



UNIVERSIDADE FEDERAL DE OURO PRETO
ESCOLA DE MINAS
COLEGIADO DO CURSO DE ENGENHARIA DE CONTROLE
E AUTOMAÇÃO - CEC AU



MATHEUS HENRIQUE LEMES FARIA

ABORDAGEM MULTI OBJETIVO PARA A SELEÇÃO DE *FEATURES*
EM REPRESENTAÇÕES PROFUNDAS DA ÍRIS

MONOGRAFIA DE GRADUAÇÃO EM ENGENHARIA DE CONTROLE E
AUTOMAÇÃO

Ouro Preto, 2021

MATHEUS HENRIQUE LEMES FARIA

**ABORDAGEM MULTIOBJETIVO PARA A SELEÇÃO DE *FEATURES*
EM REPRESENTAÇÕES PROFUNDAS DA ÍRIS**

Monografia apresentada ao Curso de Engenharia de Controle e Automação da Universidade Federal de Ouro Preto como parte dos requisitos para a obtenção do Grau de Engenheiro de Controle e Automação.

Orientador: Prof. Pedro Henrique Lopes Silva, Me.

Coorientadora: Profa. Adrielle de Carvalho Santana, Dr.Sc.

**Ouro Preto
Escola de Minas – UFOP
2021**

SISBIN - SISTEMA DE BIBLIOTECAS E INFORMAÇÃO

F224a Faria, Matheus Henrique Lemes.

Abordagem multiobjetivo para a seleção de features em representações profundas da íris. [manuscrito] / Matheus Henrique Lemes Faria. - 2022.

45 f.: il.: color., gráf., tab..

Orientador: Prof. Me. Pedro Lopes.

Coorientadora: Profa. Dra. Adrielle Santana.

Monografia (Bacharelado). Universidade Federal de Ouro Preto. Escola de Minas. Graduação em Engenharia de Controle e Automação .

1. Otimização multiobjetivo. 2. Algoritmos genéticos. 3. Biometria. I. Lopes, Pedro. II. Santana, Adrielle. III. Universidade Federal de Ouro Preto. IV. Título.

CDU 57.087.1

Bibliotecário(a) Responsável: Angela Maria Raimundo - SIAPE: 1.644.803



FOLHA DE APROVAÇÃO

Matheus Henrique Lemes Faria

Abordagem Multiobjetivo para a Seleção de Features em Representações Profundas da Íris

Monografia apresentada ao Curso de Engenharia de Controle e Automação da Universidade Federal de Ouro Preto como requisito parcial para obtenção do título de Engenheiro de Controle e Automação

Aprovada em 07 de janeiro de 2022

Membros da banca

Prof. Me. Pedro Henrique Lopes Silva - Orientador

Profa. Dra. Adrielle de Carvalho Santana – Coorientadora

Prof. Dr. Gladston Juliano Prates Moreira - Professor Convidado

Prof. Dr. Eduardo Jose da Silva Luz - Professor Convidado

Pedro Henrique Lopes Silva, orientador do trabalho, aprovou a versão final e autorizou seu depósito na Biblioteca Digital de Trabalhos de Conclusão de Curso da UFOP em 20/01/2022



Documento assinado eletronicamente por **Pedro Henrique Lopes Silva, PROFESSOR DE MAGISTERIO SUPERIOR**, em 20/01/2022, às 19:06, conforme horário oficial de Brasília, com fundamento no art. 6º, § 1º, do [Decreto nº 8.539, de 8 de outubro de 2015](#).



A autenticidade deste documento pode ser conferida no site http://sei.ufop.br/sei/controlador_externo.php?acao=documento_conferir&id_orgao_acesso_externo=0, informando o código verificador **0264866** e o código CRC **636F8196**.

Este trabalho é dedicado primeiramente a Deus, pois sem Ele nada é possível, aos meus pais, namorada e irmão, por me apoiarem e me guiarem até mais essa conquista

AGRADECIMENTOS

Agradeço ...

À Deus pela minha vida, pelas bênçãos, fé, saúde, família e amigos proporcionados à mim, e por iluminar mais uma conquista na minha vida.

Aos meus pais, Helmar e Maricelma, que apesar dos problemas sempre me fortaleceram e buscaram sempre o melhor pra mim. Ao meu irmão.

Arthur pelo apoio e amizade verdadeira ao longo desse caminho.

À minha namorada e companheira de vida, Maria Fernanda, por aguentar minhas lamúrias, escutando meus problemas e sempre estar ao meu lado me incentivando e apoiando nos momentos em que pensei abandonar tudo e desistir.

Ao Zenóbio, pela oportunidade de bolsa no Arquivo Central da UFOP, pois sem essa oportunidade dificilmente conseguiria me manter na universidade.

Ao meu orientador Pedro Silva, e coorientadora, Adrielle Santana, pela dedicação à essa monografia e ensinamentos passados a mim.

“A mente que se abre a uma nova ideia jamais voltará ao seu tamanho original.”
(Albert Einstein)

RESUMO

Sistemas Biométricos sofrem bastante com grande volume de informações e com a quantidade de dados que são utilizados para identificar ou verificar uma pessoa. Consequentemente a demanda por poder de processamento e tempo é maior. Com isso, é de fundamental importância selecionar somente as características mais relevantes para descrever a pessoa sem que o sistema perca seu poder de decisão. Portanto, para resolver esse problema nesse trabalho propõe-se uma otimização multi-objetivo com os algoritmos NSGA-II e R-NSGA-II sobre as características extraídas das imagens da competição NICE.II. O foco desta monografia é minimizar a quantidade de características selecionadas e maximizar a decidibilidade (métrica utilizada na literatura para medir o poder de decisão de um sistema) utilizando como métricas de avaliação da otimização o *Equal Error Rate* (EER) e o Hiper-volume. Ao final deste trabalho, obteve-se uma decidibilidade maior e uma redução na quantidade de características selecionadas para ambos os algoritmos genéticos quando comparado ao uso de todas as características. Quando se compara o desempenho dos dois algoritmos, o R-NSGA-II foi levemente superior ao NSGA-II devido sua natureza de possuir um ponto de referência fornecido pelo programador durante a otimização.

Palavras-chaves: Otimização. Multi-objetivo. Biometria. Íris.

ABSTRACT

Biometric Systems identified with a large volume of information and with the amount of data that are used for or a person. Consequently, the demand for processing power and time is greater. With this, it is of fundamental importance to select only the most relevant characteristics to describe the person without losing his decision power. Therefore, for this work project, a multi-objective optimization with resources II and R-NSGA- is solved on the characteristics of the NICE.II competition problems. The focus of this photograph is to minimize the amount of selected features and maximize the decision-making capacity in the literature to measure the decision-making power of a system) using the EqualError Rate(ER) and the Hipervolume as optimization evaluation measures. At the end of this work, a major decision and a reduction in the quality of selected traits was obtained for when both aspects of use when using all genetic traits compared to all traits. When comparing the performance of the two algorithms, R-NSGA-II was slightly superior to NSGA-II due to its nature of having a different reference point by the programmer during optimization.

Key-words: Optimization. Multi-objective. Biometry. Iris.

LISTA DE ILUSTRAÇÕES

Figura 1 – Demonstrativo das etapas dos três diferentes tipos de Sistemas Biométricos.	17
Figura 2 – Representação gráfica da fronteira-Pareto (linha preta) para os quatro diferentes tipos de otimização com dois objetivos.	21
Figura 3 – Fluxograma da execução de um Algoritmo Genético e suas características básicas.	22
Figura 4 – Representação gráfica do <i>crowding distance</i>	24
Figura 5 – Etapas de avaliação do NSGA-II; sendo a primeira, a partir da população são selecionados os indivíduos da fronteira-pareto por meio da ordenação de não dominância; segunda etapa, a partir dos indivíduos selecionados utiliza-se a ordenação de <i>crowding distance</i> como critério de desempate	25
Figura 6 – Efeito da variação de ϵ nas soluções	27
Figura 7 – Distribuições de genuínas e de impostores	28
Figura 8 – Gráfico de 2 curvas ROC, evidenciando o ponto do EER.	29
Figura 9 – Exemplo Hiper-volume.	30
Figura 10 – Etapas do desenvolvimento	32
Figura 11 – Exemplos de imagens da base de dados UBIRIS.V2	33
Figura 12 – Descritor usado para a extração de características da íris	34
Figura 13 – Gráfico da Seleção de Características Utilizando o PCA, eixo x são as características selecionadas e o eixo y é a inclinação da curva. De modo que o número de características em que a inclinação é próxima de zero, melhor descreve o conjunto de dados	38
Figura 14 – Evolução da fronteira-pareto ao longo das gerações utilizando o NSGA-II	39
Figura 15 – Hiper-volume da otimização ao longo das gerações utilizando o NSGA-II.	39
Figura 16 – Evolução da fronteira-pareto ao longo das gerações utilizando o R-NSGA-II	41
Figura 17 – Hiper-volume da otimização ao longo das gerações utilizando o R-NSGA-II.	41

LISTA DE TABELAS

Tabela 1 – Comparativo entre os principais tipos de biometria.	18
Tabela 2 – Pontos da Fronteira-pareto obtido com NSGA-II.	38
Tabela 3 – Pontos da Fronteira-pareto obtido com R-NSGA-II.	40

LISTA DE ABREVIATURAS E SIGLAS

CNN	Convolutional Neural Network
DCNN	Deep Convolutional Neural Network
DID	Deep Eye Descriptor
DSCNN	Deep Separable Convolutional Neural Network
EER	Equal Error Rate
FAR	False Acceptance Rate
FN	False Negatives
FP	False Positives
FPR	False PositiveRate
FRR	False Rejection Rate
GA	Generic Algorithm
HV	Hiper-volume
IA	Inteligência Artificial
JSON	JavaScript Object Notation
KNN	K-Nearest Neighbors
NICE.II	Noisy Iris Challenge Evaluation
NSGA-II	Non-domination Based Genetic Algorithm
PCA	Principal Component Analysis
PSO	Particle Swarm Optimization
PSC	Problema de Seleção de Características
RNSGA-II	Reference point Non-domination Based Genetic Algorithm
ROC	Receiver Operating Characteristic
SDM	Supervised Descent Method
SGD	Stochastic Gradient Descent

SPEA2	Strength Pareto Evolutionary Algorithm
TN	True Negatives
TP	True Positives
TPR	True Positive Rate

LISTA DE SÍMBOLOS

$\in$	Pertence
Σ	Somatório
$\prod$	Produtório

SUMÁRIO

1	INTRODUÇÃO	15
1.1	Formulação do Problema	15
1.2	Objetivos	19
1.3	Estrutura do trabalho	19
2	REFERENCIAL TEÓRICO	20
2.1	Abordagens Multi-objetivo	20
2.1.1	<i>NSGA-II</i>	21
2.1.2	<i>R-NSGA-II</i>	25
2.2	Métricas	27
2.2.1	<i>Decidibilidade</i>	27
2.2.2	<i>Equal Error Rate (EER)</i>	28
2.2.3	<i>Hiper-volume</i>	29
2.3	Trabalhos Relacionados	30
3	DESENVOLVIMENTO	32
3.1	Base de dados	33
3.2	Extração das Características	34
3.3	Problema da Seleção de Características	34
3.4	Metodologia proposta	35
4	EXPERIMENTOS E RESULTADOS	37
4.1	NSGA-II	37
4.2	R-NSGA-II	40
5	CONCLUSÃO	42
5.1	Trabalhos Futuros	42
	REFERÊNCIAS	43

1 INTRODUÇÃO

1.1 Formulação do Problema

Atualmente, é notável que a cada dia um número imensurável de dados são postados e circulam na Internet. São dados de ações negociadas na bolsa de valores, dados de empresas, pesquisas e artigos científicos, bem como vídeos, fotos postadas em redes sociais, notícias, entre outros. Quanto a sua origem, esses dados podem ser emitidos desde celulares, tablets, computadores, ou até mesmo de sistemas embarcados com sensores que coletam informações de equipamentos.

Essas informações são classificadas em públicas e pessoais, de acordo com [Greenberg, Boyle e Laberge \(1999\)](#). Os dados públicos são os artefatos cuja propriedade pertence a um grupo, podendo ser acessado e manipulado pelos indivíduos que compõe o mesmo. As fichas dos funcionários de uma empresa são exemplos de dados públicos pertencentes à um grupo (a empresa), outro exemplo mais abrangente é uma foto publicada em uma rede social em que o grupo de usuários da rede social podem ter acesso a esses dados. Já os dados classificados como pessoais, se diferem por terem sido criados e serem de propriedade de um só indivíduo e, portanto, devem ser acessados e manipulados por somente essa pessoa como, por exemplo, os dados e senhas bancárias. Contudo, no cenário atual da expansão da internet e dos dados que circulam nela, existem certas fragilidades quanto ao acesso restrito dos dados que deveriam ser pessoais.

Uma das principais metodologias de se restringir ou dificultar o acesso aos dados pessoais é a utilização de usuário e senha, mesmo que essa metodologia apresente diversas vulnerabilidades ([SILVA et al., 2018](#)). Comumente, as senhas dos usuários em questão, são um tanto quanto simplistas ou estão associadas à algo da rotina do indivíduo, de modo que por muitas vezes, por algum descuido, o usuário pode acabar compartilhando e divulgando as suas senhas em sites *fakes* ou maliciosos. Além disso, existem ainda métodos de *brute-force* (força-bruta) para acessar esses dados.

Visando resolver esses problemas pode-se utilizar técnicas biométricas, onde a modalidade biométrica ou biometria é reconhecida como um identificador biológico único de cada ser vivo ([SCHNEIER, 1999](#)). Analisando etimologicamente, a palavra biometria se origina do grego *bios* - vida, e *metricos* - medida. Ao analisar historicamente, nota-se que ao lado de pinturas rupestres foram encontradas marcas de mãos que possivelmente serviriam como a assinatura ou uma marca biométrica para identificar o artista; e na antiga Babilônia, 500 a.C. as transações comerciais eram firmadas utilizando marcas de impressões digitais ([DUARTE, 2010](#)).

Para que determinada parte do corpo humano, ou padrões comportamentais, possam ser utilizadas como biometria deve-se cumprir quatro requisitos básicos, que foram previamente

definidos por [Jain, Ross e Prabhakar \(2004\)](#):

- **Distinguível:** ser diferenciável entre os indivíduos.
- **Imutável:** não pode alterar seus padrões ao longo do tempo, ou alterar o mínimo possível.
- **Mensurável:** ser possível medir ou ser quantificada.
- **Universal:** deve ser presente em todos indivíduos do grupo.

Desse modo, pode-se destacar que por meio de técnicas computacionais e a utilização de modalidades biométricas é possível identificar e também reconhecer um indivíduo, o que possibilita restringir o acesso aos dados somente a pessoa aos quais eles pertencem. Neste cenário surgem os Sistemas Biométricos, os quais utilizam de partes do corpo para realização do reconhecimento, sendo a impressão digital ([VERMA; CHAKRABORTY, 2019](#)), íris ([WANG et al., 2012](#)) e a face ([JIANG; LEARNED-MILLER, 2017](#); [DING; TAO, 2015](#)) as mais comumente utilizadas. Ao reconhecer tais padrões torna-se possível diferenciar o indivíduo dos demais. Entretanto, um sistema biométrico não limita-se à somente características físicas pois pode-se utilizar também padrões comportamentais como, por exemplo, a forma de andar ou a união de ambas ([SAEED, 2012](#)).

Além de restringir e proteger o acesso aos dados pessoais, os Sistemas Biométricos podem ser aplicados desde o uso em sistemas de acesso de uma portaria em um condomínio particular, até mesmo seu uso no reconhecimento de indivíduos em aeroportos e fronteiras no combate ao terrorismo, por exemplo ([SILVA et al., 2018](#)).

[Prabhakar, Pankanti e Jain \(2003\)](#) definem que para o desenvolvimento de um Sistema Biométrico deve-se seguir cinco etapas. A primeira trata-se da aquisição dos dados da modalidade biométrica por meio de sensores específicos. Como exemplo, para modalidade da íris utiliza-se câmeras fotográficas e/ou câmeras de infravermelho. Na segunda etapa, deve-se extrair os padrões dos dados coletados no sensor por meio de técnicas computacionais de reconhecimento de padrões, resultando em um vetor de características ou *features*. Na terceira etapa, deve-se salvar esse vetor de características biométricas do indivíduo em um banco de dados para ser usado futuramente. A quarta etapa é a de classificação em que deve-se realizar um comparativo entre os vetores de *features* e as informações fornecidas pelo usuário. A quinta e última etapa é onde deve-se tomar a decisão de, baseado na classificação, aceitar ou rejeitar o indivíduo ou realizar a identificação do mesmo, dependendo do tipo o Sistema Biométrico.

Ainda em seu trabalho, [Prabhakar, Pankanti e Jain \(2003\)](#) classificam os Sistemas Biométricos em três diferentes tipos mostrados na Figura 1. O primeiro tipo: o de registro ou *enrollment* é a primeira interação do usuário com o sistema em que são coletados e salvos suas respectivas modalidades biométricas, esse tipo se difere por ser executado somente até a terceira etapa das citadas anteriormente. O segundo tipo, são os sistemas de verificação, onde

é fornecido o vetor de características e uma identidade para esse vetor então, a partir, disso o processo compara essas informações recebidas com outras características referente a identidade previamente salvas em um banco de dados. Um grande exemplo desse tipo sistema seria na portaria da condomínios, em que dado uma entrada biométrica (impressão digital) só é permitido entrar se o indivíduo for morador do mesmo. Outro excelente exemplo, são os caixas eletrônicos em que é fornecido o cartão (a identificação) e a biometria e então o sistema identifica o respectivo usuário. O último tipo é o de Identificação que consiste em receber uma entrada de um vetor de características e, baseado no banco de dados salvo, classificar / identificar quem é aquele indivíduo, evitando que o mesmo assuma várias identidades. Um exemplo de utilização desse tipo de sistema é o de reconhecimento de um possível foragido da justiça por meio da análise das faces encontradas nas imagens de câmeras de segurança de um aeroporto.

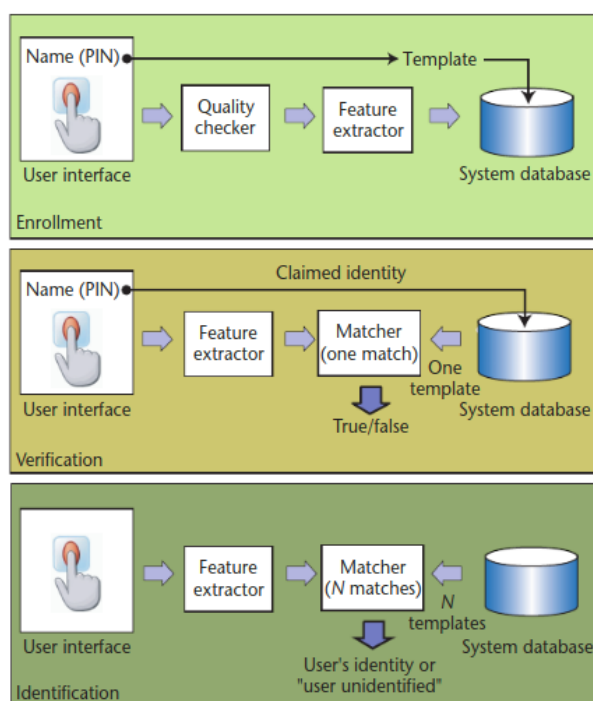


Figura 1 – Demonstrativo das etapas dos três diferentes tipos de Sistemas Biométricos.

Fonte: (PRABHAKAR; PANKANTI; JAIN, 2003)

Os Sistemas Biométricos ainda apresentam dois diferentes tipos de cenário: o de galeria aberta, ou *open-world*, o qual permite adicionar novos indivíduos ao sistema sem que haja a necessidade de realizar um novo treinamento do modelo (LUZ et al., 2017). O outro cenário é o de galeria fechada, ou *closed-world*, no qual existe a necessidade de retreinar o modelo ao adicionar um novo indivíduo (PROENÇA; NEVES, 2017).

Nota-se então que diferentes tipos de biometrias estão associadas à diferentes tipos de sistemas biométricos, visto que não existe uma modalidade biométrica universal e perfeita, cada uma possui seus pontos fortes e fracos (PRABHAKAR; PANKANTI; JAIN, 2003). A escolha

da modalidade biométrica depende muito da necessidade e da funcionalidade ao qual o sistema biométrico atenderá.

Ao analisar a Tabela 1, destaca-se como uma das mais promissoras a biometria da íris devido ao fato de possuir menos barreiras para sua universalidade do que seu principal concorrente, as impressões digitais, e por dificilmente sofrer alterações ao longo do tempo, pois possuem o humor aquoso e a córnea como proteção natural (MOAZED, 2020). Outro quesito que destaca-se é o fato da íris ser totalmente distinguível, visto que ela apresenta um fator de unicidade - parâmetro que determina quão único é aquele valor dentro de um conjunto de dados, na ordem de 10^{72} (MALHAS; RENFORTH, 2015).

Tabela 1 – Comparativo entre os principais tipos de modalidades biométricas.

Biometria	Impressão Digital	Face	Geometria da Mão	Íris	Voz
Barreiras para universalidade	Fissuras gastas; Comprometimento da mão ou do dedo	Nenhum	Comprometimento da mão	Comprometimento dos olhos	Comprometimento da fala
Distinção	Alto	Baixo	Médio	Alto	Baixo
Permanência	Alto	Médio	Médio	Alto	Baixo
Coletividade	Médio	Alto	Alto	Médio	Médio
Performance	Alto	Baixo	Médio	Alto	Baixo
Aceitabilidade	Médio	Alto	Médio	Baixo	Alto
Potencial de Evasão	Baixo	Alto	Médio	Baixo	Alto

Fonte: (PRABHAKAR; PANKANTI; JAIN, 2003) - Traduzido pelo autor

Além disso, a formação da íris não está totalmente relacionada com fatores genéticos e o ambiente em que o embrião se desenvolve. Isso permite que gêmeos univitelinos possuam íris diferentes (SUN; TAN, 2009), uma vez que o ambiente interfere nas características da íris até a conclusão da formação de sua proteção do humor aquoso (MOAZED, 2020) (cinco anos de idade aproximadamente). Devido a esses motivos, o presente trabalho deu foco para a modalidade biométrica da íris.

Estão cada vez mais presentes exemplos de sistemas que utilizam Inteligência Artificial (IA), aplicando métodos de *Deep Learning* em Sistemas Biométricos de identificação e verificação como, por exemplo, nos *smartphones* da Apple com detecção facial e na linha Galaxy da Samsung com detecção da íris, em que o sistema de login e senha está praticamente em desuso nesses celulares.

Na literatura encontra-se também diversos exemplos do uso da íris. Para a presente monografia vale destacar o artigo de Wang et al. (2012), o qual utilizou modelo de Adaboost (VAGHELA; GANATRA; THAKKAR, 2009; YIN; LIU; HAN, 2005) como classificador para o banco de dados com imagens de íris NICE: II. O trabalho de Silva et al. (2018) apresenta um novo estado da arte para detecção por meio da íris, o qual se utiliza de uma *Convolutional*

Neural Network para extração das características das imagens de íris, o qual apresenta, uma abordagem mono-objetivo para otimizar a taxa de reconhecimento ou decidibilidade. Entretanto, nesses trabalhos, o vetor de características/*features* extraído são enormes, com mais de duzentas características. Consequentemente o custo computacional para analisar todas essas *features* também é maior e também é necessário um espaço maior no banco de dados para comportar todas essas informações. Considerando esse cenário, o presente trabalho propõe selecionar somente as melhores características sem prejudicar a decidibilidade do sistema, adotando uma abordagem multi-objetivo visando otimização que maximize a decidibilidade enquanto minimiza a quantidade de *features*.

1.2 Objetivos

O principal objetivo dessa monografia é dar continuidade ao trabalho de [Silva et al. \(2018\)](#), no âmbito da detecção da íris, implementando sob uma perspectiva multi-objetivo com a aplicação de algoritmos como NSGA-II (*Non-domination Based Genetic Algorithm*) e o RNSGA-II (*Reference Point Non-domination Based Genetic Algorithm*) para maximizar a decidibilidade do sistema e minimizar a quantidade de *features*, de modo que selecione somente as características mais significativas sem reduzir o poder de decisão do sistema. Ao final, obteve-se um comparativo entre os dois algoritmos multi-objetivo, as *features* selecionadas pelo modelo e o impacto delas na decidibilidade.

Como objetivos específicos tem-se:

- Realizar uma revisão bibliográfica sobre Otimização Multi-objetivo.
- Realizar uma revisão bibliográfica sobre algoritmos de Otimização Multi-objetivo.
- Implementação do código do algoritmo em Python.
- Comparar os dois algoritmos propostos (NSGA-II e RNSGA-II).
- Escrever a monografia.

1.3 Estrutura do trabalho

A organização dessa monografia se inicia apresentando no Capítulo 2, a fundamentação teórica e a revisão da literatura com os elementos conceituais para o desenvolvimento do trabalho. No Capítulo 3 são descritos a metodologia e os experimentos realizados. Posteriormente, é apresentado o resultado do desenvolvimento e dos experimentos realizados no Capítulo 4. Por fim, a conclusão e trabalhos futuros são abordados no Capítulo 5.

2 REFERENCIAL TEÓRICO

Este capítulo contempla os conceitos relacionados com a presente monografia. Na Seção 2.1 são descritos conceitos e algoritmos de otimização multi-objetivo. Na Seção 2.2 são apresentadas as métricas que são utilizadas para avaliar as abordagens criadas e para comparação com trabalhos da literatura. Por fim, são apresentados os trabalhos relacionados com abordagens multi-objetivo para sistemas biométricos e trabalhos que utilizam o banco de dados NICE II para reconhecimento de íris, na Seção 2.3.

2.1 Abordagens Multi-objetivo

Os problemas reais existentes que estão associados geralmente à otimização de diversos objetivos os quais devem ser atingidos simultaneamente e na maioria das vezes são conflitantes (TICONA, 2008). Isso implica que não existe uma única solução que otimize e atenda perfeitamente à todas metas. Desse modo, para essa classe de problemas existe um conjunto com várias soluções eficientes.

Problemas com essas propriedades são chamados de problemas de otimização multi-objetivo, justamente por compreenderem a minimização ou maximização simultânea de um conjunto de objetivos, os quais podem estar sujeitos a um conjunto de restrições. A escolha da melhor ou mais viável solução dentre o conjunto é de responsabilidade do analista do problema, o qual deve ponderar dentre os objetivos globais (ARROYO, 2002).

Uma das formas de definir formalmente um problema de otimização multi-objetivo genérico é a que Deb (2001) propõe em seu livro. O problema pode ser definido como sendo:

$$\text{minimizar } f(x) = (f_1(x), f_2(x), \dots, f_M(x)), x \in R^N \quad (2.1)$$

$$\text{sujeito a } \begin{cases} h(x) = (h_1(x), h_2(x), \dots, h_t(x)) = a \\ g(x) = (g_1(x), g_2(x), \dots, g_u(x)) \leq b \end{cases} \quad (2.2)$$

onde $f(x)$ são as funções objetivo para o problema, $h(x)$ representa as restrições de igualdade e $g(x)$ representa as restrições de desigualdade do problema.

Na otimização multi-objetivo, o conjunto dos objetivos S representa o conjunto de soluções viáveis e factíveis do problema. São conhecidas também como soluções Pareto-ótimas. Esse conceito de Pareto-ótimo é definido por Pareto (1896), no qual um vetor de soluções s é Pareto-ótimo se não existir um outro vetor viável s^* que possa melhorar algum objetivo, sem piorar nenhum outro objetivo, ou seja, s é um Pareto-ótimo somente se s dominar s^* .

Para um problema de minimização, tem-se o conceito de dominância (PARETO, 1896):

- s domina s^* se, e somente se, $s_j \leq s_j^* \ \forall j$ e $s_j < s_j^*$ para algum j ;
- s e s^* são indiferentes ou possuem o mesmo grau de dominância se, e somente se, s não domina s^* e s^* não domina s .

Na Figura 2 é possível observar graficamente a fronteira-Pareto e como a mesma se altera de acordo com os diferentes tipos de objetivo do problema: (i) problemas em que se deseja maximizar a primeira função objetivo e minimizar a segunda função objetivo; (ii) minimizar a primeira função objetivo e maximizar a segunda função objetivo; (iii) minimizar as duas funções objetivo; e (iv) maximizar as duas funções objetivo. Para a presente monografia, devido as funções objetivos serem minimizar a quantidade de *features* selecionadas e maximizar a decidibilidade ao final da otimização a fronteira-Pareto deve-se respeitar o padrão Min-Max, como pode ser observado na Figura 2.

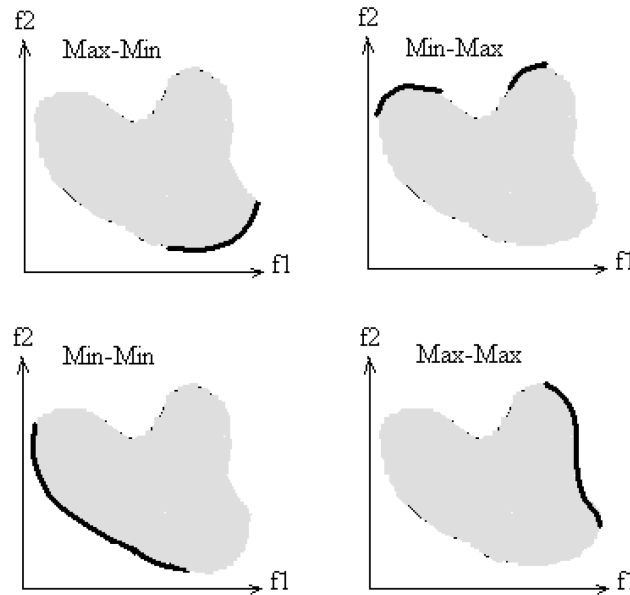


Figura 2 – Representação gráfica da fronteira-Pareto (linha preta) para os quatro diferentes tipos de otimização com dois objetivos.

Fonte: (RAHNAMAYAN et al., 2020)

2.1.1 NSGA-II

O algoritmo NSGA II (*Non-domination Based Genetic Algorithm*) é um Algoritmo Genético (*Generic Algorithm - GA*) multi-objetivo desenvolvido por Deb et al. (2002). A Figura 3 ilustra a estrutura básica de um GA com a definição de uma população e sua respectiva avaliação, e seus operadores de Seleção, Cruzamento e Mutação.

Para realizar a seleção esse algoritmo baseia-se em uma ordenação elitista das soluções por dominância (*Pareto ranking*), conhecido como ordenação de não dominância (*non dominated*

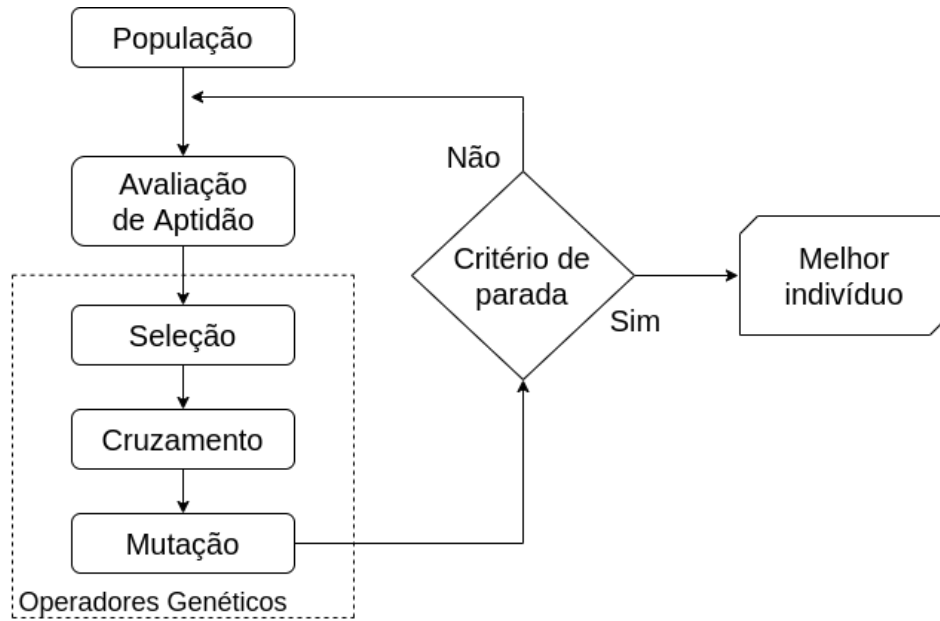


Figura 3 – Fluxograma da execu  o de um Algoritmo Gen tico e suas caracter sticas b sicas.

Fonte: (MARISA et al., 2021)

sorting). Essa ordena  o consiste em classificar indiv duos S de um conjunto I em diversos n veis ou *rankings* $N_{i_1}; N_{i_2}; \dots; N_{i_d}$ em que d   o n mero de n veis N_i de domin ncia, de modo que N_{i_1} domina os conjuntos $N_{i_2}; N_{i_3}; \dots; N_{i_d}$, o n vel N_{i_2} domina os conjuntos $N_{i_3}; N_{i_4}; \dots; N_{i_d}$ e assim sucessivamente, definindo os graus de domin ncia.

Desse modo, os indiv duos dominantes ou n o dominados de todo o conjunto I est o no n vel N_{i_1} . Ao passo que o n vel N_{i_2} cont m todos indiv duos que dominam o conjunto $I - N_{i_1}$. Analogamente isso se aplica a todos os demais n veis at  chegar no n vel N_{i_d} no qual est o contidos os indiv duos que s o dominados por todos os outros.

O Algoritmo 1 representa a l gica de programa  o que define o *non dominated sorting*. Inicialmente as linhas 2 a 16 servem para calcular o Ids , conjunto de indiv duos que s o dominados pelo indiv duo S ; e o Nds   o n mero de indiv duos que dominam o indiv duo S , isso para cada indiv duo S do conjunto de solu  es I . De modo que os indiv duos que possuem $Nds = 0$ est o contidos no n vel N_{i_1} , sendo ent o um Pareto- timo at  o momento.

Posteriormente, nas linhas 17 a 30, percorre-se o conjunto de indiv duos dominados Ids para cada indiv duo S de N_{i_d} . O contador Nds de cada indiv duo   dimin  do em 1 unidade, em que se o $Nds = 0$ a solu  o S pertence ao n vel corrente. Ent o o algoritmo repete-se at  que todos os indiv duos sejam classificados de acordo com seu n vel N_{i_d} .

Com essa ordena  o completa, o pr ximo passo   calcular a *crowding distance*, a qual consiste na m dia da dist ncia entre dois indiv duos adjacentes de cada individuo da popula  o para todos os objetivos. Desse modo, torna-se poss vel ranquear os indiv duos quanto   sua distribui  o no conjunto solu  o, priorizando os indiv duos mais "espalhados", servindo como

Algorithm 1 Ordenação de não dominância

```

1: Sejam  $S \in I$  e  $p \in I$ ;
2: for  $S$  do
3:    $Ids = \emptyset$ ;
4:    $Nds = 0$ ;
5:   for  $p$  do
6:     if  $S$  domina  $p$  then
7:        $Ids = Ids \cup p$ ;
8:     end if
9:     if  $p$  domina  $S$  then
10:       $Nds = Nds + 1$ ;
11:    end if
12:  end for
13:  if  $Nds = 0$  then
14:     $Ni_1 = Ni_1 \cup S$ ;
15:  end if
16: end for
17:  $d = 1$ ;
18: while  $Nid \neq \emptyset$  do
19:    $aux = \emptyset$ 
20:   for  $S \in Nid$  do
21:     for  $p \in Ids$  do
22:        $Ndp = Ndp - 1$ 
23:       if  $Ndp = 0$  then
24:          $aux = aux \cup p$ ;
25:       end if
26:     end for
27:   end for
28:    $d = d + 1$ ;
29:    $Nid = aux$ ;
30: end while
31: return  $Ni_1, Ni_2, \dots, Ni_d$ 

```

critério de desempate. Com isso, garante-se que o conjunto encontrado seja o mais próximo possível do conjunto Pareto-ótimo.

Algorithm 2 crowding distance

```

1:  $dist_s = \emptyset$ 
2: for  $obj$  do
3:   classificar  $obj$  por  $f_{obj}$ ;
4:    $dist_0 = dist_{NS-1} = \infty$ ;
5:   for  $S = 1$  to  $N - 1$  do
6:      $dist_s = dist_s + f_{obj}(S + 1) - f_{obj}(S - 1)$ ;
7:   end for
8: end for
9: return  $dist_s$ ;

```

O Algoritmo 2 apresenta como calcular $dist_s$, que é o valor do crowding distance do indivíduo S do conjunto I , NS é o número de soluções contidas em I e f_{obj} é o valor da obj-ésima função objetivo da solução S .

A representação gráfica do *crowding distance* pode ser observada na Figura 4, em que representa-se uma estimativa do tamanho do maior "cubóide", determinado por meio do cálculo do *crowding distance* em que $S - 1$ representa o ponto da fronteira-pareto anterior ao ponto avaliado e $S + 1$ o ponto da fronteira-pareto posterior. Permite-se que os indivíduos melhor espalhados se tornem melhores candidatos para a próxima geração, garantindo uma melhor diversidade de soluções.

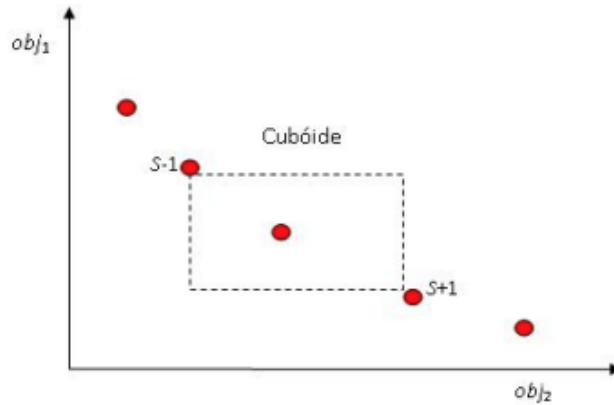


Figura 4 – Representação gráfica do *crowding distance*.

Fonte: (BLANK; DEB, 2020)

O NSGA-II utiliza um processo de seleção por torneio, o qual considera o *rank* e a aptidão (*fitness*) de cada indivíduo S , os quais são representados pela fronteira Ni_d e distância de multidão $dist_s$, respectivamente. Esse processo é observado na Figura 5. Desse modo o indivíduo a ser escolhido para ser utilizado para gerar descendentes será aquele a possuir um menor valor de *rank*, determinado pela ordenação de não dominância, ou seja o indivíduo S será escolhido se possuir um *rank* menor que p ($rank_S < rank_p$). Caso os indivíduos S e p possuírem o mesmo *rank*, então o critério de desempate é aquele que possuir a maior distância de multidão, de modo que o indivíduo S será um melhor candidato em relação ao indivíduo p se $rank_S = rank_p$ e $dist_S > dist_p$.

O funcionamento do NSGA-II é descrito no Algoritmo 3. Primeiramente é criada uma população inicial P_0 com NS indivíduos e em seguida são ordenados de acordo com o Algoritmo 1. Então, forma-se uma nova população G com a aplicação de operadores genéticos tradicionais.

Devido o NSGA-II possui critério de seleção elitista, somente os indivíduos mais aptos formarão a próxima geração, os quais são definidos pelos Algoritmos 2 e 1. Esse processo repete-se até que atinja o número máximo de gerações predefinido (GER_{max}). Há também possibilidade de o analista definir outro critério de parada.

Os Algoritmos 1, 2, 3 descritos na presente monografia foram baseados no trabalho de Silva, Souza e Guimarães (2015).

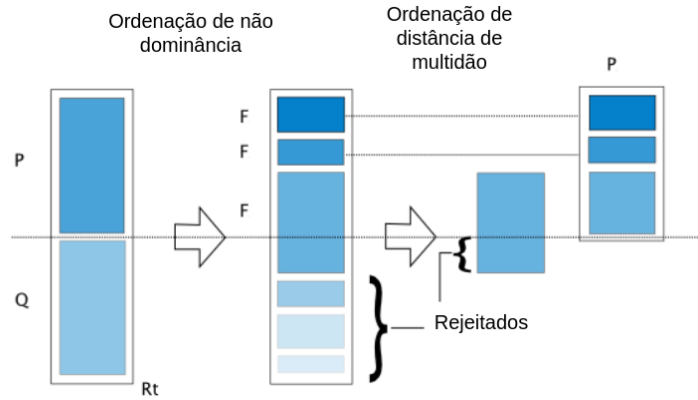


Figura 5 – Etapas de avaliação do NSGA-II; sendo a primeira, a partir da população seleciona-se os indivíduos da fronteira-pareto por meio da ordenação de não dominância; e segunda, a partir dos indivíduos selecionados utiliza-se a ordenação de *crowding distance* como critério de desempate

Fonte: (BLANK; DEB, 2020)

Algorithm 3 NSGA II

```

1:  $G_0 \leftarrow \emptyset$ ;
2: Crie uma população inicial  $P_0$ ;
3: Aplique operadores genéticos sobre  $P_0$  para criar  $G_0$ 
4: for  $g = 0$  to  $GER_{max}$  do
5:   Ordene  $R_g = P_g \cup G_g$ ;
6:    $d = d + 1$ ;
7:   while  $|P_{g+1} + Ni_d| \leq NS$  do
8:     Calcule crowding distance para  $Ni_d$ ;
9:      $P_{g+1} = P_{g+1} \cup Ni_d$ ;
10:     $d = d + 1$ ;
11:  end while
12:  Calcule crowding distance para  $Ni_d$ ;
13:  Classifique  $Ni_d$ ;
14:  Copie as primeiras  $NS - |P_{g+1}|$  soluções de  $Ni_d$  para  $P_{g+1}$ ;
15:  Gere uma nova geração  $G_{g+1}$  aplicando os operadores genéticos em  $P_{g+1}$ ;
16: end for
17: return  $P_{final}$ ;

```

2.1.2 R-NSGA-II

O trabalho (DEB; JAYAVELMURUGAN, 2006) propõe uma nova abordagem para algoritmos multi-objetivo, o qual agrega a preferência do usuário utilizando um método de ponto de referência. O R-NSGA-II como foi denominado pelo autor, é uma versão modificada do NSGA-II (citado na seção anterior). O que difere ambos é que o primeiro tem o objetivo de guiar a otimização de acordo com o ponto de referência fornecido pelo analista, caso contrário ele terá um funcionamento idêntico ao NSGA-II. Outro aspecto em que se diferem, é a métrica

de *crowdig distance* que é modificada, para o R-NSGA-II passa a ser denominada “distância de preferência” pois passa a representar o quanto as soluções estão próximas do(s) ponto(s) de referência, sendo utilizada também como critério de desempate entre as soluções da fronteira-pareto para respectiva geração.

O funcionamento da métrica de distância de preferência é da seguinte maneira:

- Realiza-se o cálculo da distância euclidiana normalizada (Equação 2.3) para cada uma das soluções da fronteira em relação a cada ponto de referência. Em seguida classifica-se essas soluções em ordem crescente de acordo com a distância calculada.
- Realiza-se o cálculo para todos pontos de referência da distância de preferência de uma determinada solução que é a classificação mínima que lhe foi atribuída entre todos os diferentes *ranks*, de modo que as soluções que possuem menores valores de distância de preferência serão as preferidas pelo método de seleção para formação da nova população (DEB; JAYAVELMURUGAN, 2006).

$$Dist(x, g) = \sqrt{\sum_{i=1}^M w_i \left(\frac{f_i(x) - f_g(x)}{f_i^{max} - f_i^{min}} \right)^2} \quad (2.3)$$

em que:

- M : número máximo de objetivos.
- w_i : vetor de pesos.
- x : solução considerada.
- g : ponto de referência especificado.
- f_i^{max} e f_i^{min} : valores máximo e mínimo para o i -ésimo objetivo.

Com a finalidade de controlar a diversidade de soluções selecionadas que estão próximas aos pontos de referência, o R-NSGA-II conta com a implementação de uma estratégia de seleção chamada *ε-clearing*. Esse processo consiste em, primeiramente escolher uma solução aleatoriamente no conjunto de soluções não dominadas. O próximo passo é agrupar todas soluções que possuem uma distância euclidiana normalizada nos valores objetivos menores que ϵ , estas recebem, em seguida, valores altíssimos de valor de distância de preferência a fim de excluir a possibilidade dessas soluções serem escolhidas para permanecer nas próximas gerações. Então escolhe-se novamente outra solução do conjunto de soluções não dominadas, diferente da escolhida anteriormente, e executa-se o mesmo processo até que se tenha percorrido todo o conjunto de soluções não dominadas.

O valor de ϵ é um parâmetro, informado pelo analista podendo ser alterado de acordo com a aplicação e pode ser diferente para cada objetivo. Na Figura 6 é possível notar o efeito da variação do ϵ nas soluções, em que f_1 e f_2 são as funções objetivos e cada linha representa uma fronteira-pareto com seus respectivos pontos (“bolinhas”) para diferentes valores de ϵ .

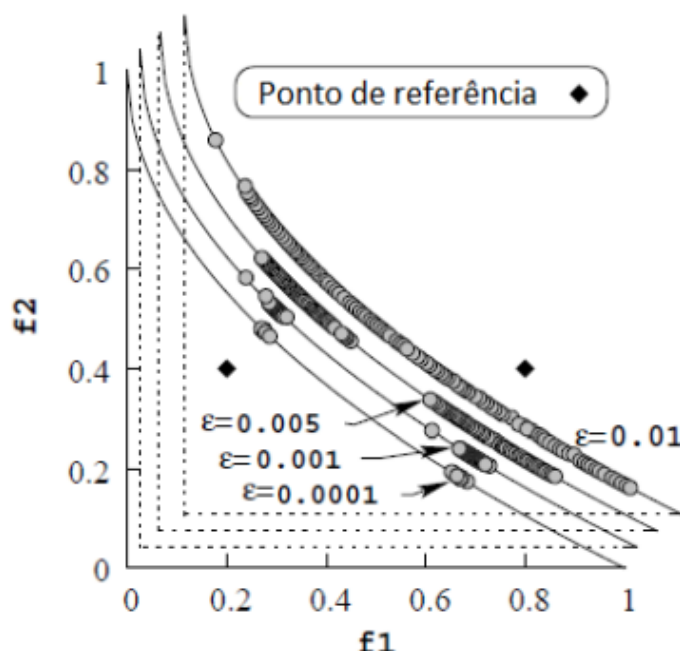


Figura 6 – Representação da influência do fator ϵ nas soluções da fronteira-pareto no R-NSGA-II.

Fonte: (DEB; JAYAVELMURUGAN, 2006)

2.2 Métricas

Nessa seção são apresentadas as métricas usadas para avaliar o desempenho da otimização.

2.2.1 Decidibilidade

Dentre as características extraídas de um indivíduo por meio de uma modalidade biométrica nenhuma é idêntica, mesmo sendo extraída da mesma pessoa (MOAZED, 2020). Isso ocorre devido a influência da iluminação do ambiente, ruídos, entre outras interferências do ambiente. Além disso, durante a aquisição dos dados biométricos pode ocorrer da pessoa estar em outra posição ou do indivíduo ter alterado alguma característica como, por exemplo, corte de cabelo ou algum ferimento no dedo. Os pares de amostras coletadas de uma pessoa são nomeadas como genuínas ou *intra-class* e para os pares de amostras dos diferentes indivíduos nomeia-se de impostores ou *inter-class*.

Desse modo, diversos trabalhos como de Daugman (2001), Daugman (2004), Wang et al. (2012) e Silva et al. (2018), utilizam como métrica de avaliação para o reconhecimento de íris a

decidibilidade (d), a qual define quão distante estão as distribuições das amostras *intra-class* das amostras *inter-class*.

Para melhor compreensão, a Figura 7 mostra à esquerda a distribuição das amostras coletadas de um mesmo indivíduo, genuínos, *same* ou *intra-class*, e a direita à distribuição das amostras de indivíduos de classes diferentes, impostores, *different* ou *inter-class*. A decidibilidade, é responsável por medir o quão distante estão essas duas distribuições, por meio da equação

$$d = \frac{|\mu_E - \mu_I|}{\sqrt{\frac{1}{2}(\sigma_E^2 + \sigma_I^2)}} \quad (2.4)$$

onde μ_E e μ_I são as médias das distribuições *inter-class* e *intra-class* respectivamente, e σ_E e σ_I , os desvios padrão as distribuições *inter-class* e *intra-class* respectivamente.

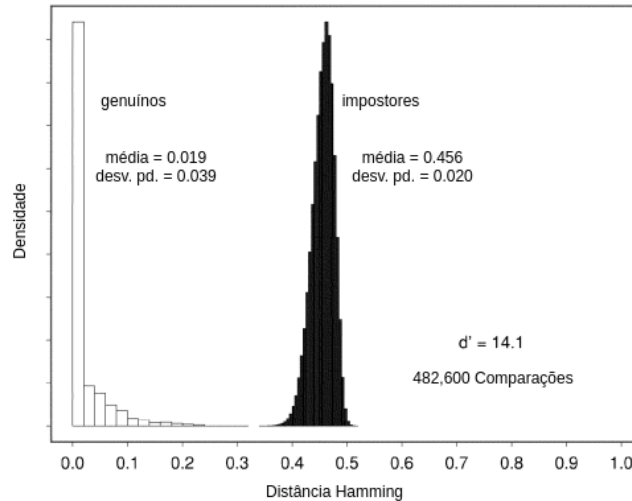


Figura 7 – Distribuições de amostras genuínas e de amostras de impostores para o cálculo da Decidibilidade.

Fonte: (DAUGMAN, 2001)

O cenário ideal é a distância entre os pares genuínos ser a menor possível e inversamente, a distância entre os pares impostores a maior possível. Desse modo, o sistema biométrico consegue distinguir melhor os dados do indivíduo dos demais dados, ou seja, é um sistema com maior facilidade de tomar alguma decisão.

2.2.2 Equal Error Rate (EER)

A Taxa de Erro Igual, ou do inglês *Equal Error Rate* (EER) corresponde ao ponto onde as taxas de falsa aceitação (FAR - *False Acceptance Rates*) e as taxas falsa rejeição (FRR - *False Rejection Rates*) são iguais e pode ser obtido por meio do gráfico das curvas de Características operacionais do receptor, ou do inglês Receiver Operating Characteristic - ROC como pode ser

observado na Figura 8. Na primeira curva ROC, em vermelho, utilizou-se a taxa de verdadeiros positivos (TPR - *True Positive Rate*), em que é descrito por $TPR = TP/(TP + FN)$ e a segunda curva ROC, tracejada em azul, utilizou-se a taxa de falso positivos (FPR - *False Positive Rate*), descrito como $FPR = FN/(FN + TP)$. TP são os dados verdadeiros positivos ou *True Positives* - dados classificados corretamente como positivos; FN são os dados falso negativos ou *False Negatives* - dados classificados erroneamente como negativos; TN são os dados verdadeiros negativos ou *True Negatives* - são os dados corretamente classificados como negativos e FP são os dados falso positivos ou *False Positives* - são os dados erroneamente classificados como positivos.

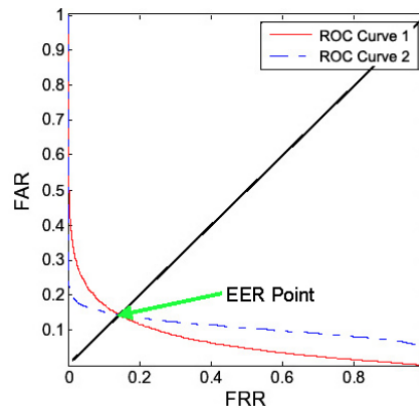


Figura 8 – Gráfico de 2 curvas ROC, em que o ponto coincidente de ambas é denominado o EER.

Fonte: (DU, 2007)

2.2.3 Hiper-volume

O Hiper-Volume(HV) é uma das métrica de avaliação de algoritmos multi-objetivos mais comumente utilizadas (FONSECA; PAQUETE; LÓPEZ-IBÁÑEZ, 2006). Nele, calcula-se o volume da região coberta entre os pontos das soluções na fronteira de Pareto P e um ponto de referência W , de modo que para cada solução $S \in P$, constitui-se um hipercubo v_s sob referência o ponto W .

O ponto W é encontrado por meio do vetor com os piores valores da função objetivo. A união de todos os hipercubos calculados representa o HV, de acordo com a Equação 2.5.

$$HV = \sum_{S \in P} v_s \quad (2.5)$$

A interpretação dessa métrica se dá por um alto valor de HV indicando que houve um espalhamento entre as soluções de P ao passo que houve também uma melhor convergência das soluções ao longo das gerações. O HV é ilustrado na Figura 9 representado pela área em cinza, em que o ponto r é o ponto de referência para calcular o hiper-volume, os eixos x e y são os eixos dos respectivos objetivos e os pontos p_1 até p_4 são os pontos da fronteira-pareto.

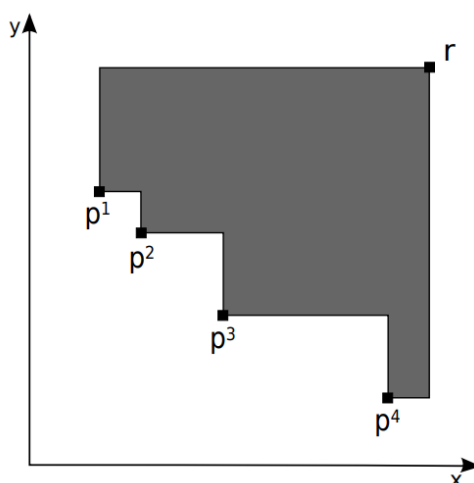


Figura 9 – Exemplificação gráfica do Hiper-volume para 2 dimensões, sendo o mesmo representado pela área em cinza.

Fonte: (GUERREIRO; FONSECA; PAQUETE, 2021)

2.3 Trabalhos Relacionados

Um dos trabalhos mais recentes na literatura que relaciona otimização multi-objetivo com sistemas biométricos é o proposto por Li et al. (2021). Esse trabalho não utiliza somente uma modalidade biométrica mas sim a fusão de duas: fala e a face. É um sistema multimodal, onde não só a face mas também a fala carregam muitas informações a respeito da emoção humana. Então os pesquisadores propõe um método de reconhecimento de emoção multimodal baseado na fusão da fala e expressão facial, utilizam-se de uma rede neural convolucional profunda ou, em inglês, *deep convolutional neural network* (DCNN), e uma rede neural convolucional profunda separável ou, em inglês, *deep separable convolutional neural network* (DSCNN) para os respectivos reconhecimentos da fala e da face. Os autores utilizam dois coeficientes (ω_1 e ω_2) para combinar linearmente o resultado das duas redes neurais, tal que $\omega_1 + \omega_2 = 1$ e os mesmos precisam ser otimizados. Além da otimização desses coeficientes o trabalho ainda tem por objetivo otimizar a acurácia e a taxa de *recall* do sistema ao longo das iterações da rede neural. A taxa de *recall* é a divisão do número de previsões corretas feitas pelo sistema pelo total de previsões feitas pelo sistema. Devido essa grande quantidade de objetivos (ω_1 , ω_2 , acurácia e a taxa de *recall*) o autor utiliza os algoritmos desenvolvidos especialmente para problemas com muitos objetivos como o NSGA-III (DEB; JAIN, 2014), MOEA/DD (LI et al., 2015), HypE (BADER; ZITZLER, 2011), e PEAS (REN; HUANG; HAO, 2015) para otimizar os coeficientes.

Sun e Zhou (2015) propõe em seu artigo um abordagem relacionada à um sistema biométrico com modalidade facial. Uma vez que esses sistemas se comportaram perfeitamente em computadores, o principal objetivo do autor foi diminuir a complexidade e grande quantidade de *features* desse tipo de sistema biométrico para melhor adaptação e desempenho em apare-

lhos celulares, haja vista, que em aparelhos celulares há uma limitação em relação ao poder processamento e duração da energia da bateria quando comparado à computadores. O autor conseguiu melhorar o algoritmo de otimização do Método de Descida Supervisionada ou, do inglês, *Supervised Descent Method* (SDM) (XIONG; TORRE, 2015) sob o aspecto de reduzir a dimensão de *features* necessários para localizar e identificar as características dentro do banco de dados.

Outro trabalho recente é o proposto por Narayan, G e M (2020), o qual é apresentado um Sistema Biométrico completo para classificação utilizando impressões digitais. O autor apresenta todas as etapas para se desenvolver um sistema, captura dos dados, faz um pre-processamento desses dados para em seguida realizar a extração das *features* e salvá-las em um banco de dados. Por fim, utiliza-se de um algoritmo de aprendizado de máquina *K-nearest neighbors* (KNN) (ZHANG; LI, 2021) como classificador do sistema. Esse trabalho segue todas etapas descritas no capítulo de Introdução e definida por Prabhakar, Pankanti e Jain (2003). Narayan, G e M (2020) ainda apontam como trabalhos futuros uma abordagem multi-objetivo, visando também selecionar somente as características extraídas das impressões digitais mais relevantes e ainda não prejudicar o desempenho do sistema, semelhante ao que é proposto neste trabalho.

Na literatura, dentre os trabalhos que têm como seu enfoque principal a biometria por meio da íris, o trabalho proposto por Daugman (2004) é um dos pioneiros. Seu artigo apresenta singularidades dos padrões de íris, características que tornam esse tipo de biometria um dos melhores no quesito de distinção. Ele ainda é pioneiro em desenvolver propriedades estatísticas, sendo a principal a decidibilidade a qual é um dos objetivos de otimização da presente monografia.

Vale destacar também o trabalho proposto por Wang et al. (2012), o qual utiliza o banco de dados de íris NICE-II (o mesmo utilizado na presente monografia). O autor utilizou uma segmentação binária e integro-diferencial para determinar o limite pupilar e externo da íris, aliado à um modelo de *Deep Learning* chamado Adaboost (VAGHELA; GANATRA; THAKKAR, 2009) como classificador para o banco de dados e utiliza como métrica de avaliação a decidibilidade e o *Equal Error Rate* sendo seus resultados foram respectivamente 1.82 e 19%.

Finalmente o trabalho ao qual essa monografia dá continuidade é o proposto por Silva et al. (2018), o qual apresenta um novo estado da arte para reconhecimento da íris e da região periocular. Utiliza-se o algoritmo de Otimização de Enxame de Partículas ou do inglês *Particle Swarm Optimization* (PSO) (KENNEDY; OBAIAHNAHATTI, 1997), que é baseado em algoritmos genéricos, mas o que o difere dos demais é que ele considera o comportamento do grupo e o seu próprio para atualizar seus pesos. Sua função é selecionar as características mais relevantes extraídas do modelo de CNN sendo, dessa forma, um problema mono-objetivo. Seus resultados para a modalidade da íris foram a decidibilidade de 2.2, EER de 14,68%, com a seleção de 256 *features*.

3 DESENVOLVIMENTO

No presente capítulo descreve-se a metodologia aplicada no desenvolvimento do trabalho. Apresenta-se o problema da seleção de características (*features*), a base de dados utilizada e como foi a implementação e lógica utilizada no desenvolvimento dos experimentos.

A Figura 10 apresenta as etapas do desenvolvimento desta monografia, em que primeiramente foi necessário compreender e descrever a base de dados NICE.II. A segunda etapa foi extrair as características da base de dados utilizando a metodologia do trabalho de Silva et al. (2018). Em seguida apresentou-se uma descrição do problema o qual essa monografia se propõe a resolver e a última etapa é a implementação dos algoritmos da otimização, informando quais parâmetros foram utilizados e qual metodologia utilizada pra criar a primeira população do GA.

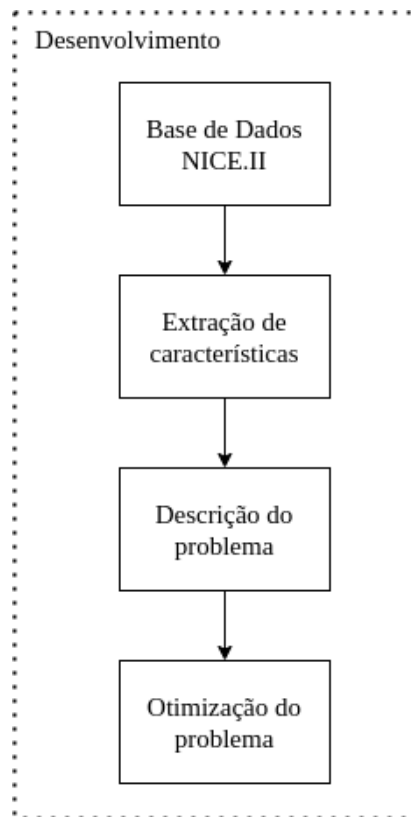


Figura 10 – Etapas do desenvolvimento desta monografia.

Fonte: produzida pelo autor

Os códigos desenvolvidos nessa monografia podem ser acessados por meio do repositório no link: <https://github.com/matheuslemesUFOP/multi-objective-iris-descriptor.git>

3.1 Base de dados

Proença et al. (2010) foram pioneiros no desenvolvimento de uma base de dados com imagens da região da íris. O foco principal era avaliar quais são os impactos da aquisição de imagens de íris em ambientes não-controlados. Chamada de UBIRIS.V2 possui 1.102 imagens de 261 indivíduos diferentes, de maneira que para simular tal cenário não-controlado as imagens foram capturadas de diferentes distâncias, ângulos e condições de iluminação. As imagens possuem resolução de 800×700 , 72 ~ dpi e cores em 24 bits. A Figura 11 apresenta exemplos das imagens presentes na UBIRIS.V2.



Figura 11 – Exemplos de imagens da base de dados UBIRIS.V2 .

Fonte: (PROENÇA et al., 2010)

A partir dessa base de dados, os mesmos criadores desenvolveram a competição NICE - *Noisy Iris Challenge Evaluation*, a qual foi sub-dividida. Na sua primeira fase, a NICE.I, tem por foco o desafio de segmentação e na segunda, a NICE.II, é voltada para o processo de reconhecimento. A presente monografia utiliza somente a segunda etapa da competição, em que a extração das características se deu pela metodologia do trabalho de Silva et al. (2018). Entretanto, o foco do presente trabalho não é esse, mas sim utilizar o menor número de características para o reconhecimento sem que o sistema piore seu desempenho e seu poder de decisão.

A competição NICE.II possui uma base de treinamento com 1000 imagens de 171 indivíduos diferentes em que não há balanceamento ou distribuição uniforme entre as amostras por indivíduo. A base de teste, ou avaliação também possui 1000 imagens, entretanto são somente 141 indivíduos diferentes. Os indivíduos da base de treinamento não necessariamente estão presentes na base de testes, ou seja, há indivíduos que estão da base de treino mas não estão na base de teste e vice-versa.

3.2 Extração das Características

Devido ao foco dessa monografia não ser a extração das características, utilizou-se os mesmos dados extraídos no trabalho do [Silva et al. \(2018\)](#), em que as imagens da competição NICE.II são submetidas a um processo de normalização e representação em coordenadas polares (*Daugman Rubber Sheet*).

Além de normalizadas as imagens da íris foram redimensionadas para 224 x 224 para adequar ao tamanho da entrada da rede neural convolucional utilizada. Para arquitetura base utilizada foi a VGG ([PARKHI; VEDALDI; ZISSERMAN, 2015](#)) criada para face. Os pesos da rede foram otimizados utilizando a estratégia de *Stochastic Gradient Descent* ([THEODORIDIS, 2015](#)), mas para o contexto de íris.

Após o treinamento, a última camada utilizada para classificação é retirada, de modo que a saída seja o vetor de características extraído de tamanho 256. O modelo foi denominado como *Deep Eye Descriptor - DID*, conforme mostrado na Figura 12.

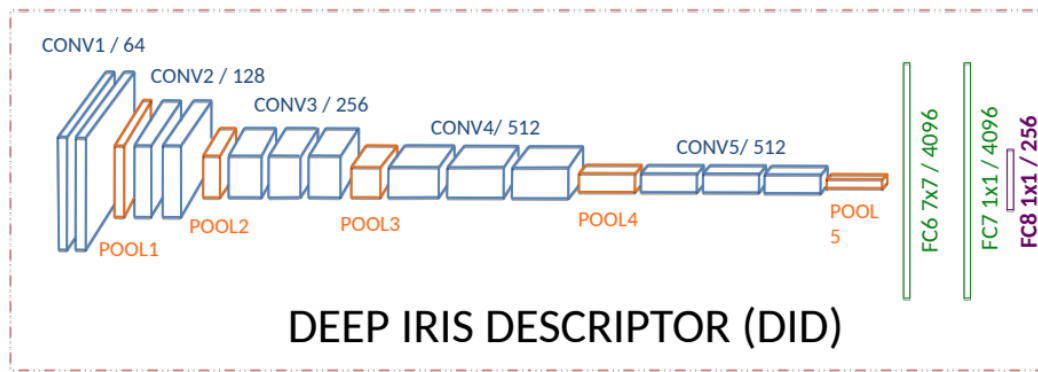


Figura 12 – Descritor usado para a extração de características da íris.

Fonte: ([SILVA et al., 2018](#))

3.3 Problema da Seleção de Características

O Problema de Seleção de Características (PSC) é definido por um conjunto de características ou *features* VF , de tamanho n que descreve um indivíduo I , de modo que a seleção de *features* se dá por um vetor de escolha binário $[0 \text{ ou } 1]$ X de mesmo tamanho n . A multiplicação do vetor VF pelo vetor X indica quais características foram selecionadas.

Quanto maior a quantidade de características n para descrever o indivíduo I , necessita-se de um maior espaço de armazenamento e também um custo de processamento maior. Para resolver esses problemas é fundamental selecionar somente as *features* mais relevantes e que melhor descrevem o indivíduo I .

Com uma menor quantidade de características, possui-se uma menor quantidade de informações que por consequência, torna mais difícil o reconhecimento de um indivíduo por parte do sistema e, consequentemente, seu poder de decisão é menor, ou seja, sua decidibilidade (descrita na subseção 2.2.1), tende a diminuir. Por isso é que a presente monografia tem como foco minimizar a quantidade de *features* ao mesmo tempo que maximiza a decidibilidade, ou seja, reduzir o custo computacional sem que prejudique a taxa de reconhecimento do sistema. Devido a esses dois objetivos serem conflitantes, tem-se um problema de otimização multi-objetivo.

A formulação matemática do problema multi-objetivo abordado pode ser definida por:

$$\text{Minimizar} \quad \sum_i x_i \quad (3.1)$$

$$\text{Maximizar} \quad \text{Decidibilidade} \left(\prod_{i,j} x_i \text{features}_{i,j} \right) \quad (3.2)$$

Sujeito a:

$$\sum_{i=1}^n x_i \geq 1, x_i \in [0, 1] \quad (3.3)$$

onde as Equações 3.1 e 3.2 representam, respectivamente, as funções objetivos que consistem em minimizar a quantidade de *features* usadas e maximizar a decidibilidade de acordo com as características selecionadas. A restrição na Equação 3.3 registra que pelo menos uma característica deve ser selecionada e já é solucionada durante a construção da primeira população que será explicada na seção seguinte.

3.4 Metodologia proposta

Na implementação do código utilizou-se a biblioteca do Python Pymoo (BLANK; DEB, 2020). Nessa biblioteca a lógica descrita nas subseções 2.1.1 e 2.1.2 para os algoritmos multi-objetivo NSGA-II e R-NSGA-II já estão implementadas, bastando definir os respectivos parâmetros. Criou-se então uma classe para o problema abordado (subseção 3.3) definindo o número de objetivos n_{obj} sendo 2, número de variáveis (n_{var}) sendo 256, pois é relacionado com a quantidade total de *features*, e com x sendo uma variável binária uma vez que é um problema de decisão binária.

Conforme descrito na seção anterior, como funções objetivos tem-se f_1 para minimizar a quantidade de *features* selecionadas e f_2 para maximizar a decidibilidade. Para critério de parada da otimização utilizou-se a quantidade de gerações, a qual, em vista do tempo de execução e convergência dos algoritmos definiu-se com 50.

De modo a facilitar a reprodução dos experimentos realizados, criou-se a primeira população “manualmente”. Adotando a metodologia de que o primeiro indivíduo seja um vetor

binário com todas as características, o segundo indivíduo exclui uma característica aleatoriamente, atribuindo 0 a essa posição, para o terceiro indivíduo exclui-se duas características aleatoriamente, e assim sucessivamente até preencher toda população. De acordo com a representação abaixo:

$$\begin{bmatrix} 1 & 1 & 1 & 1 & \dots & 1 \\ 1 & 1 & 0 & 1 & \dots & 1 \\ 1 & 0 & 1 & 1 & \dots & 0 \\ 0 & 1 & 1 & 0 & \dots & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \end{bmatrix} \quad (3.4)$$

Caso a população seja maior que o número de *features* e para que não haja um indivíduo com todas *features* eliminadas, para os indivíduos cujo índice é múltiplo do número total de *features* (256) são escolhidas todas as características, ou seja, atribui-se 1 à todos elementos do vetor e se reinicia a lógica descrita anteriormente.

Essa metodologia foi elaborada de modo que a primeira população tenha pelo menos 1 indivíduo com cada possibilidade de somatório da quantidade de características selecionadas, de 1 até 256 características, a fim de orientar melhor a seleção e a evolução dos indivíduos ao longo das gerações. Visto isso, para os experimentos utilizou-se uma população de 512 indivíduos, de modo que se tenha pelo menos 2 indivíduos com cada possibilidade de características selecionadas.

Os parâmetros básicos citados anteriormente e a metodologia de criação da primeira população foram utilizados para os dois algoritmos: NSGA-II e R-NSGA-II. Entretanto o R-NSGA-II necessita de alguns outros parâmetros, explicados na seção 2.1.2. Adotou-se para o respectivo algoritmo, como ponto de referência, 100 para o somatório da quantidade de *features* selecionadas e 2 para a decidibilidade, fator *epsilon* (ϵ) igual a 0.01, para que os indivíduos da fronteira-pareto não fiquem muito espaçados, e pesos iguais para ambos objetivos, ou seja, 0.5 para ambos, pois os dois objetivos possuem igual importância.

Ao longo das gerações é calculado o hiper-volume considerando a fronteira-pareto da respectiva geração e com ponto de referência $W = [256, 10]$, respectivamente para o somatório da quantidade de *features* selecionadas e para decidibilidade. Durante a implementação do algoritmo, para maximizar a decidibilidade, houve a necessidade de minimizar a decidibilidade negativa, devido a limitação do *framework* utilizado. Isso justifica o ponto referente a decidibilidade (10) para calcular o hiper-volume ter sido positivo.

Utilizou-se os dados de treino da NICE.II para calibrar ambos os algoritmos e a seleção feita sobre esses dados é aplicada nos dados de teste para se obter os resultados finais, apresentados no próximo capítulo. Ao final da otimização, obedecendo o critério de parada (o número de gerações), calcula-se o EER de cada ponto da fronteira-pareto da otimização. Todos esses resultados são salvos em um arquivo de texto JSON (*JavaScript Object Notation*).

4 EXPERIMENTOS E RESULTADOS

Este capítulo apresentará os resultados obtidos após execução dos algoritmos de otimização com a metodologia descrita no capítulo anterior.

Todos experimentos foram realizados em um notebook Vaio com processador Intel(R) Core(TM) i5-8250U CPU @ 1.60GHz 4 cores, 12 GB DDR4 de memória RAM e todos algoritmos implementados utilizando a linguagem de programação Python juntamente com o *framework* Pymoo (BLANK; DEB, 2020).

Com a finalidade de se conseguir um ponto-base para análise e comparação da otimização com NSGA-II e com R-NSGA-II, utilizou-se o método de Análise de Componentes Principais ou *Principal Component Analysis* (PCA). É um análise estatística em que seu objetivo consiste em encontrar um meio de condensar a informação contida em várias variáveis originais em um conjunto menor de componentes, que para o problema descrito são as *features*, com uma perda mínima de informação (XANTHOPOULOS; PARDALOS; TRAFALIS, 2013).

Para implementação do PCA, utilizou-se o *framework* em Python do Sklearn (PEDREGOSA et al., 2011) e o resultado foi que para descrever pelo menos 90% do conjunto de *features* é necessário aproximadamente 50 características. Fica mais evidente, ao analisar a Figura 13, em que o gráfico se aproxima de uma reta em aproximadamente 50 *features*. De modo que o número de características em que a inclinação é mais próxima de zero é a quantidade que melhor descreve todo o conjunto de dados.

Calculando a decidibilidade e o EER com as *features* selecionadas pelo PCA, obteve-se 1,4376 e 21,82% respectivamente.

4.1 NSGA-II

O algoritmo do NSGA-II foi implementado na linguagem Python utilizando a biblioteca do Pymoo (BLANK; DEB, 2020), com uma população de 512 indivíduos e 50 gerações com critério de parada. O tempo total da otimização foi de 8 horas e 26 minutos.

O resultado foi a obtenção da seguinte fronteira-pareto da quantidade de *features* e decidibilidade com seus respectivos EER, apresentados na Tabela 2:

Observa-se uma redução das *features* em aproximadamente 80%. Uma notória redução e ainda uma melhora na decidibilidade e no EER do sistema, visto que selecionando todas as 256 características a decidibilidade é de 1,5051 e seu EER é de 19,96%. Vale ressaltar que há um ganho sobre a aplicação do PCA também.

É possível observar como foi a evolução da fronteira-pareto do NSGA-II ao longo das gerações na Figura 14 abaixo. Nota-se uma melhora constante na evolução dos indivíduos do

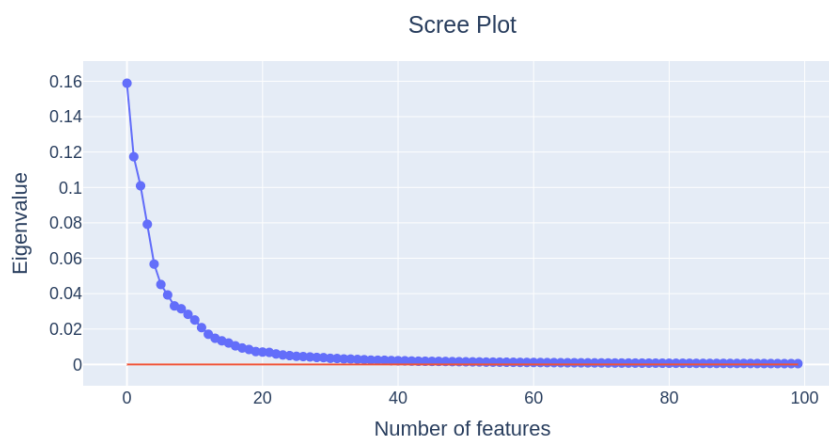


Figura 13 – Gráfico da Seleção de Características Utilizando o PCA, eixo x são as características selecionadas e o eixo y é a inclinação da curva. De modo que o número de características em que a inclinação é próxima de zero, melhor descreve o conjunto de dados

Fonte: produzido pelo autor

Tabela 2 – Pontos da Fronteira-pareto obtido com NSGA-II.

Quantidade de <i>features</i>	Decidibilidade	EER
36	1,4626	20,47%
37	1,5213	19,94%
38	1,5226	19,95%
39	1,5372	19,56%
40	1,5425	19,66%
41	1,5492	19,55%
43	1,6043	19,11%
45	1,6060	19,07%
46	1,6068	18,94%
50	1,6274	18,67%
51	1,6631	17,96%
53	1,6641	17,97%
55	1,6761	18,34%
65	1,6761	18,05%
76	1,6763	18,11%
82	1,6789	18,11%

Fonte: produzido pelo autor.

algoritmo genético.

Quanto ao hiper-volume da otimização o resultado obtido foi de 2559,82. Analisando a Figura 15, observando a evolução do hiper-volume ao longo das gerações nota-se que essa melhoria foi devido a fronteira-pareto, ou seja, conforme a fronteira-pareto melhorou o hiper-volume também melhorou.

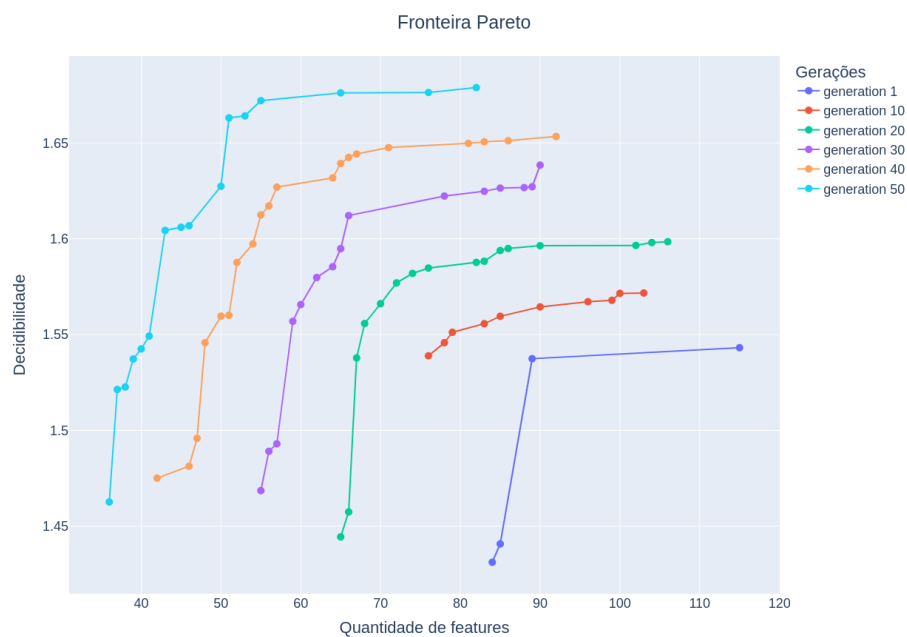


Figura 14 – Evolução da fronteira-pareto ao longo das gerações utilizando o NSGA-II.

Fonte: produzido pelo autor

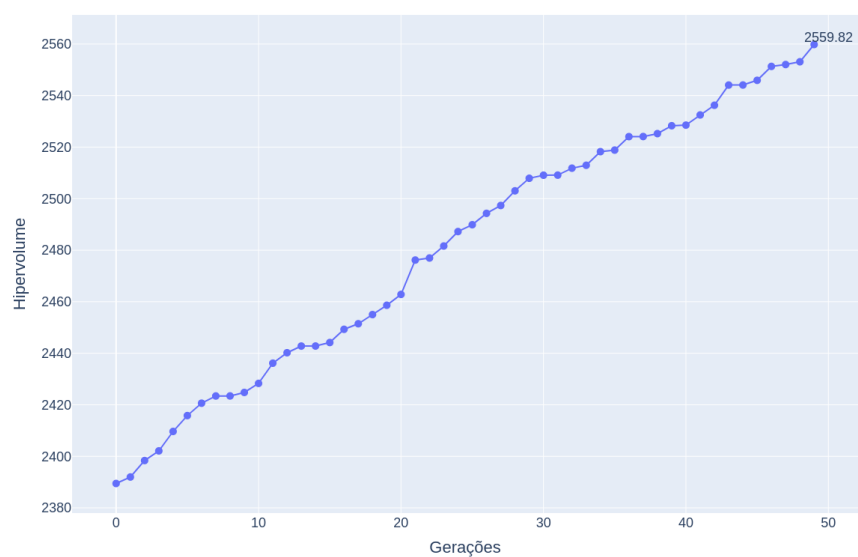


Figura 15 – Hiper-volume da otimização ao longo das gerações utilizando o NSGA-II.

Fonte: produzido pelo autor

4.2 R-NSGA-II

Para o algoritmo R-NSGA-II utilizou-se a mesma biblioteca em Python Pymoo, utilizada anteriormente no NSGA-II e os mesmos parâmetros utilizados anteriormente com a população de 512 indivíduos e 50 gerações, sendo o número de gerações utilizado como critério de parada. O tempo total de otimização foi um pouco maior se comparado ao do NSGA-II, sendo de 10 horas e 54 minutos.

Selecionando todas as características o R-NSGA-II conseguiu uma decidibilidade de 1,5051 e ERR de 19,96%, resultado quase idêntico ao NSGA-II para esse caso. Como resultado final obteve-se uma redução da quantidade de *features* de aproximadamente 81.64% e decidibilidade média de 1.7954 (resultados pouco superiores aos do NSGA-II), com a seguinte fronteira-pareto da quantidade de *features* e decidibilidade com seus respectivos EER, apresentados na Tabela 3:

Tabela 3 – Pontos da Fronteira-pareto obtido com R-NSGA-II.

Quantidade de Features	Decidibilidade	EER
45	1,7924	19,81%
46	1,7960	19,56%
47	1,7963	19,52%
52	1,7972	19,46%

Fonte: produzido pelo autor

A evolução da fronteira-pareto ao longo das gerações utilizando o R-NSGA-II obteve um melhor resultado se comparado ao NSGA-II, observando a Figura 16, pois a quantidade de *features* e o EER foram praticamente os mesmos e decidibilidade foi maior. Isso ocorreu devido ao ponto de referência utilizado pelo R-NSGA-II o qual orienta melhor a busca e seleção dos indivíduos. Nota-se também que a quantidade de pontos da fronteira-pareto do algoritmo R-NSGA-II é menor se comparado ao algoritmo NSGA-II, obtendo uma fronteira menos “abrangente”, isso ocorre em vista do fator ϵ usado ter sido relativamente alto ($\epsilon = 0,01$).

Quanto a análise de duração dos dois algoritmos, o R-NSGA-II demandou mais tempo devido a maior complexidade nos cálculos associados ao modelo do algoritmo, explicado na subseção 2.1.2.

O hiper-volume obtido realizando os experimentos no algoritmo do R-NSGA-II (vide Figura 17, só passou a evoluir a partir da vigésima-oitava geração e a partir desse ponto seguiu uma melhoria contínua até chegar na geração 50 com o valor de 2556,69 (muito próximo ao resultado obtido com algoritmo do NSGA-II).

Essa evolução a partir da geração 28 pode ser observada na Figura 16 pelo *gap*, espaço vazio no gráfico, entre as fronteiras-pareto da geração 20 e 30.

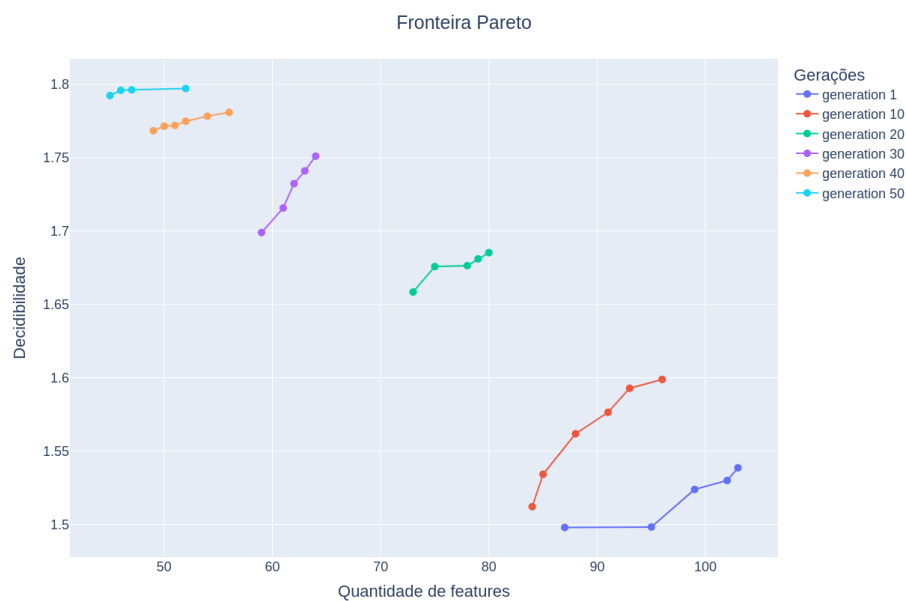


Figura 16 – Evolução da fronteira-pareto ao longo das gerações utilizando o R-NSGA-II.

Fonte: produzido pelo autor

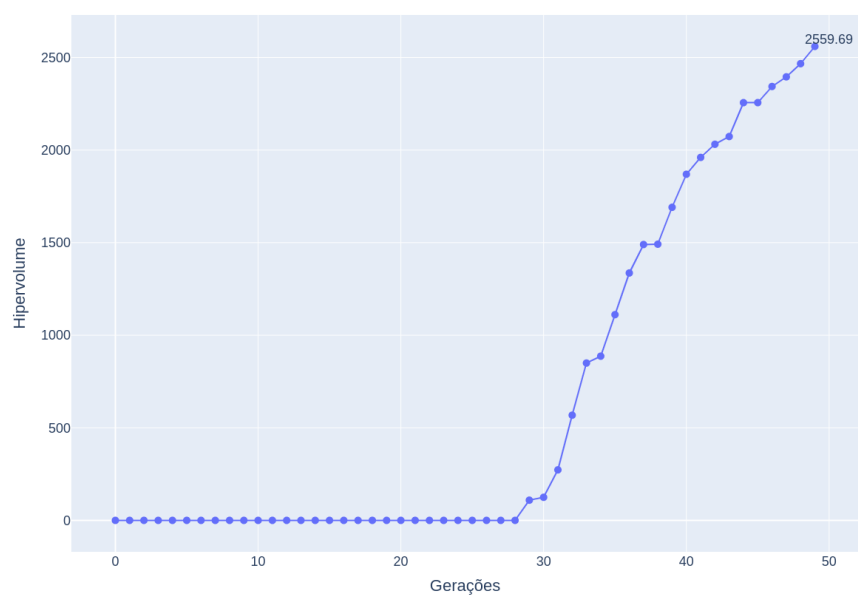


Figura 17 – Hiper-volume da otimização ao longo das gerações utilizando o R-NSGA-II.

Fonte: produzido pelo autor

5 CONCLUSÃO

A abordagem proposta apresentou resultados para um pioneiro de otimização multi-objetivo para seleção de características para a modalidade biométrica íris para os dados da competição NICE.II. O enfoque foi utilizar dois algoritmos genéticos de otimização multi-objetivo NSGA-II e R-NSGA-II para objetivos conflitantes de maximizar a decidibilidade e minimizar a quantidade de características selecionadas. Os dois algoritmos conseguiram, com menos *features*, uma decidibilidade maior que quando se comparado a otimização utilizando todas as *features* para o mesmo algoritmo, com bons parâmetros de hiper-volume e EER. O algoritmo do R-NSGA-II sobrepujou o NSGA-II considerando o a superioridade no seu resultado de decidibilidade com praticamente a mesma quantidade de características, uma vez que ele foi capaz de selecionar melhor quais características são mais relevantes graças a sua “natureza” de utilizar um ponto de referência durante a evolução da otimização. Além disso a otimização realizada em ambos algoritmos conseguiram resultados dentro do esperado em relação a quantidade de *features* quando se comparado ao PCA, mas ao analisar a decidibilidade e o EER os dois algoritmos tiveram desempenho muito superior, devido a otimizar ser capaz de selecionar melhor as características mais relevantes para descrever a biometria íris para o banco de dados utilizado.

O atual cenário de modalidades biométricas é uma área bastante promissora, haja vista a possibilidade de criação de sistemas mais robustos e confiáveis. Entretanto, a quantidade de características após sua extração é muito grande, o que é inviável devido ao custo maior de processamento e armazenamento dessas informações. Considerando esse problema, é que essa monografia se propôs a resolve-lo utilizando esses algoritmos de otimização multi-objetivo.

5.1 Trabalhos Futuros

As principais contribuições apresentadas nessa monografia tem por foco os dados a competição NICE.II. Sugere-se como trabalhos futuros utilizar outros algoritmos multi-objetivos como SPEA2 (*Strength Pareto Evolutionary Algorithm*) e o PSO (*Particle Swarm Optimization*) para realizar a otimização; avaliar melhor os parâmetros dos algoritmos, como quantidade de gerações maior como critério de parada, os parâmetros do R-NSGA-II como por exemplo, utilizar mais pontos de referência, utilizar um computador com melhor poder de processamento para realizar a otimização, além de uma nova proposta para construção da primeira população.

Outra opção seria utilizar a mesma abordagem de selecionar as características mais relevantes sem perder o poder de decisão do sistema mas, aplica-lo em outras modalidades biométricas como, por exemplo, a face ou a impressão digital.

REFERÊNCIAS

- ARROYO, J. E. C. *Heurísticas e metaheurísticas para otimização combinatória multiobjetivo*. Tese (Doutorado) — Universidade de Campinas, Campinas, Brasil, 2002. Disponível em: <<https://www.teses.usp.br/teses/disponiveis/3/3143/tde-19112004-165342/publico/Kleber-Cap3.pdf>>. Acesso em: 25 nov. 2021. Citado na página 20.
- BADER, J.; ZITZLER, E. Hype: An algorithm for fast hypervolume-based many-objective optimization. *Evolutionary computation*, 2011. v. 19, p. 45–76, 03 2011. Citado na página 30.
- BLANK, J.; DEB, K. pymoo: Multi-objective optimization in python. *IEEE Access*, 2020. v. 8, p. 89497–89509, 2020. Citado 4 vezes nas páginas 24, 25, 35 e 37.
- DAUGMAN, J. Statistical richness of visual phase information: Update on recognizing persons by iris patterns. *International Journal of Computer Vision*, 2001. v. 45, p. 25–38, 10 2001. Citado 2 vezes nas páginas 27 e 28.
- DAUGMAN, J. How iris recognition works. *Circuits and Systems for Video Technology, IEEE Transactions on*, 2004. v. 14, p. 21 – 30, 02 2004. Citado 2 vezes nas páginas 27 e 31.
- DEB, K. *Multi-Objective Optimization Using Evolutionary Algorithms*. [S.l.: s.n.], 2001. Citado na página 20.
- DEB, K.; JAIN, H. An evolutionary many-objective optimization algorithm using reference-point-based nondominated sorting approach, part i: Solving problems with box constraints. *Evolutionary Computation, IEEE Transactions on*, 2014. v. 18, p. 577–601, 08 2014. Citado na página 30.
- DEB, K.; JAYAVELMURUGAN, S. Reference point based multi-objective optimization using evolutionary algorithms. In: . [S.l.: s.n.], 2006. v. 2, p. 635–642. Citado 3 vezes nas páginas 25, 26 e 27.
- DEB, K. et al. A fast and elitist multiobjective genetic algorithm: Nsga-ii. *Evolutionary Computation, IEEE Transactions on*, 2002. v. 6, p. 182 – 197, 05 2002. Citado na página 21.
- DING, C.; TAO, D. Robust face recognition via multimodal deep face representation. *IEEE Transactions on Multimedia*, 2015. 11 2015. Citado na página 16.
- DU, Y. Rethinking the effective assessment of biometric systems. *Spie Newsroom*, 2007. 01 2007. Citado na página 29.
- DUARTE, O. C. M. B. *Histórico da Biometria*. 2010. Biometria - Histórico. Disponível em: <https://www.gta.ufrrj.br/grad/10_1/1a-versao/assinatura/index.html>. Acesso em: 18 out. 2021. Citado na página 15.
- FONSECA, C.; PAQUETE, L.; LÓPEZ-IBÁÑEZ, M. An improved dimension-sweep algorithm for the hypervolume indicator. In: . [S.l.: s.n.], 2006. p. 1157 – 1163. Citado na página 29.
- GREENBERG, S.; BOYLE, M.; LABERGE, J. Pdas and shared public displays: Making personal information public, and public information personal. *Personal Technologies*, 1999. v. 3, p. 2–3, 04 1999. Citado na página 15.

GUERREIRO, A.; FONSECA, C.; PAQUETE, L. The hypervolume indicator: Computational problems and algorithms. *ACM Computing Surveys*, 2021. v. 54, p. 1–42, 07 2021. Citado na página 30.

JAIN, A.; ROSS, A.; PRABHAKAR, S. An introduction to biometric recognition. *Circuits and Systems for Video Technology, IEEE Transactions on*, 2004. v. 14, p. 4 – 20, 02 2004. Citado na página 16.

JIANG, H.; LEARNED-MILLER, E. Face detection with the faster r-cnn. In: . [S.l.: s.n.], 2017. p. 650–657. Citado na página 16.

KENNEDY, J.; OBAIAHNAHATTI, B. A discrete binary version of the particle swarm algorithm. In: . [S.l.: s.n.], 1997. v. 5, p. 4104 – 4108 vol.5. ISBN 0-7803-4053-1. Citado na página 31.

LI, K. et al. An evolutionary many-objective optimization algorithm based on dominance and decomposition. *IEEE Transactions on Evolutionary Computation*, 2015. v. 19, p. 694–716, 09 2015. Citado na página 30.

LI, M. et al. Multimodal emotion recognition model based on a deep neural network with multiobjective optimization. *Wireless Communications and Mobile Computing*, 2021. v. 2021, p. 1–10, 08 2021. Citado na página 30.

LUZ, E. et al. Deep periocular representation aiming video surveillance. *Pattern Recognition Letters*, 2017. v. 114, 12 2017. Citado na página 17.

I. MALHAS e A. RENFORTH. *Iris recognition systems*. 2015. Citado na página 18.

MARISA, E. et al. Método de segmentação de objectos em imagens baseado em contornos activos e algoritmo genético. 2021. 11 2021. Citado na página 22.

MOAZED, K. Iris anatomy. In: _____. [S.l.: s.n.], 2020. p. 15–29. ISBN 978-3-030-45755-6. Citado 2 vezes nas páginas 18 e 27.

NARAYAN, L.; G, S.; M, S. Fingerprint recognition and its advanced features. *International Journal of Engineering Research and*, 2020. V9, 04 2020. Citado na página 31.

PARETO, V. *Cours D'Economie Politique*. 1. ed. [S.l.]: F. Rouge, 1896. Acesso em: 25 nov. 2021. Citado na página 20.

PARKHI, O.; VEDALDI, A.; ZISSERMAN, A. Deep face recognition. In: . [S.l.: s.n.], 2015. v. 1, p. 41.1–41.12. Citado na página 34.

PEDREGOSA, F. et al. Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 2011. v. 12, p. 2825–2830, 2011. Citado na página 37.

PRABHAKAR, S.; PANKANTI, S.; JAIN, A. Biometric recognition: Security and privacy concerns. *Security & Privacy, IEEE*, 2003. v. 1, p. 33 – 42, 04 2003. Citado 4 vezes nas páginas 16, 17, 18 e 31.

PROENÇA, H. et al. The UBIRIS.v2: A Database of Visible Wavelength Iris Images Captured On-the-Move and at-a-distance. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2010. v. 32, p. 1529–1535, 08 2010. Citado na página 33.

PROENÇA, H.; NEVES, J. Deep-prwis: Periocular recognition without the iris and sclera using deep learning frameworks. *IEEE Transactions on Information Forensics and Security*, 2017. PP, p. 1–1, 11 2017. Citado na página 17.

RAHNAMAYAN, S. et al. Ranking multi-metric scientific achievements using a concept of pareto optimality. *Mathematics*, 2020. v. 8, p. 956, 06 2020. Citado na página 21.

REN, H.; HUANG, X.; HAO, J. Finding robust adaptation gene regulatory networks using multi-objective genetic algorithm. *IEEE/ACM Transactions on Computational Biology and Bioinformatics*, 2015. v. 13, p. 1–1, 01 2015. Citado na página 30.

SAEED, K. Biometrics principles and important concerns. In: _____. [S.l.: s.n.], 2012. p. 3–20. ISBN 978-1-4614-5607-0. Citado na página 16.

SCHNEIER, B. The uses and abuses of biometrics. *Communications of the ACM*, 1999. v. 42, 08 1999. Citado na página 15.

SILVA, A.; SOUZA, M.; GUIMARÃES, V. Um algoritmo baseado na metaheurística lahc para resolver o problema de planejamento operacional de lavra em minas a céu aberto. In: . [S.l.: s.n.], 2015. Citado na página 24.

SILVA, P. H. et al. Multimodal feature level fusion based on particle swarm optimization with deep transfer learning. In: IEEE. *2018 IEEE Congress on Evolutionary Computation (CEC)*. [S.l.], 2018. p. 1–8. Citado 9 vezes nas páginas 15, 16, 18, 19, 27, 31, 32, 33 e 34.

SUN, Z.; TAN, T. Ordinal measures for iris recognition. *IEEE transactions on pattern analysis and machine intelligence*, 2009. v. 31, p. 2211–26, 12 2009. Citado na página 18.

SUN, Z.; ZHOU, T. A face recognition optimization algorithm based on statistical characteristics of feature points. *Beijing Jiaotong Daxue Xuebao/Journal of Beijing Jiaotong University*, 2015. v. 39, p. 128–134, 04 2015. Citado na página 30.

THEODORIDIS, S. Stochastic gradient descent. In: _____. [S.l.: s.n.], 2015. p. 161–231. ISBN 9780128015223. Citado na página 34.

TICONA, W. G. C. *Algoritmos evolutivos multi-objetivo para a reconstrução de árvores filogenéticas*. Tese (Doutorado) — Instituto de Ciências Matemáticas e de Computação, Universidade de São Paulo, São Carlos, Brasil, 2008. Disponível em: <<https://www.teses.usp.br/teses/disponiveis/55/55134/tde-02042008-142554/pt-br.php>>. Acesso em: 25 nov. 2021. Citado na página 20.

VAGHELA, V.; GANATRA, A.; THAKKAR, A. Boost a weak learner to a strong learner using ensemble system approach. 2009. 03 2009. Citado 2 vezes nas páginas 18 e 31.

VERMA, G.; CHAKRABORTY, R. A hybrid privacy preserving scheme using finger print detection in cloud environment. *Ingénierie des systèmes d'information*, 2019. v. 24, p. 343–351, 08 2019. Citado na página 16.

WANG, Q. et al. Adaboost and multi-orientation 2d gabor-based noisy iris recognition. *Pattern Recognition Letters - PRL*, 2012. v. 33, 06 2012. Citado 4 vezes nas páginas 16, 18, 27 e 31.

XANTHOPOULOS, P.; PARDALOS, P.; TRAFALIS, T. Principal component analysis. In: _____. [S.l.: s.n.], 2013. p. 21–26. ISBN 978-1-4419-9877-4. Citado na página 37.

XIONG, X.; TORRE, F. De la. Global supervised descent method. In: . [S.l.: s.n.], 2015. p. 2664–2673. Citado na página [31](#).

YIN, X.-C.; LIU, C.-P.; HAN, Z. Feature combination using boosting. *Pattern Recognition Letters*, 2005. v. 26, p. 2195–2205, 10 2005. Citado na página [18](#).

ZHANG, S.; LI, J. Knn classification with one-step computation. *IEEE Transactions on Knowledge and Data Engineering*, 2021. PP, p. 1–1, 10 2021. Citado na página [31](#).