory under a microscope. The Neubauer chamber was used to assist counting of yeast, so that six replicates of four quadrants with Brix at concentrations of 3, 6 and 12 were defined for analysis. The images obtained were cut with 03 Brix separating the images of viable yeasts non-viable as well as the background. The 2614 images obtained and separated into three classes (feasible, unfeasible and background of screen) were used for training and identification.

During validation of the BioViC software, the combination of a contour

detection technique for image segmentation with the SMOOpC classifier presented the best results and seemed to compare well with human counting. However, this last experiment with BioViC was just exploratory and needs further research.

Keywords: Bag-of-Visual-Words, Opponent Color, Yeast, Supervised Learning, *Saccharomyces cerevisiae*.

Sumário

Sumário	xiv
Lista de Figuras	xvii
Lista de Tabelas	xix
Lista de Abreviaturas	xxi
1 Introdução	1
2 Extratores de Atributos	7
2.1 Histograma de Palavras Visuais (BoVW)	7
2.1.1 Detecção e Descrição de Pontos de Interesse	8
2.1.2 Criação de um Vocabulário	9
2.1.3 Geração do Histograma	9
2.2 Algoritmo BoC (<i>Bag-of-Color</i>)	9
2.3 Algoritmo CCV (<i>Color Coherence Vectors</i>)	11
2.3.1 Calculando o CCV	12
2.3.2 Um Exemplo	12
2.4 Algoritmo CM (<i>Color Moments</i>)	14
2.5 Algoritmo OpC (<i>Opponent Color</i>)	16
3 Materiais e Métodos	17
3.1 Classificadores	17
3.1.1 SVM	17
3.1.2 Árvore de Decisão	18
3.1.3 Naive Bayes	19
3.1.4 KNN	19
3.2 Validação Cruzada	20
3.3 Banco de Imagens	20
3.4 Testes de Hipóteses	22
3.5 Métrica	23

4 Experimentos e Resultados	24
4.1 Desempenho das Técnicas	24
4.1.1 ANOVA e Friedman	25
4.1.2 Pós-teste Tukey	26
4.2 Ajuste de Parâmetro da Técnica SMOOpC	26
4.3 Discussão	29
5 Software BioViC	31
5.1 Localização de leveduras	31
5.2 Validação do Software	33
5.3 Implementação do software BioViC	34
6 Considerações Finais	40
6.1 Trabalhos Futuros	41
Referências	45

Lista de Figuras

1.1	Imagem com leveduras viáveis (quadrado de cor vermelha), e leveduras inviáveis (quadrado de cor azul) capturadas através de um microscópio utilizando a câmera de Neubauer.	2
1.2	Imagem geral do desenvolvimento deste trabalho.	6
2.1	Imagem composta por três pontos de interesse, com cores diferentes. Dois pontos vermelhos que representam pontos de interesse que foram capturados em uma região com a variação entorno do ponto de interesse do escuro para o claro. Um ponto azul que representa um ponto de interesse detectado em uma região com a variação de claro para o escuro.	8
2.2	Imagem ilustrando os passos do algoritmo BoVW. (a) correspondem a detecção e descrição de pontos de interesse, (b) corresponde a criação do vocabulário e (c) a contagem de palavras visuais e (d) corresponde a criação do histograma.	10
2.3	Imagem com todos os passos realizados pelo algoritmo BoC. . . .	11
2.4	Com os pontos de interesse detectados, uma região é definida em torno de cada ponto e a partir desta região o CCV é calculado. . .	14
2.5	Com os pontos de interesse detectados, uma região é definida em torno de cada ponto e a partir desta região a CM é calculada. . .	15
3.1	Representação de duas classes separadas por um hiperplano que divide a margem máxima.	18
3.2	Estrutura de uma árvore de decisão extraída de um jogo de voleibol.	19
3.3	Espaço KNN composto por duas classes: triângulos e retângulos. O círculo com a cor verde representa um novo dado a ser classificado.	19

3.4	Validação Cruzada em 3 pastas (subconjuntos) realizada em três repetições. Em cada repetição, dois subconjuntos são utilizados para treino e um para teste do classificador, variando a ordem dos subconjuntos.	20
3.5	Imagem de um recipiente contendo corante azul de metileno. . . .	21
3.6	Imagem da camera de Neubauer, com amostra de leveduras. . . .	21
3.7	Imagem de leveduras retiradas por microscópio através da camera de Neubauer. Esta imagem contém leveduras (as formas circulares) e fundo com os retículos da câmara de Neubauer que corresponde as linhas na vertical e na horizontal.	22
4.1	Diagrama de caixa das técnicas analisadas. Cada técnica corresponde ao classificador combinado com um extrator de atributos.	25
4.2	Diagrama de caixa das variações do SMOOpC e o desempenho obtido.	27
4.3	Diagrama de caixa das variações do SMOOpC com ajuste automático e manual.	28
4.4	Diagrama de caixa das imagens, onde cada imagem está representada pela repetição, quadrante e a Imagem. Assim a Imagem R1Q1I2 representa a Repetição 1, Quadrante 1 e Imagem 2. . . .	28
4.5	A imagem R1Q1I4 que apresentou maior desempenho com a técnica SMOOpC com dicionário com o tamanho 256.	28
4.6	A imagem R2Q1I2 que apresentou o menor desempenho com a técnica SMOOpC com dicionário com o tamanho 256.	29
4.7	Ambas imagens classificadas como viáveis.	30
4.8	Ambas imagens classificadas como viáveis.	30
5.1	Imagem da tela responsável por detectar leveduras através da técnica de detecção de círculos.	32
5.2	Imagem da tela responsável por detectar leveduras através da técnica de detecção de contornos.	32
5.3	Imagem da tela responsável por detectar leveduras através da técnica <i>Template Matching</i>	33
5.4	Imagens com leveduras sem a câmara de Neubauer.	33
5.5	Imagens com o resultado das principais técnicas utilizadas para localizar as leveduras.	37
5.6	Imagem com o resultado da identificação e contagem de leveduras viáveis.	38
5.7	Imagem com o resultado da identificação e contagem de leveduras inviáveis.	38
5.8	Imagem da tela responsável por cadastrar usuário no sistema. . .	38

5.9 Imagem da tela responsável por cadastrar usina no sistema. . . .	38
5.10 Imagem da tela responsável por contar as leveduras.	39
5.11 Imagem da tela responsável por emitir relatório.	39
5.12 Imagem com o resultado da identificação e contagem de leveduras.	39

Lista de Tabelas

2.1	Valores da intensidade dos pixels da imagem de exemplo.	12
2.2	Rotulação da intensidade dos pixels na imagem exemplo.	13
2.3	As regiões coloridas representam os componentes conectados. . .	13
2.4	Resultado do cálculo dos componentes, onde cada componente possui uma área e a cor associada.	14
2.5	Definição das cores coerentes e não-coerentes. Cada cor tem a quantidade total das áreas coerentes e não-coerentes.	14
3.1	Resultado de recortes de leveduras e fundo.	22
4.1	Valores: A-IBkBoW, B-IBkBoC, C-IBkCCV, D-IBkCM, E-IBkOpC, F-J48BoC, G-J48BoW, H-J48CCV, I-J48CM, J-J48OpC, K-NBBBoC, L-NBBBoW, M-NBCCV, N-NBCM, O-NBOpC, P-SMOBoC, Q-SMOBoW, R-SMOCCV, S-SMOCM, T-SMOOpC.	26
4.2	Matriz de confusão.	29
5.1	Tabela com os valores p-value obtidos do teste Tukey das técnicas utilizadas no software.	34

Lista de Abreviaturas

SURF *Speeded-Up Robust Features*

SIFT *Scale-Invariant Feature Transform*

BoC Extrator de atributos *Bag-of-Color*

BoW Extrator de atributos *Bag-of-Word*

CCV Extrator de atributos *Color Coherence Vectors*

CM Extrator de atributos *Color Moments*

IBk Classificador *Instance-based learning k-Nearest Neighbours* (KNN)

OpC Extrator de atributos *Opponent Color*

J48 Classificador árvore de decisão

NB Classificador *Naive Bayes*

SMO Classificador *Support Vector Machine*

IBkBoC Classificador IBk e extrator BoC

IBkBoW Classificador IBk e extrator BoW

IBkCCV Classificador IBk e extrator CCV

IBkCM Classificador IBk e extrator CM

IBkOpC Classificador IBk e extrator OpC

J48BoC Classificador J48 e extrator BoC

J48BoW Classificador J48 e extrator BoW

J48CCV Classificador J48 e extrator CCV

J48CM Classificador J48 e extrator CM

J48OpC Classificador J48 e extrator OpC

NBBoC Classificador *Naive Bayes* e extrator BoC

NBBoW Classificador *Naive Bayes* e extrator BoW

NBCCV Classificador *Naive Bayes* e extrator CCV

NBCM Classificador *Naive Bayes* e extrator CM

NBOpC Classificador *Naive Bayes* e extrator OpC

SMOBoC Classificador SMO e extrator BoC

SMOBoW Classificador SMO e extrator BoW

SMOCCV Classificador SMO e extrator CCV

SMOCM Classificador SMO e extrator CM

SMOOpC Classificador SMO e extrator OpC

D1024 Técnica SMOOpC com o dicionário no tamanho de 1024

D128 Técnica SMOOpC com o dicionário no tamanho de 128

D2048 Técnica SMOOpC com o dicionário no tamanho de 2048

D256 Técnica SMOOpC com o dicionário no tamanho de 256

D512 Técnica SMOOpC com o dicionário no tamanho de 512

Introdução

Em um panorama estadual, a cultura da cana-de-açúcar foi inserida no Mato Grosso do Sul na década de 1980, após a adoção do Programa Nacional do Alcool (Proálcool). Com o objetivo de produzir álcool e açúcar, grandes áreas cultiváveis (em torno de 396,2 mil hectares em 2010/2011) e que antes eram utilizadas pela pecuária, passaram a ser utilizadas para o cultivo da cana. Devido ao aumento da área plantada e a elevação do número de indústrias de beneficiamento da cana, houve um reflexo na busca por tecnologias que melhorassem a qualidade e a produtividade nos processos de produção (Quinta et al., 2010). Entre os produtos extraídos da cana pela indústria, destacam-se o etanol e o açúcar.

A produção do etanol é caracterizada pela fermentação do mosto, resultado da diluição do caldo da cana com água. Neste processo as leveduras da espécie *Saccharomyces cerevisiae* são adicionadas ao mosto, com o propósito de produzir o etanol por fermentação. Este processo ocorre, com o consumo do açúcar do mosto pelas leveduras e consequente a produção de etanol e o gás-carbono. Para reduzir custos no processo de produção, a fermentação “MelleBoinot” é adotada devido a sua característica de reciclagem de leveduras. Nessa reciclagem, as leveduras são reutilizadas em fermentações consecutivas. O reaproveitamento ou reutilização das leveduras reduz custo no processo de produção (Boinot, 1939).

A qualidade da produção do etanol está diretamente relacionada à viabilidade das leveduras. Portanto, para garantir a produção, é necessário efetuar o controle microbiológico em laboratório, considerando que as leveduras viáveis são responsáveis pela fermentação e as leveduras inviáveis não desempenham o papel da fermentação como deveriam (Stratford, 1996).

Nas atividades relacionadas ao controle microbiológico, amostras são retiradas dos tanques de mosto, e examinadas em laboratório por um técnico responsável. Esta atividade consiste na identificação e contagem das leveduras, de maneira visual e com o auxílio de um microscópio. Por ser uma atividade repetitiva e visual, esta tarefa é suscetível a erros, pois este processo pode ser cansativo e subjetivo. Para facilitar a identificação das leveduras, as amostras são misturadas em água e corante azul de metileno, que colore em azul as leveduras inviáveis (Ceccato-Antonini, 2011). Podemos visualizar na Figura 1.1 dois tipos de leveduras. A viáveis que são marcadas por um quadrado cujas bordas possuem a cor vermelha e as leveduras inviáveis estão marcadas por um quadrado com as bordas em azul.

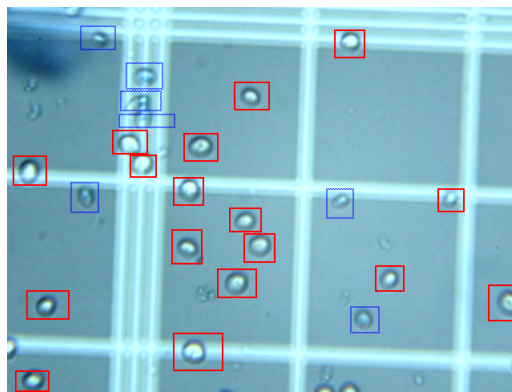


Figura 1.1: Imagem com leveduras viáveis (quadrado de cor vermelha), e leveduras inviáveis (quadrado de cor azul) capturadas através de um microscópio utilizando a câmera de Neubauer.

As atividades como a identificação e a contagem das leveduras, são tarefas repetitivas que podem ser realizadas automaticamente por meio de computadores. Como estas análises são feitas por meio de imagens, é possível automatizá-las por intermédio da visão computacional e aprendizagem supervisionada.

Através de técnicas de visão computacional é possível extrair atributos de imagens. Estes atributos estão relacionados à informação que podemos obter como: cor, forma, textura e etc. Estas informações podem ser utilizadas para realizar a identificação e/ou o reconhecimento de imagens, por meio de classificadores obtidos pelos algoritmos de aprendizagem supervisionada. A literatura apresenta algumas abordagens que utilizam a visão computacional para a contagem de leveduras e células.

Silva et al. (2012) propuseram um sistema para a contagem e a classificação de leveduras. As principais técnicas de visão computacional utilizadas foram: a transformada de Hough para localizar as leveduras na imagem, e a extração de atributos referentes à forma e cor para efetuar a identificação. Os classificadores utilizados para identificar as leveduras através dos atributos

foram: árvore de decisão, k-vizinhos mais próximos e redes neurais. O banco de imagens foi obtido em laboratório, através de experimentos realizados com a fermentação do caldo da cana com leveduras nas concentrações Brix 3, 6 e 12. Em cada concentração de açúcar as amostras foram retiradas e analisadas através da câmera de Neubauer. As imagens foram obtidas através de 6 repetições com 4 quadrantes, e em cada quadrante 4 imagens foram capturadas. Os resultados obtidos mostraram ser parecidos com a classificação manual realizada por um humano.

Outro trabalho com leveduras é o de Mongelo et al. (2011) que desenvolveram um software para efetuar a contagem e a identificação das leveduras em imagens. A técnica utilizada foi o algoritmo K-curvatura em regiões de interesse (ROI) para a extração de atributos referentes à forma e a árvore de decisão para identificação. O banco de imagens utilizado foi construído através de imagens coletadas por microscópio. Os autores realizaram os experimentos com a fermentação do caldo de cana em laboratório, conforme a avaliação feita em usinas de cana-de-açúcar. Os resultados obtidos pelos autores demonstraram que a aplicação não conseguiu diferenciar corretamente as leveduras viáveis e inviáveis. Isto porque, a extração de atributos utilizada foi em relação a forma enquanto que as leveduras são diferenciadas através da cor.

Schier (2011) utilizaram e avaliaram a técnica FRT (*Fast Radial Transform*) para contar colônias de leveduras viáveis em placas de Petri. Esta técnica foi utilizada para localizar leveduras na imagem, onde o maior problema foi contar as leveduras de forma separada, uma vez que cada colônia tinha um grupo de leveduras que muitas vezes se apresentavam sobrepostas umas as outras. Os testes foram realizados com 245 imagens contendo colônias de leveduras de diferentes formas. Os resultados mostraram que a técnica FRT apresentou uma taxa de erro menor que 0.04, para a contagem de leveduras em relação as colônias analisadas.

Coelho et al. (2009) apresentam uma análise de 8 técnicas para efetuar a segmentação com o intuito de localizar células em imagens coletadas de microscópio. As técnicas utilizadas foram: *AS Manual*, *RC Threshold*, *Otsu Threshold*, *Mean Threshold*, *Watershed (direct)*, *Watershed (gradient)*, *Active Masks* e *Merging Algorithm*. O objetivo desta análise foi encontrar a melhor técnica para segmentar a imagem de modo a destacar as células, facilitando a contagem e a identificação. Dois bancos de imagens foram utilizados (U2OS e NIH3T3), totalizando em 4009 imagens microscópicas de células. Para fazer a avaliação, os autores realizaram uma segmentação manual (*AS Manual*) de 97 imagens microscópica de células, para avaliar as segmentações resultantes das outras técnicas. Os resultados obtidos demonstraram que a técnica

Merging Algorithm apresentou os melhores resultados em todas as métricas que foram utilizadas para comparar o desempenho.

Entre as diferentes técnicas de extração de atributos, o algoritmo denominado Histograma de Palavras Visuais (*Bag-of-Visual-Words-BoVW*) é uma técnica de extração de atributos resultante de um vetor de ocorrência de um vocabulário constituído por palavras visuais. Uma palavra visual é um centroide calculado pelo algoritmo k-médias, que corresponde a média entre os pontos parecidos na imagem. Estes pontos são regiões localizadas por detetores de pontos de interesse, como por exemplo, o algoritmo SURF (Csurka et al., 2004).

Neste trabalho utilizamos o algoritmo BoVW para a extração de atributos, por ser uma técnica cada vez mais difundida em muitos trabalhos que relatam o uso na classificação de imagens. Este algoritmo, inicialmente foi proposto para imagens em níveis de cinza. Entretanto, a cor é uma característica muito importante nas imagens de leveduras com corante. Devido a sua importância, diversos estudos estão sendo direcionados para atribuir informação relacionada à cor no BoVW. Van Weijer and Khan (2013) desenvolveram um estudo ressaltando as vantagens e desvantagens dos dois principais métodos que são: *Late fusion* e *Early Fusion*. O primeiro consiste na adição de informação relativa à cor depois da quantização do histograma gerado pelo BoVW e o segundo método funciona de forma oposta, ou seja, a informação de cor é adicionada antes da quantização do histograma criado pelo BoVW. Com o método *Late fusion*, a identificação de um objeto é feita de maneira independente da cor, entretanto, a cor continua sendo uma informação a mais, por exemplo: é possível identificar imagens de pessoas independentes de cor da pele, mesmo que a cor seja importante para distinguir pessoas de objetos. Já no método *Early Fusion* a cor está totalmente conectada a identificação, por exemplo: é possível identificar imagens de pessoas que sejam de uma mesma cor.

Botterill et al. (2008) descreveram uma aplicação para o reconhecimento de cenas e objetos em tempo real para auxiliar na navegação de robôs. As técnicas utilizadas foram o BoVW em conjunto com o histograma de cor obtido do espaço de cor HSV. Os resultados obtidos foram relacionados ao tempo necessário para o reconhecimento, onde a taxa de tempo de reconhecimento se manteve em 0.0036 segundos por cada imagem, permitindo o reconhecimento em tempo real da cena.

O objetivo deste trabalho foi avaliar o desempenho obtido pelo BoVW através da adição de informação de cor e desenvolver uma aplicação para a identificação e contagem de leveduras, denominada BioViC. No que se refere à informação de cor foram utilizadas 4 técnicas: CCV Pass et al. (1997), CM Bahri and Zouaki (2013), BoC Wengert et al. (2011) e OpC van de Sande et al.

(2008). O CCV extrai informação de cor através de regiões ou aglomerados de uma mesma cor. O CM extrai informação de cor através da média e variância aplicada em cada imagem. O BoC é um histograma de cor, cujo objetivo é extrair a frequência de determinadas cores. OpC é uma variante que aplica o BoVW em cada canal de cor. No Capítulo 2 essas técnicas são abordadas em mais detalhes.

Na etapa de classificação, utilizamos as seguintes técnicas de aprendizagem supervisionada: árvore de decisão, Naives Bayes, Máquina de vetores de suporte e KNN. A árvore de decisão utilizada por Murthy (1998). O Naives Bayes utilizado por Kohavi (1996). KNN é utilizado por Salama et al. (2012). Máquina de vetores de suporte encontramos no trabalho de Madzarov et al. (2009).

Estas técnicas de aprendizagem supervisionada foram utilizadas com os diferentes extratores de atributos variantes do BoVW, que apresentamos anteriormente. Desta forma utilizamos a visão computacional como extrator de atributos e os algoritmos de aprendizagem supervisionada para gerar classificadores a partir dos atributos e efetuar a tarefa de identificação das leveduras. Na etapa de extração de atributos, foi utilizado um banco de imagens com 2614 imagens recortadas manualmente, definidas em 3 classes (fundo sem leveduras, levedura inviável e viável).

A avaliação de desempenho, neste trabalho, foi realizada através do teste de hipótese ANOVA, a fim de avaliarmos o quão diferentes ou não são as variantes do BoVW. O valor-p obtido através das técnicas foi $2e^{-16}$, o que indica que as técnicas são bem diferentes uma das outras, levando em consideração a porcentagem de classificação correta. A técnica que apresentou o maior desempenho segundo o ANOVA foi o OpC com o classificador SVM. Um segundo experimento foi realizado com o OpC e SVM, e o resultado mostrou que o dicionário com o tamanho 256 apresentou o melhor resultado, uma vez que, utilizamos dicionários maiores com o tamanho de 2048 nos experimentos.

Com a avaliação das técnicas de identificação de leveduras foi desenvolvida uma ferramenta para identificação automática de leveduras e disponibilizada pelo software BioViC. A ferramenta de identificação de leveduras necessita de imagens de possíveis leveduras para realizar a tarefa de extração de atributos e a identificação, e como as imagens retiradas em microscópio podem conter várias leveduras é necessário localizá-las. Para localizar as leveduras algumas técnicas de visão computacional como: detecção de contornos, detecção de círculos e Template matching (casamento de modelos). Foram utilizadas e avaliadas em relação à contagem de leveduras realizada por humanos. Na Figura 1.2 podemos visualizar uma imagem obtida em microscópio e em seguida a localização das leveduras e por fim a identificação das leveduras é

realizada.

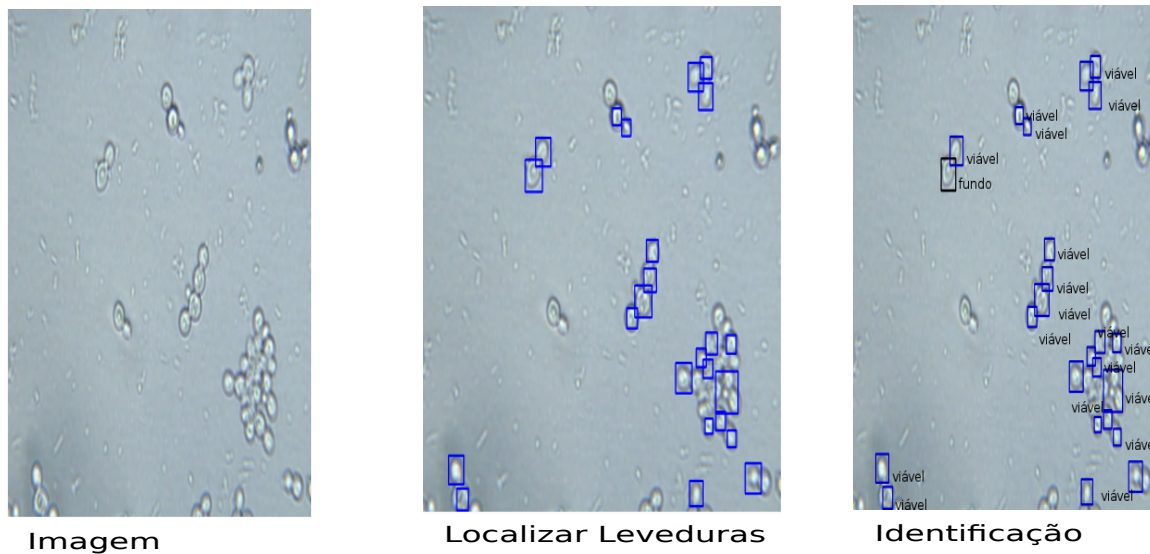


Figura 1.2: Imagem geral do desenvolvimento deste trabalho.

O Capítulo 2 descreve as técnicas de visão computacional utilizadas como extrator de atributos. O Capítulo 3 aborda os classificadores utilizados. O Capítulo 4 detalha o banco de dados e o teste de hipóteses utilizado para a comparação das técnicas. O Capítulo 5 descreve a implementação do software e avaliação de técnicas para localizar as leveduras e as considerações finais são descritas no Capítulo 6.

Extratores de Atributos

Para efetuar a identificação de imagens, é necessário extrair atributos que são utilizados na tarefa de distinguir e comparar imagens. As principais técnicas de visão computacional utilizadas neste trabalho para a extração de atributos são variações do algoritmo BoVW (*Bag-of-Visual-Words*) que utilizam informação de cor como fator adicional, pois, o algoritmo BoVW é aplicado somente em imagens em níveis de cinza.

A Seção 2.1 descreve o algoritmo BoVW, que é um extrator de atributos utilizado em imagens em nível de cinza, na Seção 2.2 o algoritmo BoC é apresentado e nas Seções 2.3 , 2.4 e 2.5 são apresentadas variações aplicadas ao BoVW. Todas estas técnicas foram implementadas utilizando a biblioteca OpenCV ¹ (em linguagem C++) e ImageJ ² (em linguagem java).

2.1 *Histograma de Palavras Visuais (BoVW)*

O histograma de palavras (*Bag-of-Words* - BoW) é um algoritmo utilizado no campo de reconhecimento de textos. Esta técnica extrai um histograma a partir da contagem das ocorrências de palavras-chave em um texto. Palavras-chave é um conjunto de palavras utilizadas para distinguir um texto, de forma a identificar, como por exemplo, o tipo do texto: tecnológico, poético, romântico e etc. Os valores do histograma resultante do BoW são utilizados por classificadores para efetuar a identificação.

Assim como no reconhecimento de textos, o algoritmo BoW passou a ser aplicado em imagens e denominado como Histograma de Palavras Visuais

¹<http://opencv.org/>

²<http://imagej.nih.gov/ij/>

(*Bag-of-Visual-Words*). Como o nome diz, as palavras passaram a ser visuais. Uma palavra visual é um ponto centroeide obtido pelo algoritmo k-médias, que representa um grupo de pontos de interesse semelhantes (Csurka et al., 2004). Na Figura 2.1 podemos visualizar pequenos pontos coloridos de vermelho e azul, que representam pontos de interesse. Os pontos vermelhos representam os pontos de interesse identificados em uma região com fundo claro, enquanto que os pontos azuis são pontos de interesse identificados em uma região com o fundo escuro.

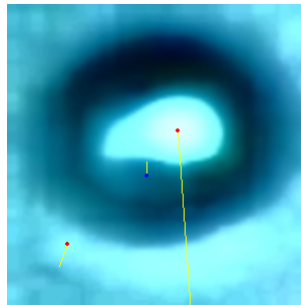


Figura 2.1: Imagem composta por três pontos de interesse, com cores diferentes. Dois pontos vermelhos que representam pontos de interesse que foram capturados em uma região com a variação entorno do ponto de interesse do escuro para o claro. Um ponto azul que representa um ponto de interesse detectado em uma região com a variação de claro para o escuro.

O algoritmo BoVW é composto pelos seguintes passos: detecção e descrição de pontos de interesse, criação do vocabulário e geração do histograma. Cada um destes passos são descritos nas subseções abaixo.

2.1.1 *Detecção e Descrição de Pontos de Interesse*

Em imagens, os pontos de interesse são pontos que se destacam e que podem ser utilizados em tarefas de reconhecimento. Estes pontos são regiões onde existe uma variação intensa dos pixels da vizinhança em diversas direções. Para encontrá-los, algoritmos como SURF Bay et al. (2008) e o SIFT (*Scale Invariant Feature Transform*) Alhwarin et al. (2008), foram propostos. Estes algoritmos detectam e descrevem os pontos de interesse através de vetores numéricos contendo valores referentes a direção das variações que ocorrem em torno do ponto.

Inúmeros trabalhos utilizam o algoritmo SURF como detector e também descritor de pontos de interesse, como o trabalho de Valgren and Lilienthal (2007), que investigaram o desempenho do SURF e do SIFT no reconhecimento de imagens de paisagens, retiradas no decorrer de um ano. Os resultados demonstraram que o SURF manteve o mesmo desempenho do SIFT, porém, com o tempo de processamento menor.

As principais características do SURF estão relacionadas com invariância à rotação, oclusão e intensidade de luz. Neste trabalho utilizamos o algoritmo SURF na fase de detecção e descrição de pontos de interesse no BoVW, por ser um algoritmo muito aplicado e também por ser mais rápido se comparado com o SIFT na detecção de pontos de interesse.

2.1.2 Criação de um Vocabulário

Como dito na Seção 2.1.1, cada ponto de interesse é descrito por um vetor contendo valores referentes a direção das variações que ocorrem em torno do ponto.

O vocabulário é obtido do conjunto de palavras visuais definidas através do agrupamento realizado pelo algoritmo k-médias de todo o conjunto de pontos de interesse extraídos pelo algoritmo detector e descritor de pontos de interesse em imagens Beltrame (2010). Desta forma, k-palavras são definidas para representarem todo o conjunto como um vocabulário, e utilizadas para criar um histograma. Estas k-palavras são definidas como palavras visuais.

2.1.3 Geração do Histograma

O histograma é o resultado da contagem das ocorrências das palavras visuais na imagem, usando o vocabulário definido na Seção 2.1.2. Novamente as palavras visuais são encontradas na imagem através do agrupamento dos pontos de interesse encontrados por um algoritmo detector e descritor de pontos de interesse, e então utilizadas para efetuar a contagem das ocorrências, resultando em um histograma. O resultado do histograma é utilizado como um vetor de atributos correspondentes as ocorrências das palavras visuais. Na Figura 2.2 podemos visualizar os passos do algoritmo BoVW.

2.2 Algoritmo BoC (Bag-of-Color)

O algoritmo BoC é uma técnica de extração de atributos em imagens, também conhecido como extrator de assinatura de cor (Wengert et al., 2011). Este algoritmo pode extrair informação relacionada à cor de forma global ou local. Na forma global os atributos estão relacionados com toda a imagem. Já na forma local os atributos estão relacionados com as regiões capturadas por algoritmos de extração de pontos de interesse (Wengert et al., 2011).

O espaço de cor utilizado pelo algoritmo BoC é o **CIELab** (*Commission Internationale de l'Eclairage, Luminance, red-green axis, blue-yellow axis*), pois é mais consistente com a métrica Euclidiana utilizada para efetuar a medida de distância $\sqrt{\sum_{i=1}^n (p_i + q_i)^2}$, onde n é o número de elementos, p e q são os elemen-

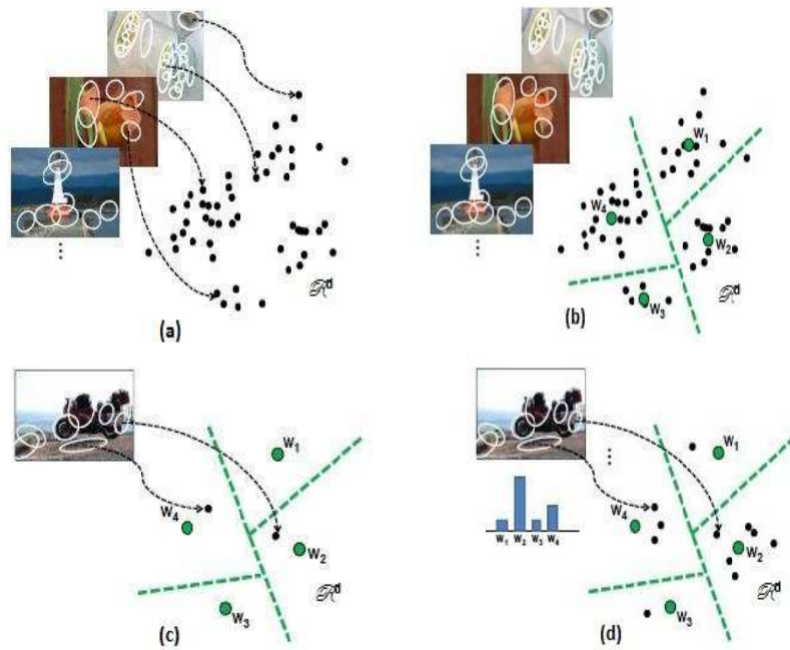


Figura 2.2: Imagem ilustrando os passos do algoritmo BoVW. (a) correspondem a detecção e descrição de pontos de interesse, (b) corresponde a criação do vocabulário e (c) a contagem de palavras visuais e (d) corresponde a criação do histograma.

tos. Dado um conjunto de imagens, o algoritmo pode ser representado com os seguintes passos:

1. Redimensionar cada imagem para 256x256 pixels e converter para o espaço de cor **CIELab**.
2. Para cada imagem redimensionada e convertida, realizar a divisão em blocos de 16x16 pixels.
3. Para cada bloco, realizar uma contagem para verificar qual é a cor de maior ocorrência. Caso a maior ocorrência seja menor que cinco, uma cor é escolhida aleatoriamente do bloco. Os resultados obtidos são colocados em um vetor também denominado como descritor, com 256 posições, onde cada posição corresponde a uma determinada cor de maior ocorrência.
4. Após extrair um conjunto de vetores de cor (descritores) é necessário encontrar as k-cores médias, que representam todo o conjunto de cores. Para isto, o algoritmo k-médias é utilizado para encontrar as k-cores que definiram o dicionário denominado **C**.
5. Um histograma de cor é obtido a partir do passo anterior, considerando um conjunto novo de imagens redimensionadas para 128x128 com um

número fixo de ≈ 16384 pixels. Para cada pixel é efetuado o cálculo da distância Euclidiana de cada pixel ao dicionário **C** e o incremento de cada **bin** (corresponde a contagem de uma intensidade de pixel) correspondente no histograma.

6. Com o resultado do cálculo do histograma extraído através do BoC é efetuada a normalização através de duas técnicas: *power-law* e *normalization*. A técnica *power-law* normaliza os valores do vetor resultante do histograma, dado por $x_i = \sqrt[n]{x_i}$, onde x_i corresponde valor de cada bin e n representa a quantidade de bins. A técnica *normalization* normaliza os valores do vetor, atualizando-o com a seguinte equação:

$$x_i = \frac{x_i}{\sum_{i=1}^n x_i} \quad (2.1)$$

Na Figura 2.3 podemos visualizar todos os passos que são realizados pelo algoritmo BoC.

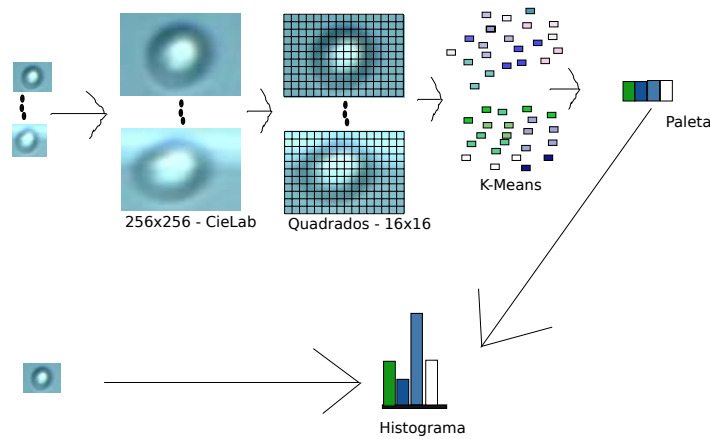


Figura 2.3: Imagem com todos os passos realizados pelo algoritmo BoC.

2.3 Algoritmo CCV (Color Coherence Vectors)

O algoritmo CCV utiliza o conceito de coerência para a extração de atributos referentes à cor. A coerência refere-se ao grau de similaridade apresentada em uma região de pixels. O grau de similaridade é definido como regiões que são representadas por uma mesma cor. O CCV obtém um histograma mais robusto em relação à informação espacial, que é um dos principais problemas encontrados em histogramas (Pass et al., 1997).

Esta técnica pode ser aplicada em uma imagem de maneira global ou local. A maneira global consiste em aplicar o algoritmo em toda a imagem. Já na

maneira local, o CCV é aplicado na região em torno de cada ponto de interesse capturado por algum detector de pontos de interesse, como por exemplo, o SURF.

2.3.1 Calculando o CCV

Os passos utilizados no CCV são os seguintes:

- Suavizar a imagem através do cálculo da média da vizinhança (os 8 vizinhos mais próximo de um determinado pixel) à cada pixel.
- Encontrar e rotular os componentes conectados.
- Efetuar a contagem de todos os pixels localizados em cada componente.
- Definir e aplicar um limiar para classificar a quantidade de elementos coerentes e não-coerentes para cada cor.

O espaço de cor utilizado é o RGB, como cada canal deste espaço utiliza 8 bits, utilizamos 2-bits para representar cada canal o que resulta em 64 cores, isto porque, foi reduzida a quantidade de cores e consequentemente obtido um conjunto de cores uniformes. Na Equação 2.2 podemos visualizar que cada canal do espaço RGB que é representado por 8 bits (expoente) passa a ser representado por 2 bits, o resultado final corresponde a 64 cores, reduzindo a quantidade de cor.

$$(RGB) = (2^8 \times 2^8 \times 2^8) \longrightarrow (2^2 \times 2^2 \times 2^2) = 64 \quad (2.2)$$

Na Seção 2.3.2 um exemplo de funcionamento é apresentado, neste exemplo utilizamos somente um canal.

2.3.2 Um Exemplo

Considere uma imagem com 6x6 pixels, a Tabela 2.1 corresponde a um canal do espaço **RGB** dessa imagem. Neste exemplo cada pixel possui um valor entre 10 à 48, escolhidos aleatoriamente.

22	10	21	22	15	16
24	21	13	20	14	17
23	17	48	23	17	16
25	25	22	14	15	14
27	22	12	11	17	18
24	21	10	12	15	19

Tabela 2.1: Valores da intensidade dos pixels da imagem de exemplo.

Neste exemplo cada pixel é rotulado usando a seguinte regra: cor 1 é definida para o intervalo (1 à 19), cor 2 é definida para o intervalo (20 à 38) e etc. Estes intervalos foram definidos no trabalho de Pass et al. (1997). A rotulação dos pixels da Tabela 2.1 pode ser visualizada na Tabela 2.2.

2	1	2	2	1	1
2	2	1	2	1	1
2	1	3	2	1	1
2	2	2	1	1	1
2	2	1	1	1	1
2	2	1	1	1	1

Tabela 2.2: Rotulação da intensidade dos pixels na imagem exemplo.

Após a rotulação é possível encontrar os componentes conectados. Um componente conectado é um conjunto de pixels que definem uma determinada região. Para encontrar e contar os componentes conectados, é necessário observar a região em torno de cada pixel e rotular todos os vizinhos que possuam a mesma cor. Na Tabela 2.3 podemos visualizar os componentes rotulados e coloridos.

B	C	B	B	A	A
B	B	C	B	A	A
B	C	D	B	A	A
B	B	B	A	A	A
B	B	A	A	A	A
B	B	A	A	A	A

Tabela 2.3: As regiões coloridas representam os componentes conectados.

A área de cada componente conectado é calculada através da soma da quantidade de pixels pertencentes ao componente conectado. Uma vez feita a contagem de todos os componentes, um limiar é definido como um valor em torno de 1% do tamanho total da imagem. Neste exemplo o valor 14 foi definido, assim valores maiores e/ou iguais são considerados como coerentes e menores são considerados como não-coerentes. Na Tabela 2.4 podemos visualizar a contagem de cada componente e a sua relação com os rótulos definidos como cor.

Assumimos que para uma cor ser coerente a sua frequência em cada região (componente) deve ser maior ou igual a um limiar $\gamma = 14$, caso contrário é considerada como não-coerente. Na Tabela 2.5 podemos visualizar o resultado final com os pares de cores coerentes e não-coerentes, este resultado consiste na soma para cada cor, suas parcelas de coerentes e não-coerentes.

Componente	A	B	C	D
Cor	1	2	1	3
Área	17	15	3	1

Tabela 2.4: Resultado do cálculo dos componentes, onde cada componente possui uma área e a cor associada.

Cor	1	2	3
Coerente	17	15	0
Não-coerente	3	0	1

Tabela 2.5: Definição das cores coerentes e não-coerentes. Cada cor tem a quantidade total das áreas coerentes e não-coerentes.

Desta forma assumimos que existem 3 cores que podem ser descritas por um par de valores, onde para cada par o primeiro valor corresponde a coerente e o segundo valor corresponde a não-coerente:

$$\langle (17,3), (15,0), (0,1) \rangle \quad (2.3)$$

O vetor resultante com os pares coerentes e não-coerentes, foi adicionado ao conjunto de atributos extraídos pelo algoritmo BoVW. Assim, a informação de cor foi extraída da região com dimensão de 6x6 pixels entorno de cada ponto de interesse localizado pelo algoritmo. Na Figura 2.4 podemos visualizar os passos do algoritmo CCV.

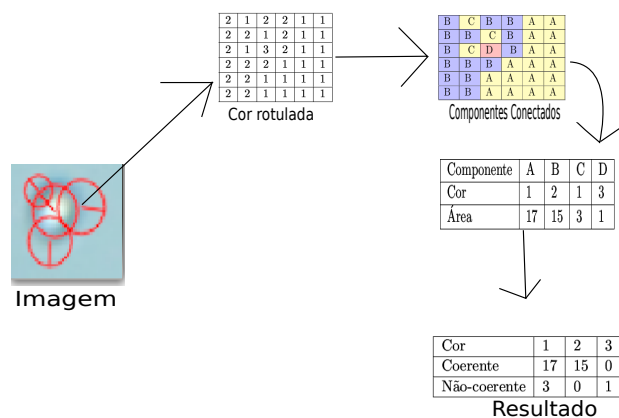


Figura 2.4: Com os pontos de interesse detectados, uma região é definida em torno de cada ponto e a partir desta região o CCV é calculado.

2.4 Algoritmo CM (Color Moments)

O algoritmo CM foi proposto por Bahri and Zouaki (2013) como uma extensão do algoritmo SURF, denominada *SURF-Color Moments* (corresponde ao

SURF + *Color Moments*). Esta extensão adiciona a informação de cor ao algoritmo SURF. A informação de cor é extraída na região entorno de cada ponto de interesse localizado pelo SURF, através do cálculo da média e da variância dos valores dos pixels.

As informações extraídas pelo CM foram definidas como momentos de primeira ordem (média) e momentos de segunda ordem (variância). As Equações 2.4 e 2.5 representam respectivamente a média e a variância, dado que P_{ij} é a intensidade do pixel j do i -ésimo canal, e N é o número total de pixels.

$$E_i = \frac{1}{N} \sum_{j=1}^N P_{ij} \quad (2.4)$$

$$\sigma_i = \left[\frac{1}{N} \sum_{j=1}^N (P_{ij} - E_i)^2 \right]^{\frac{1}{2}} \quad (2.5)$$

Como podemos visualizar na Figura 2.5, para cada ponto de interesse detectado pelo SURF, as estatísticas são calculadas em regiões de 5×5 segundo o artigo de Bahri and Zouaki (2013). Um vetor é preenchido com os resultados dos cálculos da média e da variância para cada canal da imagem, assim em cada ponto de interesse detectado na imagem no espaço RGB teremos um vetor de 6 posições. Este vetor é concatenado ao vetor que descreve o ponto de interesse.

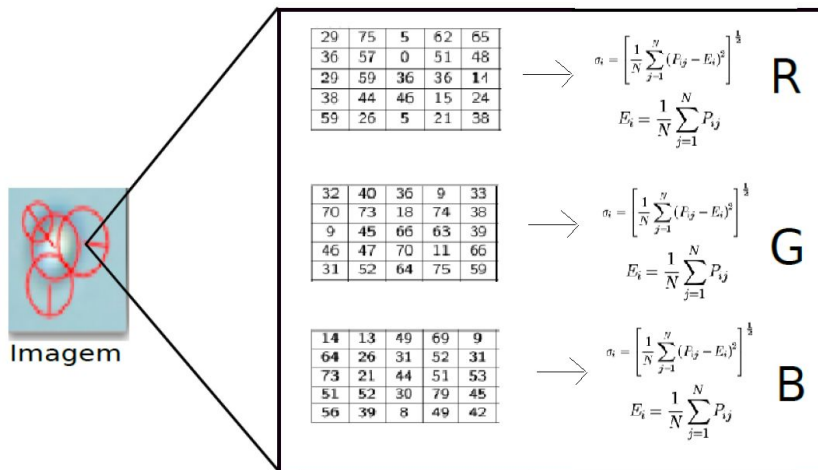


Figura 2.5: Com os pontos de interesse detectados, uma região é definida em torno de cada ponto e a partir desta região a CM é calculada.

2.5 Algoritmo OpC (Opponent Color)

Em van de Sande et al. (2008) é abordada a utilização do SURF em imagens coloridas, com o objetivo de utilizar a informação de cor associada aos pontos de interesse detectados. Esta abordagem é definida como OpC. A técnica OpC é uma extensão do algoritmo SURF. A principal mudança que o algoritmo OpC apresenta em relação ao SURF, está na matriz Hessiana que é responsável pela detecção dos pontos de interesse. O valor do determinante dessa matriz é utilizado como uma referência de um possível ponto de interesse, uma vez, que quanto maior for este valor, maiores serão as mudanças entorno de cada ponto. Podemos visualizar a representação de uma matriz Hessiana através da Equação 2.6.

No algoritmo SURF a matriz Hessiana é aplicada em imagens com um único canal, enquanto que, no OpC a matriz Hessiana é aplicada em imagens com três canais. O somatório do resultado da matriz Hessiana aplicada em cada canal é utilizado para encontrar os pontos de interesse.

$$H(x, \sigma) = \begin{vmatrix} L_{x,x}(x, \sigma) & L_{x,y}(x, \sigma) \\ L_{x,y}(x, \sigma) & L_{y,y}(x, \sigma) \end{vmatrix} \quad (2.6)$$

O algoritmo OpC antes de efetuar a detecção dos pontos de interesse aplica um ajuste nos canais da imagem através da técnica *Boosting Color*. A técnica *Boosting Color* utiliza o conceito de cores oponentes.

O sistema visual humano possui três tipos de células cones que são responsáveis por capturar a informação de cor. Cada cone é mais sensível a uma determinada cor primária, porém no sistema visual cada cor é registrada através dos três cones juntos. Isto quer dizer que, uma determinada cor não pode ser registrada somente por um cone, e sim através dos três juntos. Existem três canais oponentes: vermelho-verde, azul-amarelo e preto-branco. Os canais oponentes definem as cores oponentes.

Na Equação 2.7 temos o canal O_3 que representa a intensidade, e os canais O_1 e O_2 que representam respectivamente os pares oponentes vermelho-verde e amarelo-azul.

$$\begin{pmatrix} O_1 \\ O_2 \\ O_3 \end{pmatrix} = \begin{pmatrix} \frac{R-G}{\sqrt{2}} \\ \frac{R+G-2B}{\sqrt{6}} \\ \frac{R+G+B}{\sqrt{3}} \end{pmatrix} \quad (2.7)$$

Materiais e Métodos

Os experimentos foram realizados utilizando os seguintes extratores de atributos: BoVW, BoC, e variantes do BoVW (CCV, CM e OpC). Na seção 3.1 são descritos os principais classificadores utilizados. Na Seção 3.2 é definida a principal técnica de amostragem utilizada. O banco de imagens utilizado é descrito na Seção 3.3. A métrica utilizada para comparar os experimentos é relatada na Seção 3.5.

3.1 *Classificadores*

As subseções seguintes descrevem os principais algoritmos denominados classificadores, que utilizamos para efetuar a identificação das leveduras através dos atributos extraídos pelas técnicas abordadas no Capítulo 2. Estes classificadores foram utilizados no ambiente Weka¹. O Weka é um software desenvolvido em 1993, pela universidade da Nova Zelândia para analisar informações de bases de dados referentes a agricultura. Este software implementa diversos algoritmos estudados na área de inteligência artificial. Apesar de ser inicialmente desenvolvido em linguagem C, o Weka passou a ser escrito, utilizado e disponibilizado em linguagem java, para ser usado em multiplataformas (windows, linux, mac e etc.) e está disponível sob a licença GNU.

3.1.1 *SVM*

Witten and Frank (2005) definem SVM (*Support Vector Machines*) como algoritmos de vetores de suporte, por se tratar de vários algoritmos e não de

¹<http://www.cs.waikato.ac.nz/ml/weka/>

máquinas. Este classificador corresponde a classe de classificadores lineares. O SVM separa as classes em hiperplanos e busca a margem máxima entre os mesmos. Dadas duas classes linearmente separáveis a margem máxima é definida como a região que separa as duas classes. Os vetores de suporte correspondem aos dados mais próximos do hiperplano da margem máxima.

Na Figura 3.1 podemos visualizar duas classes, onde uma corresponde aos círculos conectados de cor cinza, e outra classe corresponde aos círculos conectados de cor azul. A linha pontilhada corresponde à distância entre dois vetores de suporte que estão mais próximos do hiperplano da margem máxima.

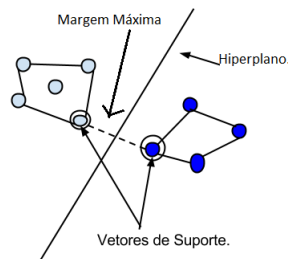


Figura 3.1: Representação de duas classes separadas por um hiperplano que divide a margem máxima.

Em uma classe pode existir um ou mais vetores de suporte, que são utilizados para separar as classes. A margem máxima é a região que separa as classes e é calculada a partir da distância entre os vetores de suporte entre classes diferentes. O algoritmo SVM a princípio foi aplicado em problemas que envolviam duas classes, porém foi estendido para múltiplas classes (Platt, 1999).

3.1.2 Árvore de Decisão

A árvore de decisão é apresentada sob a forma de uma estrutura hierárquica construída recursivamente, através da divisão e conquista. A árvore de decisão é composta por nós, cujos nós denominados folhas representam as classes e os nós internos representam as decisões a serem tomadas. Este classificador é utilizado em muitos trabalhos como o de Salama et al. (2012) que utilizou na identificação de células de câncer de mama.

Na Figura 3.2, podemos visualizar um exemplo de árvore de decisão para encontrar o perfil de um jogador de voleibol, onde a raiz com o rótulo idade classifica o jogador pela idade, de tal forma que, menor que 14 anos não pode jogar, já maior e com a altura maior que 1,70 poderá entrar no perfil de um jogador. Os quadrados representam as classes (nós folhas), e as elipses representam os nós internos.

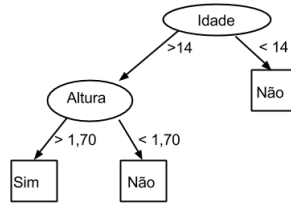


Figura 3.2: Estrutura de uma árvore de decisão extraída de um jogo de voleibol.

3.1.3 Naive Bayes

Este classificador utiliza o conceito da probabilidade condicional de Bayes e assume que os atributos são independentes. Para os dados reais e independentes, este classificador apresenta bons resultados. Entretanto, o seu desempenho está muito relacionado com a dependência dos dados, caindo drasticamente para os dados redundantes Witten and Frank (2005).

3.1.4 KNN

No KNN (*K-Nearest Neighbor*) a classificação é feita através da similaridade dos elementos. As aproximações locais são calculadas, ou seja, cada ponto é classificado pelos seus k -vizinhos mais próximos, através de medidas de distância como, por exemplo, a distância Euclidiana (Salama et al., 2012).

Se consideramos $X = (x_1, x_2, \dots, x_n)$ e $Y = (y_1, y_2, \dots, y_n)$ dois pontos, a medida Euclidiana será definida pela fórmula logo abaixo:

$$d(x, y) = \sqrt{(x_1 - y_1)^2 + (x_2 - y_2)^2 + \dots + (x_n - y_n)^2} \quad (3.1)$$

Na Figura 3.3 podemos observar duas classes representadas por quadrados de cor azul e triângulos de cor amarela. Um círculo com a cor verde representa um novo dado a ser classificado. Quando o valor de k é igual a 3 o círculo tem como vizinhança dois triângulos e um quadrado, e então é classificado como triângulo. Quando o valor de k é igual a 5 teremos 5 quadrados e 2 triângulos, sendo o círculo classificado como quadrado.

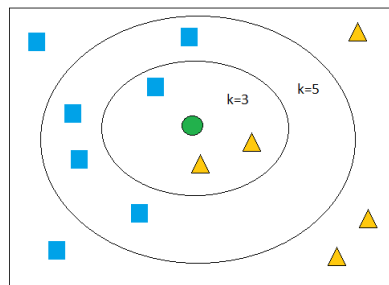


Figura 3.3: Espaço KNN composto por duas classes: triângulos e retângulos. O círculo com a cor verde representa um novo dado a ser classificado.

3.2 Validação Cruzada

A técnica Validação Cruzada permite distribuir um conjunto de dados em um número fixo de partições com quantidades aproximadamente iguais. Uma quantidade das partições é utilizada no treinamento de um classificador enquanto o restante fica para teste, sendo o número maior de partições para o treinamento. Na Figura 3.4 podemos visualizar um exemplo de Validação Cruzada. Os dados são repartidos em três partições (subconjuntos) e são realizadas três repetições, variando com duas partições para treino e uma partição para teste, com o classificador.

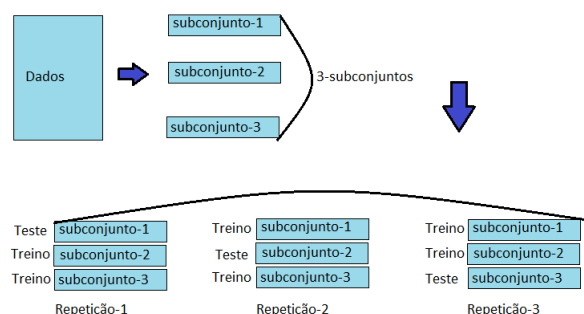


Figura 3.4: Validação Cruzada em 3 pastas (subconjuntos) realizada em três repetições. Em cada repetição, dois subconjuntos são utilizados para treino e um para teste do classificador, variando a ordem dos subconjuntos.

3.3 Banco de Imagens

As imagens utilizadas neste trabalho foram extraídas do banco de imagens de leveduras². Este banco foi construído pelo grupo INOVISÃO (Desenvolvimento e Inovação em Visão Computacional). A fermentação foi obtida através da adição de leveduras *Saccharomyces cerevisae* ao mosto na concentração de 1% (p/v)³. O mosto foi ajustado em 12° Brix sendo também utilizadas as amostragens quando o valor do Brix reduziu-se para 6 e 3. O Brix corresponde a medida utilizada para avaliar a concentração de açúcar no caldo da cana. Sendo definido por peso/peso de sólidos solúveis numa solução impura.

Para efetuar a contagem, as amostras contendo as leveduras foram misturadas em corante azul de metileno, que coloriu com a cor azul as leveduras com baixa atividade fisiológica. A câmera de Neubauer⁴ foi utilizada para auxiliar na contagem das leveduras. Esta câmera possui pequenos quadrados que

²<http://biovic.weebly.com/bancos-de-imagens.html>

³Por volume.

⁴Câmera retangular grossa utilizada em laboratório para auxiliar na contagem de células por unidade de volume em suspensão.

facilitam o processo de contagem visual das leveduras, além de permitir que cada amostra analisada tenha seu volume já previamente definido. As imagens obtidas foram retiradas através do microscópio óptico com o aumento de 400X. Na Figura 3.5 podemos visualizar um recipiente contendo corante azul de metileno e na Figura 3.6 podemos visualizar uma câmera de Neubauer.



Figura 3.5: Imagem de um recipiente contendo corante azul de metileno.

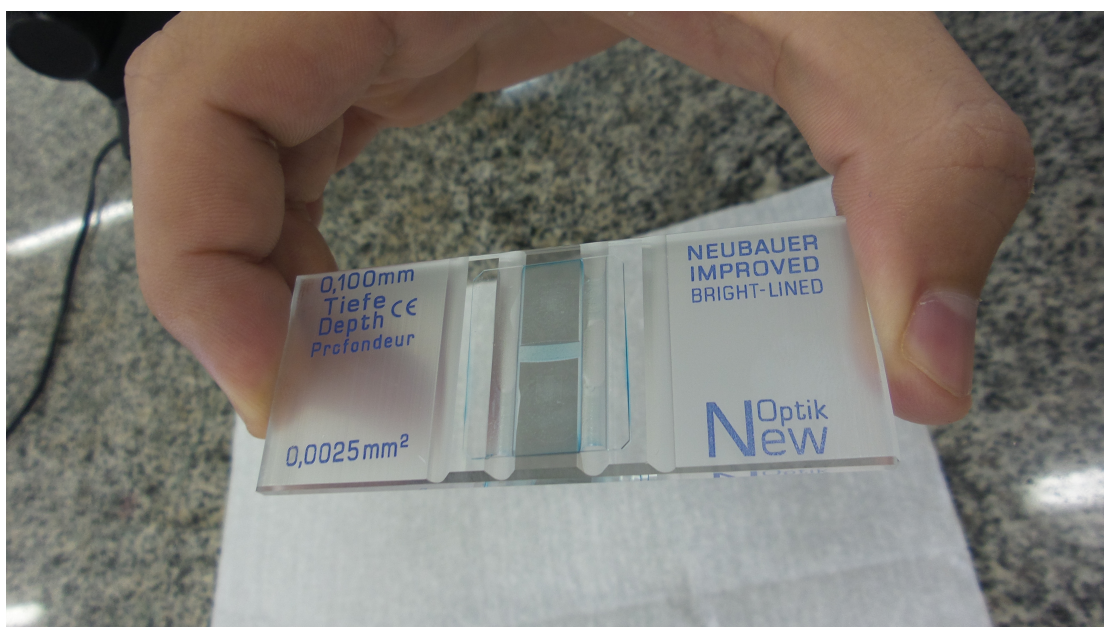


Figura 3.6: Imagem da camera de Neubauer, com amostra de leveduras.

Três tipos de concentração de açúcar (Brix 11, 6 e 3) foram utilizados na amostragem em que as imagens foram obtidas através de microscópio. Para cada valor de Brix foram coletadas 6 amostras (repetições), onde para cada repetição foram utilizados 6 quadrantes da câmera de Neubauer, e para cada

quadrante, 4 imagens foram retiradas. A Figura 3.7 representa uma imagem na concentração Brix 3, na repetição 1 representando o quadrante 1. Esta imagem contém leveduras viáveis e inviáveis.

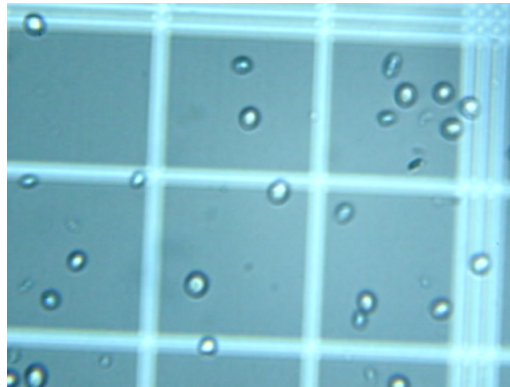


Figura 3.7: Imagem de leveduras retiradas por microscópio através da camera de Neubauer. Esta imagem contém leveduras (as formas circulares) e fundo com os retículos da câmera de Neubauer que corresponde as linhas na vertical e na horizontal.

Foram recortadas 2614 imagens manualmente de 30 imagens obtidas do Brix 03, e definidas 3 classes de imagens que são: inviável, viável e fundo. A classe viável corresponde às leveduras com alta atividade biológica e que são responsáveis pela fermentação e produção do etanol. A classe inviável corresponde as leveduras com baixa atividade biológica, já a classe fundo corresponde ao fundo da imagem, que no caso da câmera de Neubauer corresponderam aos retângulos delimitadores e também restante da imagem que não possui nenhuma levedura. Na Figura 3.1 podemos visualizar exemplos de recortes, com leveduras viáveis, inviáveis e fundo.

Viável	Inviável	Fundo
		
		

Tabela 3.1: Resultado de recortes de leveduras e fundo.

3.4 Testes de Hipóteses

Os principais testes de hipóteses que utilizamos neste trabalho foram o ANOVA e o Friedman. Os pós-testes utilizados foram o Tukey usado depois do ANOVA e Wilcoxon usado depois do Friedman.

- ANOVA (Análise de Variância): é uma técnica que permite examinar as diferenças entre as médias de populações, para verificar o quão próximas

ou não estão das médias. Para utilizar o ANOVA é necessário assumir que as amostras são aleatórias e independentes. Desta forma, observamos se as diferenças amostrais são reais (realmente ocorrem) ou casuais (mera decorrência) (Larson, 2008).

- Friedman: é um teste não paramétrico, que faz poucas ou nenhuma suposição sobre os dados que estão sendo utilizados. Este teste é uma alternativa ao ANOVA, quando não existe nenhuma suposição sobre a variância (Friedman, 1937).
- Tukey: é um pós-teste para comparar as médias, servindo de complemento ao ANOVA. Ele efetua múltiplas comparações entre pares de médias para o conjunto de dados.
- Wilcoxon: é um pós-teste para comparar às médias, servindo de complemento ao Friedman. Ele efetua múltiplas comparações entre pares de médias para o conjunto de dados.

3.5 Métrica

A principal métrica utilizada para a comparação das técnicas foi a porcentagem de classificação correta. A porcentagem de classificação correta corresponde ao número de imagens identificadas corretamente de todas as classes, dividida pelo total de imagens.

Experimentos e Resultados

As técnicas de extração de atributos descritas no Capítulo 2 foram utilizadas com os classificadores relatados no Capítulo 3. Os experimentos foram realizados com o objetivo de encontrar os melhores resultados que estão descritos de forma que, a Seção 4.1 apresenta o desempenho de todas as técnicas e a Seção 4.2 descreve um novo experimento com a técnica que apresentou o maior desempenho. Os problemas encontrados estão descritos na Seção 4.3. A métrica utilizada para a comparação das técnicas foi a porcentagem de classificação correta.

4.1 *Desempenho das Técnicas*

Foi utilizado o banco de imagens da Seção 3.3 do Capítulo 3, e aplicadas as técnicas de extração de atributos com os classificadores estudados: KNN, Árvore de decisão, NB e SMO que são denominados respectivamente como: IBk, J48, Naive Bayes e SMO. O desempenho foi analisado através do teste de hipótese ANOVA e Friedman, para comparar o desempenho das técnicas com relação a métrica de porcentagem de classificação correta. O tamanho do dicionário utilizado pelo BoW e as variantes que usam a informação de cor, foi definido inicialmente com o valor 512, pois foi um valor entre 64 e 1024, utilizado em outros trabalhos.

Foram realizadas combinações de classificadores obtidos no Weka (IBk definido como KNN, J48 uma implementação do algoritmo C4.5 que é uma árvore de decisão, NB que é o Naive Bayes e SMO que é uma implementação otimizada do SVM) com os extratores de atributos (BoC, BoW, CCV, CM e OpC), assim por exemplo, a sigla IBkBoC indica a análise do classificador IBk em

conjunto com o extrator de atributo BoC. Analogamente, são definidas as seguintes siglas: IBkBoC, IBkBoW, IBkCCV, IBkCM, IBkOpC, J48BoC, J48BoW, J48CCV, J48CM, J48OpC, NBBoC, NBBoW, NBCCV, NBkCM, NBkOpC, SMOBoC, SMOBoW, SMOCCV, SMOCM e SMOOpC.

A Figura 4.1 representa o diagrama de caixa obtido do software R¹. Neste diagrama podemos visualizar o desempenho de cada técnica através da comparação das medianas definidas pela faixa mais escura localizada em cada caixa. Observando o diagrama, podemos verificar que a técnica SMOOpC apresentou a mediana com o maior desempenho em relação as outras medianas. Esta técnica corresponde ao classificador SMO com o extrator de atributos *Opponent Color*. Podemos observar alguns comportamentos incomuns, como as combinações do extrator de atributos BoC com os classificadores, quase que apresentam os melhores desempenhos, exceto com o classificador SMO. As combinações do classificador IBk com os extratores de atributos, apresentam quase que todos os piores resultados. As combinações de extratores de atributos com o classificador SMO, quase que apresentam os melhores resultados.

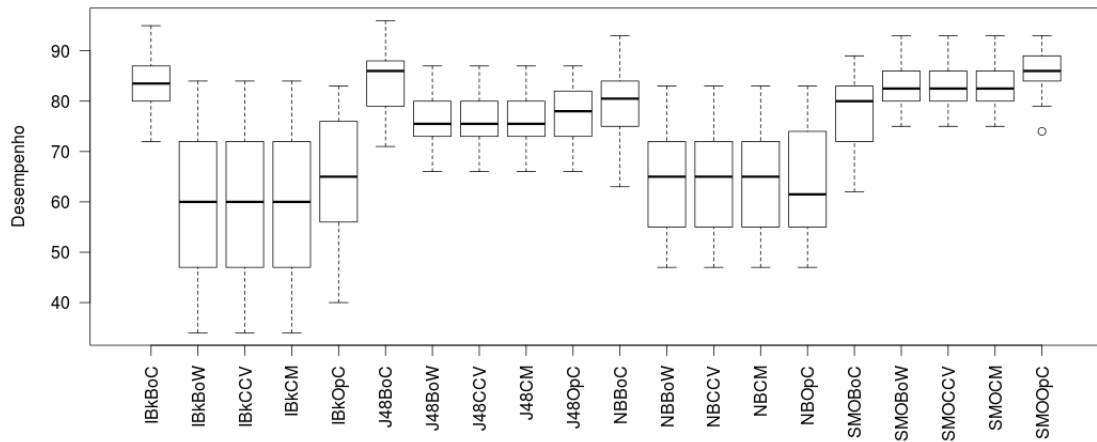


Figura 4.1: Diagrama de caixa das técnicas analisadas. Cada técnica corresponde ao classificador combinado com um extrator de atributos.

4.1.1 ANOVA e Friedman

Para validar os resultados e mostrar que as diferenças amostrais são significativas através de hipóteses, aplicamos a análise de variância e obtivemos o valor- $p < 2e^{-16}$. Este resultado indica que a hipótese nula pode ser descartada, ou seja, as médias indicam que existe uma diferença estatística entre as técnicas.

¹Software estatístico, mais informações em: <http://www.r-project.org/>

O teste de Friedman similar ao ANOVA, também foi aplicado. O resultado obtido a partir do valor-p foi $< 2.2e^{-16}$, o que descarta a hipótese nula. Novamente, o resultado indicou uma diferença entre as técnicas utilizadas.

4.1.2 Pós-teste Tukey

Na Tabela 4.1 podemos visualizar os resultados do teste de Tukey, neste teste as técnicas são comparadas em pares. Nesta tabela podemos visualizar cada célula com o valor-p. As células são coloridas conforme o valor-p, assim, quanto mais escura for a cor azul, maior será o valor-p. Os valores menores que 0.05 indicam o descarte da hipótese nula, isto é, existe uma diferença estatisticamente significativa entre as técnicas. Os valores 0.00 indicam um valor-p muito baixo.

Podemos visualizar na Tabela 4.1 que a técnica SMOOpC(T) que apresentou o maior desempenho em relação as outras técnicas, aparece semelhante as técnicas: IBkBoC, J48BoC, SMOBoW, SMOCV e SMOCM. Isto porque, estas técnicas apresentaram um desempenho maior em alguns conjunto de imagens do que outros, porém, a técnica SMOOpC manteve um maior desempenho no resultado geral.

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T
A																				
B	0.00																			
C	1.00	0.00																		
D	1.00	0.00	1.00																	
E	0.36	0.00	0.36	0.36																
F	0.00	0.99	0.00	0.00	0.00															
G	0.00	0.02	0.00	0.00	0.00	0.00														
H	0.00	0.02	0.00	0.00	0.00	0.00	1.00													
I	0.00	0.02	0.00	0.00	0.00	0.00	1.00	1.00												
J	0.00	0.14	0.00	0.00	0.00	0.01	1.00	1.00	1.00											
K	0.00	0.83	0.00	0.00	0.00	0.32	0.98	0.98	0.98	0.99										
L	0.12	0.00	0.12	0.12	1.00	0.00	0.00	0.00	0.00	0.00	0.00									
M	0.12	0.00	0.12	0.12	1.00	0.00	0.00	0.00	0.00	0.00	0.00	1.00								
N	0.12	0.00	0.12	0.12	1.00	0.00	0.00	0.00	0.00	0.00	0.00	1.00	1.00							
O	0.39	0.00	0.39	0.39	1.00	0.00	0.00	0.00	0.00	0.00	0.00	1.00	1.00	1.00						
P	0.00	0.28	0.00	0.00	0.00	0.04	0.99	0.99	0.99	1.00	0.99	0.00	0.00	0.00	0.00					
Q	0.00	1.00	0.00	0.00	0.00	1.00	0.01	0.01	0.01	0.09	0.74	0.00	0.00	0.00	0.00	0.20				
R	0.00	1.00	0.00	0.00	0.00	1.00	0.01	0.01	0.01	0.09	0.74	0.00	0.00	0.00	0.00	0.20	1.00			
S	0.00	1.00	0.00	0.00	0.00	1.00	0.01	0.01	0.01	0.09	0.74	0.00	0.00	0.00	0.00	0.20	1.00	1.00		
T	0.00	0.99	0.00	0.00	0.00	0.99	0.00	0.00	0.00	0.00	0.03	0.00	0.00	0.00	0.00	0.00	0.99	0.99	0.99	

Tabela 4.1: Valores: A-IBkBoW, B-IBkBoC, C-IBkCCV, D-IBkCM, E-IBkOpC, F-J48BoC, G-J48BoW, H-J48CCV, I-J48CM, J-J48OpC, K-NBBoC, L-NBBoW, M-NBCCV, N-NBCM, O-NBOPC, P-SMOBoC, Q-SMOBoW, R-SMOCCV, S-SMOCM, T-SMOOpC.

4.2 Ajuste de Parâmetro da Técnica SMOOpC

A Seção 4.1 mostrou que a técnica SMOOpC apresentou o melhor desempenho. A técnica SMOOpC por ser uma técnica variante do algoritmo BoW,

permite a mudança do tamanho do dicionário. O tamanho do dicionário do *Opponent Color* foi ajustado para: D128, D256, D512, D1024 e D2048.

Na Figura 4.2 podemos observar que a técnica que apresentou o melhor resultado foi SMOOpC com um dicionário de 256. O valor-p obtido pelo ANOVA foi $1.58e^{-8}$, o que indica que a hipótese nula pode ser descartada, assim estas variações são bem diferentes uma das outras.

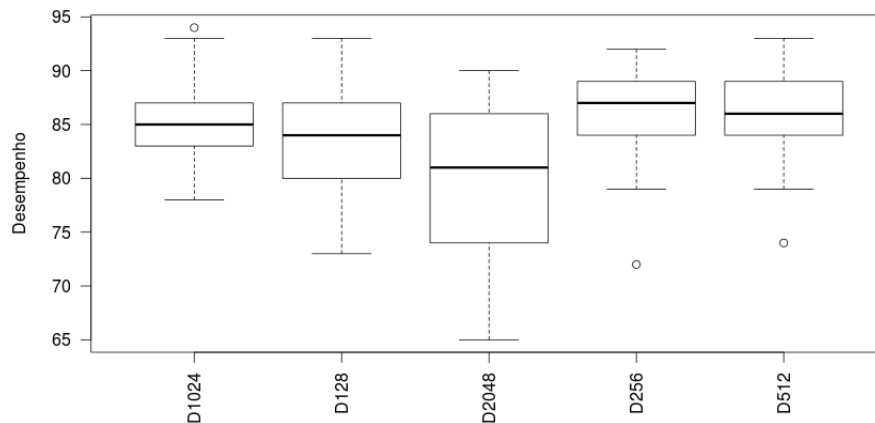


Figura 4.2: Diagrama de caixa das variações do SMOOpC e o desempenho obtido.

Com a ferramenta libSVM² foi realizado um novo experimento. Esta ferramenta fornece um meio de encontrar os valores dos parâmetros c (constante de complexidade) e g (gamma) definidos pelo SMO com o maior desempenho. Na Figura 4.3 podemos visualizar os resultados da técnica SMOOpC, variando o tamanho do dicionário com ajuste de parâmetros através do libSVM identificado com o sufixo Lib, exemplo: D256Lib e com ajuste de parâmetros realizado manualmente, exemplo D256. Os resultados demonstram que D256 continua com a mediana maior, e com o mesmo desempenho da técnica D256Lib. Mesmo com o ajuste de parâmetro automático não houve mudança nos resultados.

Na Figura 4.4 podemos visualizar os diagramas de caixa representando o desempenho da técnica SMOOpC em todas as imagens do banco de imagens descrito na Seção 3.3. Podemos observar que existe uma variação no desempenho em cada imagem.

Na Figura 4.5 podemos observar a imagem que a técnica SMOOpC apresentou o maior desempenho, nesta imagem existem poucas leveduras inviáveis. Já a Figura 4.6 é o exemplo que apresenta o menor desempenho. A principal diferença entre estas imagens é o número de leveduras inviáveis, isto porque, a imagem com o pior resultado apresenta um maior número de leveduras

²<http://www.csie.ntu.edu.tw/~cjlin/libsvm/>

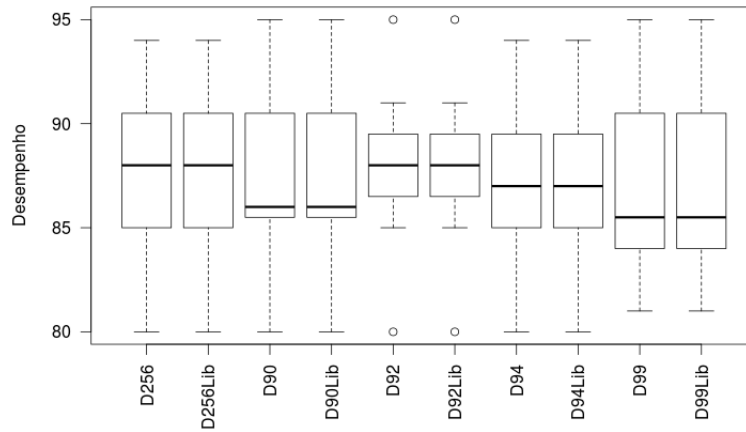


Figura 4.3: Diagrama de caixa das variações do SMOOpC com ajuste automático e manual.

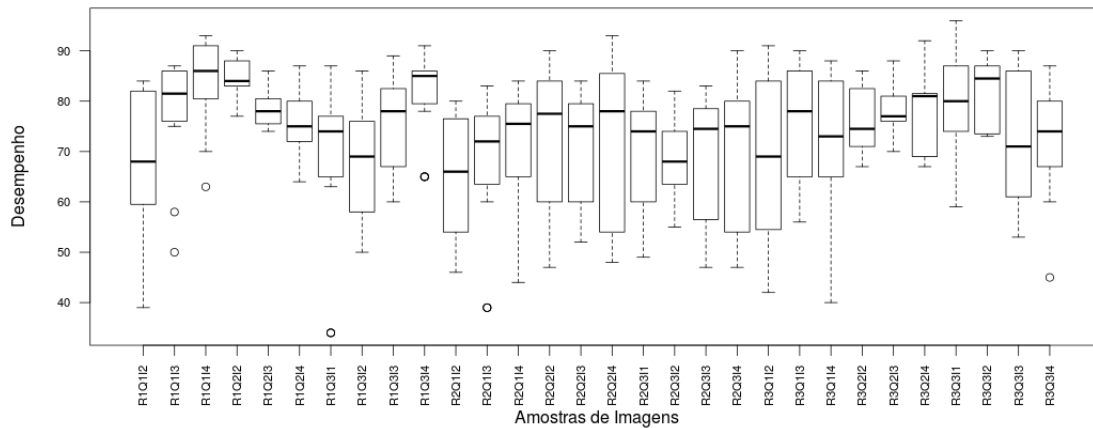


Figura 4.4: Diagrama de caixa das imagens, onde cada imagem está representada pela repetição, quadrante e a Imagem. Assim a Imagem R1Q1I2 representa a Repetição 1, Quadrante 1 e Imagem 2.

inviáveis, o que está causando uma confusão na identificação entre leveduras inviáveis e viáveis. Existindo também uma diferença de nitidez entre as imagens e o classificador identificou melhor as leveduras viáveis.

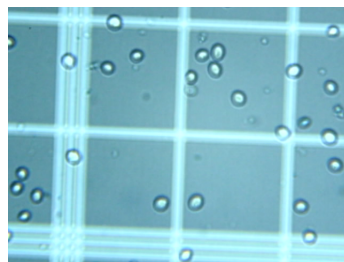


Figura 4.5: A imagem R1Q1I4 que apresentou maior desempenho com a técnica SMOOpC com dicionário com o tamanho 256.

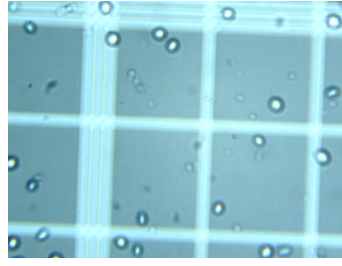


Figura 4.6: A imagem R2Q1I2 que apresentou o menor desempenho com a técnica SMOOpC com dicionário com o tamanho 256.

4.3 Discussão

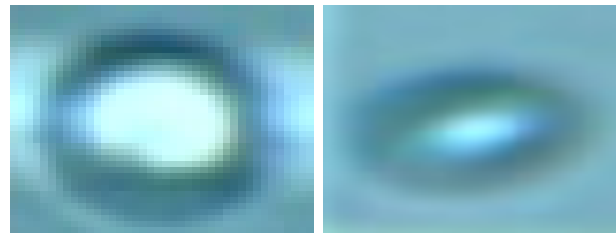
Utilizamos os extratores de atributos do Capítulo 2 e as técnicas de aprendizagem supervisionada do Capítulo 3, com o banco de imagens descrito na Seção 3.3 e o resultado obtido foi descrito na Seção 4.1. Com isto, um novo teste com ajuste de tamanho de dicionário foi realizado, com a técnica SMOOpC que apresentou o melhor resultado.

Os resultados com os ajustes de tamanho do dicionário mostraram que a técnica SMOOpC com o dicionário no tamanho 256 apresentou o maior desempenho na identificação das imagens de leveduras incolores. A matriz de confusão extraída do conjunto de treino na amostra (repetição) 1 e quadrante 2 no Brix 3, com a melhor técnica SMOOpC é apresentada na Tabela 4.2, onde as imagens de leveduras viáveis e imagens de fundo apresentaram o melhor desempenho. Ao todo, são 82 imagens identificadas como fundo e 16 imagens identificadas como leveduras viáveis.

a	b	c	<- classe
82	1	12	a = fundo
3	2	1	b = inviavel
4	0	16	c = viavel

Tabela 4.2: Matriz de confusão.

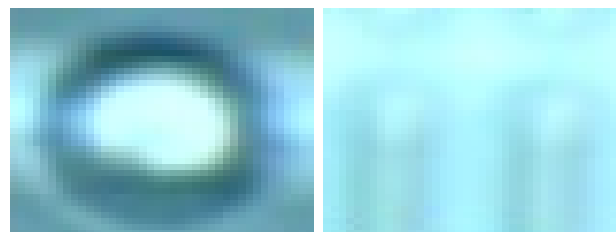
Nas Figuras 4.7a e 4.7b podemos visualizar duas leveduras que foram classificadas como viáveis, porém como podemos observar, a identificação está incorreta. Este é um dos problemas encontrados, onde as duas imagens foram confundidas, pois ambas leveduras apresentam a mesma forma, porém a cor da região central de cada levedura é o fator determinante para classificá-las. Como não temos um controle sobre as regiões detectadas pelo algoritmo *Opponent Color*, então os pontos de interesse podem ser capturados em qualquer região da imagem. Isto significa que não temos a informação espacial, e este é um dos principais problemas que encontramos em histogramas.



(a) Levedura Viável. (b) Levedura Inviável.

Figura 4.7: Ambas imagens classificadas como viáveis.

Nas Figuras 4.8a e 4.8b podemos visualizar duas imagens que foram classificadas como leveduras viáveis, porém a Figura 4.8b corresponde ao fundo. Este é um exemplo, onde o classificador confunde, pois a cor do fundo se contrasta com a cor do centro da levedura viável. Este é um exemplo, que a forma é um fator que discrimina melhor a levedura viável do fundo da imagem. Como estamos trabalhando com o algoritmo *Opponent Color*, procuramos por mudanças locais em toda imagem, deixando os atributos referentes a forma, que é uma característica importante na identificação da imagem.



(a) Levedura Viável. (b) Imagem de Fundo.

Figura 4.8: Ambas imagens classificadas como viáveis.

Os resultados demonstraram que a informação de cor adicionada ao algoritmo BoVW, melhora os resultados na identificação das leveduras. Mesmo existindo alguns problemas na identificação, o algoritmo *Opponent Color* com o classificador SMO apresentou o melhor desempenho com um dicionário de tamanho 256. Em relação ao trabalho de Quinta et al. (2010) nossa técnica representa invariante a rotação da imagem e com um desempenho acima de 85% enquanto o artigo apresentou o desempenho de 80%. Já em relação ao trabalho de Mongelo et al. (2012), nossa técnica conseguiu identificar melhor as leveduras viáveis.

Software BioViC

Técnicas para localizar as leveduras são apresentadas na Seção 5.1. Na Seção 5.2 é realizada a validação das técnicas utilizadas e a implementação do software BioViC é descrita na Seção 5.3.

5.1 *Localização de leveduras*

As imagens retiradas de cada amostra de leveduras podem conter várias leveduras. É necessário localizar as possíveis leveduras para aplicar a identificação através da técnica SMOOpC. Foram analisadas algumas técnicas de visão computacional para localizar as leveduras e que podemos destacar:

- **Detecção de Círculos:** As leveduras são visualizadas em formas ovais, através das imagens obtidas pelo microscópio. E em visão computacional, existem alguns algoritmos capazes de identificar formas geométricas, como é o caso da técnica da Transformada de Hough. Esta técnica é capaz de detectar linhas, círculos e elipses representados através de pixels em imagens (Duda and Hart, 1972). Podemos visualizar a detecção de círculos através da Figura 5.1.
- **Detecção de Contornos:** Outra forma de localizar as leveduras é através do contorno das próprias leveduras, dado que o contorno é uma fronteira que separa duas regiões em uma imagem. Existem algumas técnicas que extraem contornos através da binarização da imagem, ou seja, cada pixel que compõe uma imagem possui um valor associado a intensidade. Este valor é transformado em 0 ou 1, sendo que o valor 0 representa o fundo da imagem e o valor 1 o objeto de interesse, neste caso, o contorno(Suzuki

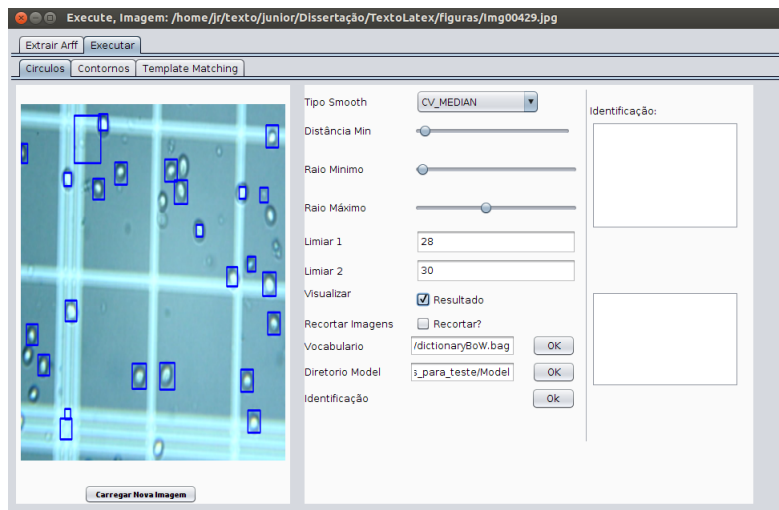


Figura 5.1: Imagem da tela responsável por detectar leveduras através da técnica de detecção de círculos.

et al., 1985). Na Figura 5.2 podemos visualizar o resultado da técnica de detecção de contornos.

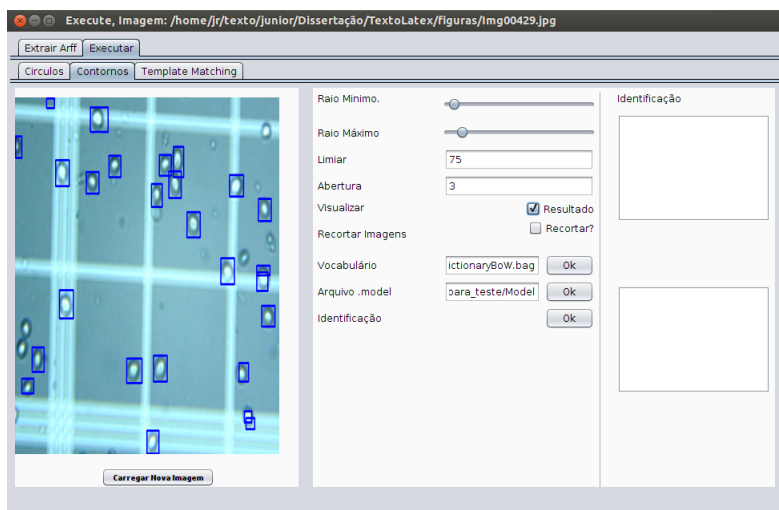


Figura 5.2: Imagem da tela responsável por detectar leveduras através da técnica de detecção de contornos.

- **Casamento de Modelos (*Template Matching*):** Esta é uma técnica em visão computacional que é responsável por efetuar a subtração de uma imagem por uma sub-imagem (*template*), através do deslocamento da sub-imagem na imagem. O valor da subtração representa a distância mínima entre a sub-imagem e a região em que a sub-imagem está sendo subtraída na imagem. A distância mínima poderá indicar uma possível combinação entre duas imagens, ou seja, a sub-imagem poderá estar presente ou não na imagem (Lewis, 1995). Esta técnica foi aplicada através de recortes de leveduras escolhidos pelo usuário e que são utilizados como sub-imagens (templates) em imagens obtidas para a contagem de

leveduras.

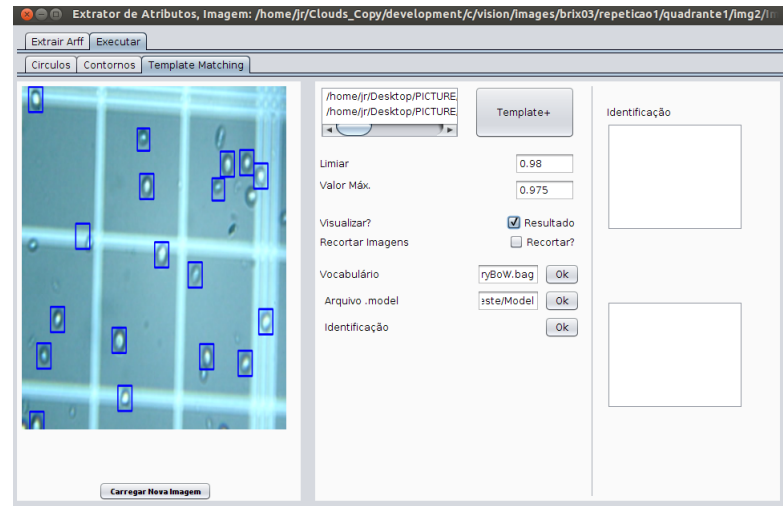
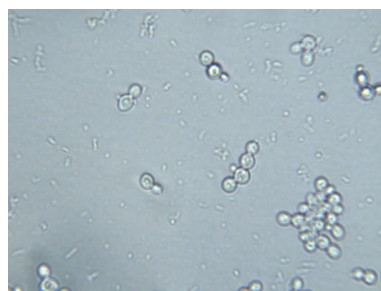


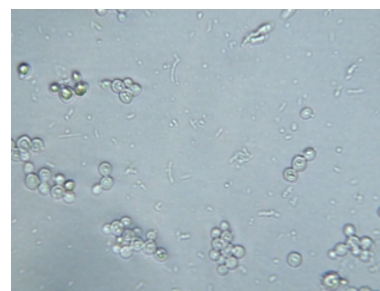
Figura 5.3: Imagem da tela responsável por detectar leveduras através da técnica *Template Matching*.

5.2 Validação do Software

Uma avaliação foi realizada para a validação do software BioViC. Para isto, as técnicas utilizadas foram: detecção de contornos, detecção de círculos, *template matching* (casamento de modelos - c.modelos) e manual. Estas técnicas foram utilizadas para localizar as possíveis leveduras para efetuar a identificação e a contagem através da técnica SMOOpC. Foram utilizadas 10 imagens com leveduras sem a câmera de Neubauer. Na Figura 5.4 podemos visualizar dois exemplos de imagens utilizadas na validação das técnicas. Na Figura 5.5 podemos visualizar todas as técnicas implementadas, cada imagem representa uma técnica e o resultado visual da contagem.



(a) imagem com leveduras



(b) imagem com leveduras.

Figura 5.4: Imagens com leveduras sem a câmera de Neubauer.

Foi realizada a contagem de leveduras viáveis em todas as 10 imagens e na Figura 5.6 podemos visualizar o resultado da contagem de leveduras viáveis. A técnica detecção de contornos apresentou uma contagem maior de leveduras viáveis em relação as outras técnicas.

Na Figura 5.7 podemos visualizar o resultado da contagem de leveduras inviáveis. A técnica detecção de contornos apresentou uma contagem maior de leveduras inviáveis em relação as outras técnicas. A técnica *template matching* apresenta os resultados mais parecidos com a contagem manual.

Para avaliarmos se existe semelhança estatística entre as técnicas, realizamos o teste de Tukey. Na Tabela 5.1 podemos visualizar o teste de Tukey entre as técnicas. Com um nível de confiança de 5%, as técnicas detecção de círculos e casamento de modelos pareceram estatisticamente iguais com a técnica manual realizada por um especialista que identificou e contou as leveduras viáveis. Nesta avaliação os pares de técnicas foram analisados através da quantidade de leveduras contadas a partir de cada técnica. Na técnica manual uma pessoa utilizou o software BioVic para marcar cada levedura identificada na imagem.

	círculos	contornos	c.modelos	manual
círculos	1	0	0	0
contornos	0.0031319	1	0	0
c.modelos	0.3228003	0.1638671	1	0
manual	0.8049983	0.0283374	0.8341115	1

Tabela 5.1: Tabela com os valores p-value obtidos do teste Tukey das técnicas utilizadas no software.

O software BioViC apresentou uma contagem de leveduras viáveis igual a contagem manual realizada por uma especialista segundo o teste de Tukey realizado.

5.3 Implementação do software BioViC

O software BioViC foi desenvolvido para ser uma ferramenta computacional que auxilia o controle microbiano. O projeto de desenvolvimento foi baseado no modelo SCRUM, que é uma metodologia ágil de desenvolvimento de software. Com o SCRUM metas e reuniões semanais são definidas, onde cada reunião tem um limite de tempo de no máximo 20 minutos e as atividades realizadas são apresentadas e discutidas (Sommerville et al., 2003). Esta metodologia permite o desenvolvimento rápido de software, apesar de não exigir uma documentação complexa e carregada de exigências e análises. O Software BioViC disponibiliza uma documentação gerada a partir do javadoc, que poderá ser utilizada para o esclarecimentos das funcionalidades implementadas.

A linguagem utilizada para desenvolver o BioViC foi o Java. Esta lingua-

gem é multiplataforma, desenvolvida para **Sun-Microsystems**¹ e vendida para a empresa **Oracle**². Por utilizar uma máquina virtual também conhecida como JVM, permite que todos os programas sejam executados em diversas plataformas e inclusive nos smartphones (Deitel et al., 2001). Assim, o BioViC pode ser executado nos sistemas operacionais: **Linux, Mac e Windows**.

Algumas extensões ou bibliotecas foram utilizadas para implementar a funcionalidade principal do BioViC que é a contagem e a identificação de leveduras em imagens obtidas de microscópio. As extensões utilizadas foram: a biblioteca OpenCV em Java e o módulo WEKA. A biblioteca OpenCV é responsável por disponibilizar técnicas implementadas de visão computacional, assim como, oferecer meios para que o usuário implemente novas técnicas. O módulo Weka oferece técnicas implementadas de aprendizagem automática para que o usuário, possa testar e utilizar na implementação de aplicações.

As funcionalidades do BioViC se destacam em:

- **Cadastrar Usuário:** esta funcionalidade permite que o usuário cadastre um novo usuário no sistema. Com o cadastro, o usuário poderá executar todas as outras funcionalidades. Podemos visualizar a tela de cadastro de usuário na Figura 5.8.
- **Cadastrar Usina:** esta funcionalidade permite que o usuário do sistema cadastre uma nova usina. Isto significa que ao gerar um relatório, o usuário terá disponível a informação da usina. A Figura 5.9 representa a tela de cadastro de usina.
- **Alterar Cadastro de usuário:** esta funcionalidade permite que o usuário do sistema altere o cadastro de outro usuário. Esta opção é utilizada caso seja necessário alterar alguma informação.
- **Alterar Cadastro de Usina:** esta funcionalidade permite que o usuário altere o cadastro da usina. Esta opção é utilizada caso seja necessário alterar alguma informação da usina.
- **Contar Leveduras:** esta funcionalidade permite que o usuário execute o módulo de contagem e identificação automática de leveduras. Este módulo é responsável por localizar e identificar as leveduras, além de permitir que um relatório seja salvo. Na Figura 5.10 podemos visualizar a tela que representa a funcionalidade de contar e identificar as leveduras.
- **Emitir Relatório:** esta funcionalidade é responsável por emitir um relatório, contendo informações da contagem e identificação de leveduras.

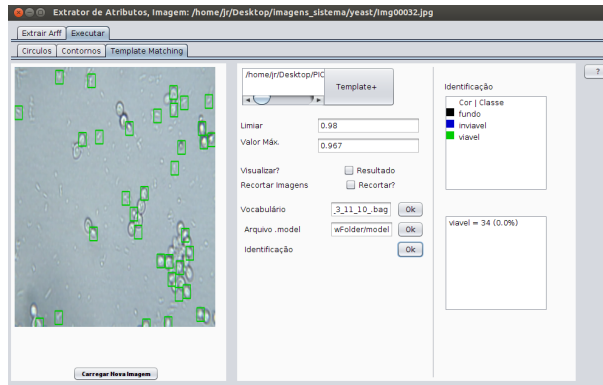
¹<http://www.oracle.com/us/sun/index.html>

²<http://www.oracle.com/index.html>

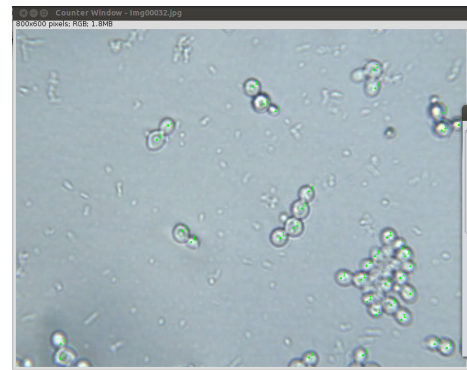
Neste relatório, as informações como: nome de usuário, nome da usina, quantidade de leveduras, data, número da amostra e etc. Na Figura 5.11 podemos visualizar a tela responsável por emitir relatório.

- Help: esta funcionalidade permite que o usuário tenha uma visão das funcionalidades do sistema.

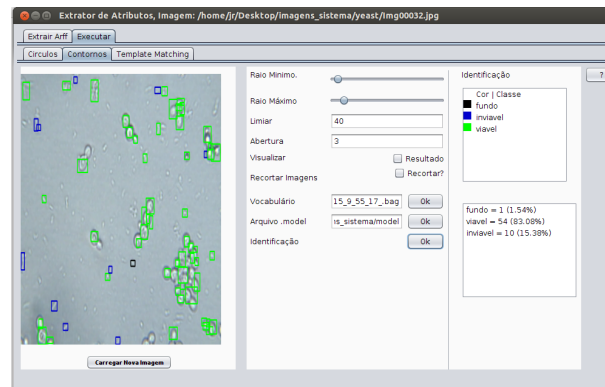
A câmera de Neubauer auxilia a contagem de leveduras de maneira visual. Este objeto possui custo razoável e está suscetível a danos através do uso incorreto. Verificamos que o software BioViC não precisa da câmera de Neubauer para facilitar a contagem e identificação, o que reduz custo e permite a contagem e identificação seja realizada mais rapidamente. Na Figura 5.12 podemos visualizar o resultado da identificação e contagem realizada sem a câmera de Neubauer.



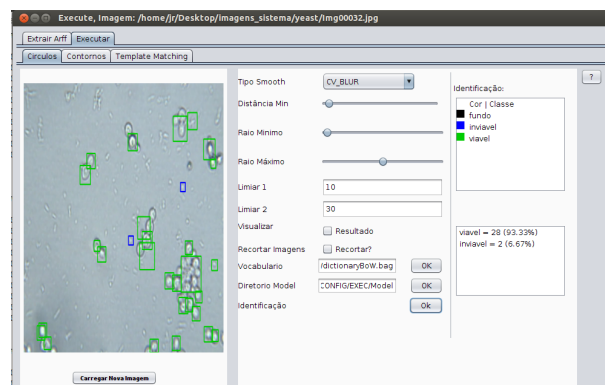
(a) Técnica Template Matching.



(b) Contagem Manual.



(c) Técnica Contornos



(d) Técnica Círculos.

Figura 5.5: Imagens com o resultado das principais técnicas utilizadas para localizar as leveduras.

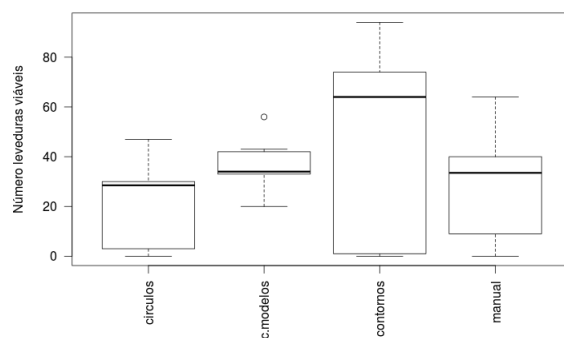


Figura 5.6: Imagem com o resultado da identificação e contagem de leveduras viáveis.

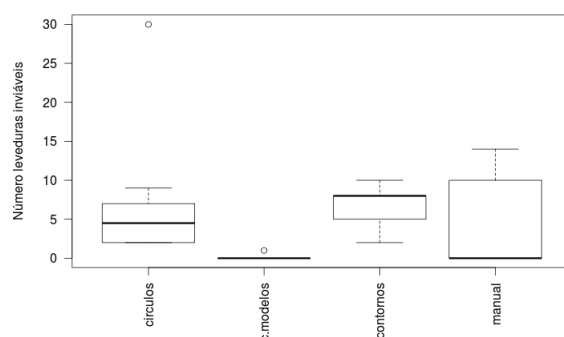


Figura 5.7: Imagem com o resultado da identificação e contagem de leveduras inviáveis.

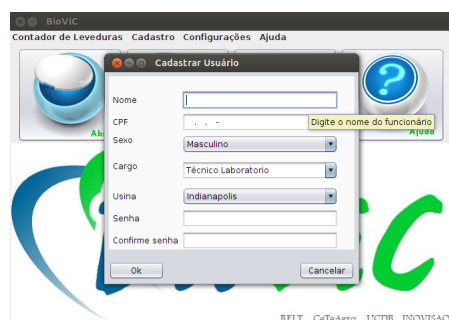


Figura 5.8: Imagem da tela responsável por cadastrar usuário no sistema.

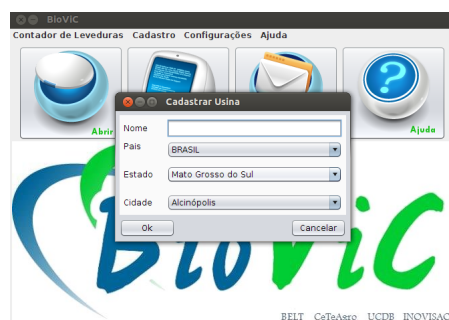


Figura 5.9: Imagem da tela responsável por cadastrar usina no sistema.



Figura 5.10: Imagem da tela responsável por contar as leveduras.

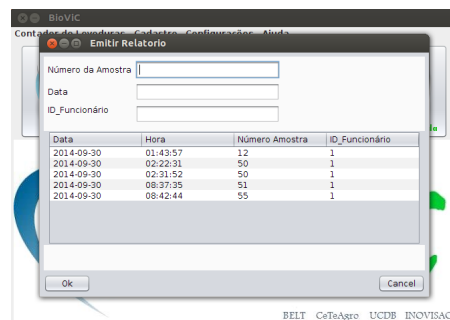


Figura 5.11: Imagem da tela responsável por emitir relatório.

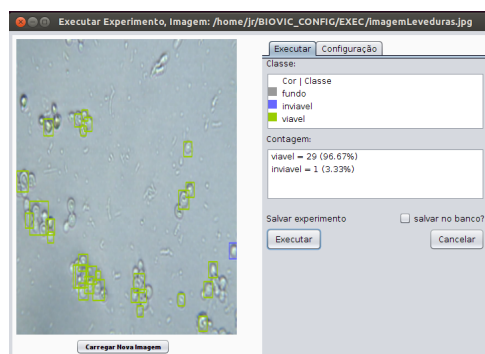


Figura 5.12: Imagem com o resultado da identificação e contagem de leveduras.

Considerações Finais

Para garantir a qualidade na produção do etanol, as atividades de contagem e classificação de leveduras viáveis e inviáveis através do corante azul de metileno são cruciais. Um modo de automatizar este processo é utilizar a visão computacional. Com isto, neste trabalho analisamos o algoritmo BoVW em conjunto de algumas técnicas que capturam a informação referente a cor, para extrair atributos que foram utilizados em classificadores. Os resultados obtidos demonstraram que a técnica de extração de atributos *Opponent Color* com o classificador SMO obteve os melhores resultados. A métrica utilizada foi a porcentagem de classificação correta. Com os testes de hipótese ANOVA e Friedman e os valores-p respectivos $2e^{-16}$ e $2.2e^{-16}$, verificamos a equivalência estatística entre as técnicas utilizadas. Porém com o pós-teste de Tukey, foi abordada uma equivalência da melhor técnica (SMOOpC) com outras técnicas. Com a matriz de confusão verificamos que a técnica SMOOpC com o dicionário de tamanho 256 identificou melhor as leveduras viáveis e o fundo, mesmo apresentando erros na identificação das leveduras inviáveis, a técnica SMOOpC apresenta um desempenho melhor em relação as outras técnicas analisadas. Esta técnica foi implementada no Software BioViC na etapa de identificação. Como as imagens possuem várias leveduras, foi avaliado técnicas de visão computacional para localizar as leveduras e então identificá-las. A técnica detecção de contornos e detecção de círculos apresentou o desempenho igual a contagem e identificação realizada manualmente. O software BioViC apresentou uma contagem de leveduras viáveis iguais com a contagem realizada por um especialista e o desempenho da técnica SMOOpC foi acima de 85% enquanto que Quinta et al. (2010) apresentam o desempenho de 80%.

6.1 *Trabalhos Futuros*

A validação do Software BioVic foi realizada de forma exploratória e não sistemática apenas para se ter uma ideia inicial de seu desempenho completo pois o foco deste trabalho foi na comparação das técnicas de extração de atributos e de classificação e não no desempenho final do software. É necessária agora a realização de experimentos melhores delineados e envolvendo uma quantidade maior de aspectos a serem analisados e um base de usuários mais representativa.

Referências Bibliográficas

- Alhwarin, F., Wang, C., Ristic-Durrant, D., e Gräser, A. (2008). Improved sift-features matching for object recognition. In *BCS Int. Acad. Conf.*, páginas 178–190. Citado na página 8.
- Bahri, A. e Zouaki, H. (2013). A surf-color moments for images retrieval based on bag-of- features. *European Journal of Computer Science and Information Technology*, 1(11-22). Citado nas páginas 4, 14, e 15.
- Bay, H., Ess, A., Tuytelaars, T., e Van Gool, L. (2008). Speeded-up robust features (surf). *Computer vision and image understanding*, 110(3):346–359. Citado na página 8.
- Beltrame, Walber Antônio Ramos, F. F. C. S. (2010). Aplicações práticas dos algoritmos de clusterização kmeans e bisecting k-means. *Departamento de Informática–Universidade Federal do Espírito Santo (UFES) Av. Fernando Ferrari*. Citado na página 9.
- Boinot (1939). Melle process of alcoholic fermentation with re-use os the yeass. *The international Sugar Journal*. Citado na página 1.
- Botterill, T., Mills, S., e Green, R. (2008). Speeded-up bag-of-words algorithm for robot localisation through scene recognition. In *Image and Vision Computing New Zealand, 2008. IVCNZ 2008. 23rd International Conference*, páginas 1–6. IEEE. Citado na página 4.
- Ceccato-Antonini, S. R. (2011). Microbiologia da fermentação alcoólica: a importância do monitoramento microbiológico em destilarias. pagina 120. Citado na página 2.
- Coelho, L. P., Shariff, A., e Murphy, R. F. (2009). Nuclear segmentation in microscope cell images: a hand-segmented dataset and comparison of algorithms. In *Biomedical Imaging: From Nano to Macro, 2009. ISBI'09. IEEE International Symposium*, páginas 518–521. IEEE. Citado na página 3.

- Csurka, G., Dance, C., Fan, L., Willamowski, J., e Bray, C. (2004). Visual categorization with bags of keypoints. In *Workshop on statistical learning in computer vision, ECCV*, volume 1, páginas 1–2. Citado nas páginas 4 e 8.
- Deitel, H. M., Deitel, P. J., e Santry, S. E. (2001). *Advanced Java 2 Platform How to Program*. Prentice Hall PTR, Upper Saddle River, NJ, USA, 1st edition. Citado na página 35.
- Duda, R. . e Hart, P. E. (1972). Use of the hough transformation to detect lines and curves in pictures. *Commun. ACM*, 15(1):11–15. Citado na página 31.
- Friedman, M. (1937). The use of ranks to avoid the assumption of normality implicit in the analysis of variance. *Journal of the American Statistical Association*, 32(200):675–701. Citado na página 23.
- Kohavi, R. (1996). Scaling up the accuracy of naive-bayes classifiers: A decision-tree hybrid. In *KDD*, páginas 202–207. Citado na página 5.
- Larson, M. G. (2008). Analysis of variance. *Circulation*, 117(1):115–121. Citado na página 23.
- Lewis, J. P. (1995). Fast template matching. In *Vision interface*, volume 95, páginas 15–19. Citado na página 32.
- Madzarov, G., Gjorgjevikj, D., e Chorbev, I. (2009). A multi-class svm classifier utilizing binary decision tree. *Informatica*, 33(2). Citado na página 5.
- Mongelo, A. I., da Silva, D. S., Quinta, L. N. B., Pistori, H., e Cereda, M. P. (2011). Validação de método baseado em visão computacional para automação de contagem de viabilidade de leveduras em indústrias alcooleiras. páginas 17–21 Outubro, Bento Gonçalves, RS. Citado na página 3.
- Murthy, S. K. (1998). Automatic construction of decision trees from data: A multi-disciplinary survey. *Data mining and knowledge discovery*, 2(4):345–389. Citado na página 5.
- Pass, G., Zabih, R., e Miller, J. (1997). Comparing images using color coherence vectors. In *Proceedings of the fourth ACM international conference on Multimedia*, páginas 65–73. ACM. Citado nas páginas 4, 11, e 13.
- Platt, J. C. (1999). Fast training of support vector machines using sequential minimal optimization. In *Advances in kernel methods*, páginas 185–208. MIT press. Citado na página 18.

- Quinta, L. N., Queiroz, J. H., Souza, K. P., Pistori, H., e Cereda, M. P. (2010). Classificação de leveduras para o controle microbiano em processos de produção de etanol. (4-7 Julho, Presidente Prudente - São Paulo). Citado na página 1.
- Salama, G. I., Abdelhalim, M., e Zeid, M. A.-e. (2012). Breast cancer diagnosis on three different datasets using multi-classifiers. *Breast Cancer (WDBC)*, 32(569):2. Citado nas páginas 5, 18, e 19.
- Schier, Jan Kovár, B. (2011). Automated counting of yeast colonies using the fast radial transform algorithm. In *BIOINFORMATICS*, páginas 22–27. Citado na página 3.
- Silva, D. S., Quinta, L. N. B., Mongelo, A. I., Cereda, M. P., e Pistori, H. (2012). Classificação de leveduras utilizando transformada de hough e aprendizagem supervisionada. In: *WVC 2012 - Workshop de Visão Computacional*, (2):27–30 Maio, Goiânia–Goiás. Citado na página 2.
- Sommerville, I., Melnikoff, S. S. S., Arakaki, R., e de Andrade Barbosa, E. (2003). *Engenharia de software*, volume 6. Addison Wesley São Paulo. Citado na página 34.
- Stratford, M. (1996). Yeast flocculation: restructuring the theories in line with recent research. *Cerevisia (Belgium)*. Citado na página 1.
- Suzuki, S. et al. (1985). Topological structural analysis of digitized binary images by border following. *Computer Vision, Graphics, and Image Processing*, 30(1):32–46. Citado na página 31.
- Valgren, C. e Lilienthal, A. J. (2007). Sift, surf and seasons: Long-term outdoor localization using local features. In *EMCR*. Citado na página 8.
- van de Sande, K. E., Gevers, T., e Snoek, C. G. (2008). Color descriptors for object category recognition. In *Conference on Colour in Graphics, Imaging, and Vision*, volume 2008, páginas 378–381. Society for Imaging Science and Technology. Citado nas páginas 4 e 16.
- Van Weijer, J. e Khan, F. S. (2013). Fusing color and shape for bag-of-words based object recognition. In *Computational Color Imaging*, páginas 25–34. Springer. Citado na página 4.
- Wengert, C., Douze, M., e Jégou, H. (2011). Bag-of-colors for improved image search. In *Proceedings of the 19th ACM international conference on Multimedia*, páginas 1437–1440. ACM. Citado nas páginas 4 e 9.

Witten, I. H. e Frank, E. (2005). *Data Mining: Practical machine learning tools and techniques*. Morgan Kaufmann. Citado nas páginas 17 e 19.