

*Combinação de Classificadores Floresta de
Caminhos Ótimos aplicados no
Reconhecimento Facial.*

Jair José da Silva

Julho / 2016

Dissertação de Mestrado em Ciência da Computação

Combinação de Classificadores Floresta de Caminhos Ótimos aplicados no Reconhecimento Facial

Esse documento corresponde ao Projeto de Pesquisa apresentado à Banca Examinadora para qualificação no curso de Mestrado em Ciência da Computação da Faculdade Campo Limpo Paulista.

Campo Limpo Paulista, 28 de Julho de 2016.

Jair José da Silva

Dr. Luis Mariano del Val Cura
Orientador

Faculdade Campo Limpo Paulista
Programa de Mestrado em Ciência da Computação

**“Combinação de Classificadores Floresta de Caminhos Ótimos
Aplicados no Reconhecimento Facial”**

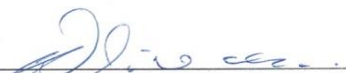
Jair José da Silva

Dissertação de Mestrado apresentado ao Programa
de Mestrado em Ciência da Computação da
Faculdade Campo Limpo Paulista, como parte dos
requisitos para a obtenção do título de Mestre em
Ciência da Computação.

Membros da Banca:



Prof. Dr. Luis Mariano del Val Cúra
(Orientador –FACCAMP)



Prof. Dr. Osvaldo Luiz de Oliveira
(FACCAMP)



Prof. Dr. Norton Trevisan Roman
(USP)

Campo Limpo Paulista, 28 de julho de 2016.

FICHA CATALOGRÁFICA

Dados Internacionais de Catalogação na Publicação (CIP)

Câmara Brasileira do Livro, São Paulo, Brasil.

Silva, Jair José da

Combinação de Classificadores Floresta de Caminhos Ótimos aplicados no reconhecimento facial / Jair José da Silva. Campo Limpo Paulista, SP: FACCAMP, 2016.

Orientador: Profº. Dr. Luis Mariano del Val Cura

Dissertação (Programa de Mestrado em Ciência da Computação) – Faculdade Campo Limpo Paulista – FACCAMP.

1. Classificadores de padrão. 2. Reconhecimento facial. 3. Optimum Path Forest - OPF. 4. Imagem facial. 5. Comparação Bayesiana. I. Del Val Cura, Luis Mariano. II. Campo Limpo Paulista. III. Título.

CDD-005.1

Agradecimentos

Ao eterno Deus que estava comigo mesmo antes de eu existir. Ele preparou o minha vida e esquadrinha meu coração. A meu orientador que no tempo todo deste trabalho se empenhou para que chegássemos até aqui, compartilhando todo a sua experiência e conhecimento sobre as matérias pertinentes ao trabalho. Aos meus irmãos, pessoas importantes na minha infância, inspiradoras na minha adolescência e orgulho dos meus dias atuais. A minhas filhas, meus tesouros, meninas de quem só tenho a me orgulhar, pela preocupação, torcida, incentivo e expectativa em meu sucesso. Aos amigos que de coração, perto ou distante, sempre torceram e se alegram com cada passo que consigo avançar. A minha esposa Tereza, quem viveu e compartilhou comigo o esforço dispendido para a realização deste trabalho, acreditando, torcendo e colaborando com tudo o quanto pôde para o meu sucesso e principalmente, aos meus eternos e queridos pais. Não puderam estar nesta minha caminhada, mas não tenho dúvida de que aplaudem, mais do que qualquer outro, o meu êxito.

Os pensamentos abaixo são para homenagear a todos estas pessoas importantes em minha vida e importantes neste tempo, em que realizei meu curso de mestrado em ciências da computação. Também para aqueles que acreditam: estamos neste mundo para aprender, pois é o conhecimento um importante tesouro. Cada um de nós tem de conquista-lo com nossos esforços.

“Deus é uma ser o ser supremo. Que apenas conhece a palavra sim.”

“O segredo dos que triunfam é: nunca desistirem.”

“Um sábio foi contratado por um rei para responder as perguntas de seus súditos.

Um dia um súdito ao fazer uma pergunta ao sábio, recebe a seguinte resposta: - Bem eu não tenho a resposta, mas vou tentar consegui-la.

O homem indignado reclama: - Como, você não é sábio coisa nenhuma! Por acaso não foi contratado para respondera todas as nossas dúvidas?

O sábio simpaticamente sorri para o jovem arrogante e responde:

- Não meu amigo! Eu estou aqui para responder as perguntas que eu souber responder. Pois se eu recebesse uma moeda para cada questão que ainda pretendo aprender; nem todo dinheiro que existe no mundo seria o suficiente.”

Resumo. Esta dissertação apresenta uma proposta de utilização de uma família de classificadores supervisionados Floresta de Caminhos Ótimos (OPF) aplicados ao reconhecimento facial. Nesta proposta, os descritores faciais utilizados são extraídos da imagem diferença de duas imagens faciais comparadas. Para decidir se as imagens comparadas pertencem ao mesmo indivíduo, o descritor da imagem diferença é classificado em duas classes: a classe que representa imagens diferença do mesmo indivíduo (classe intrapessoal) e a classe que representa imagens diferença de indivíduos diferentes (classe interpessoal). Para esta classificação, o método de decisão do OPF é adaptado e transformado em uma função que define a similaridade do descritor com cada classe. Nesta proposta, os resultados de similaridade de cada um dos classificadores da família são integrados através de algoritmos de votação. O uso de múltiplos classificadores permite sua utilização em uma implementação paralela do reconhecimento facial. Os resultados experimentais obtidos mostram a vantagem do uso de múltiplos classificadores quando comparados com os resultados de um único classificador.

Abstract. This dissertation proposes the use of a family of supervised classifiers Optimum-Path Forest (OPF), applied to face recognition problem. In this approach, facial feature vectors are extracted from the difference image of two facial images compared. To decide if the compared images belong to the same individual, each difference feature vector is classified into two classes: the class representing difference images of the same individual (intrapersonal class) and the class representing difference images of different individuals (interpersonal class). In this work, the OPF decision method is adapted as a function that defines the similarity of the feature vector with each class. The similarity values computed by classifiers of the family are integrated by voting methods. This approach can be an alternative implementation of a parallel face recognition algorithm the experimental results show that the use of multiple classifiers improves the precision if compared with one classifier approach.

Sumário

1.	Introdução	1
2.	Conceitos Básicos de Biometria	6
	Tipos de Biometrias	7
	Biometria da Impressão Digital	8
	Biometria da Retina	9
	Biometria da Íris	10
	Biometria da Face	11
	Biometria da Voz	13
	Módulos de um Sistema Biométrico	14
	Módulo Sensor	14
	Módulo Avaliação da Qualidade	14
	Módulo de Extração de Características	15
	Módulo de Tomada de Decisão	15
	Módulo de Banco de Dados	16
	Fases de um Sistema Biométrico	16
	Reconhecimento por Verificação	18
	Reconhecimento por Identificação	19
	Reconhecimento por Verificação com Registro da Amostra Original	20
	Algumas Métricas de Avaliação de Algoritmos Biométricos	22
	Considerações Finais	25
3.	Métodos de Reconhecimento Facial	27
	Reconhecimento Facial	27
	<i>Eigenfaces</i>	28
	Elastic Bunch Graph Method (<i>EBGM</i>)	31

Reconhecimento Facial Utilizando Classificadores	33
Quantidades de Classes do Classificador	35
Abordagem Bayesiana	35
Considerações Finais	38
4. Classificador Floresta de Caminhos Ótimos	39
Classificador Floresta de Caminhos Ótimos	39
Algoritmo de Treinamento	40
Algoritmo de Classificação	44
Classificador Floresta de Caminhos Ótimos Aprimorado	46
Aprimoramento com Vários Classificadores OPFs e Métodos de Votação	48
Considerações Finais	48
5. Família de Classificadores OPF para Reconhecimento Facial	49
Etapas do Processo	51
Considerações Finais	56
6. Experimentos e Análise dos Resultados	57
6.1. Metodologia	57
6.2. Experimentos com Diferentes Funções de Distância	63
6.4. Experimentos com Famílias de Classificadores	68
6.4.1. Experimentos com o Método de Votação Média Parcial	68
6.4.2. Experimentos com o Método de Votação Média Absoluta	73
6.5 Análise dos Resultados	77
7. Conclusões e Trabalhos Futuros	79
8. Referências	81
9. Apêndices A. Gráficos para a comparação do OPF com abordagem por votação utilizando várias funções de distância	86

Lista de Tabelas

Tabela 2.1. Tipos de Biometria (Bonato <i>et al.</i> , 2010).	8
Tabela 6.1. Métricas de Avaliação obtidas com classificador OPF com amostras <i>Eigenfaces</i> . Comparação dos resultados para cada função de distância.	64
Tabela 6.2. Métricas de Avaliação obtidas com classificador OPF e descritores do <i>EBGM</i> . Comparação dos resultados das funções de distâncias.	65
Tabela 6.3. Resultados com descritores gerados pelo método <i>Eigenfaces</i> e função de distância Euclidiana e votação com média parcial.	70
Tabela 6.4. Resultados com descritores gerados pelo método <i>Eigenfaces</i> e função de distância Manhattan e votação com média parcial.	70
Tabela 6.5. Resultados com descritores gerados pelo método <i>Eigenfaces</i> e função de distância Chi-Squared e votação com média parcial.	70
Tabela 6.6. Resultados com descritores gerados pelo método <i>EBGM</i> e função de distância Euclidiana e votação com média parcial.	71
Tabela 6.7. Resultados com descritores gerados pelo método <i>EBGM</i> e função de distância Manhattan e votação com média parcial.	72
Tabela 6.8. Resultados com descritores gerados pelo método <i>EBGM</i> e função de distância Chi-Squared e votação com média parcial.	72
Tabela 6.9. Resultados com descritores gerados pelo método <i>Eigenfaces</i> e função de distância Euclidiana e votação com média absoluta.	74
Tabela 6.10. Resultados com descritores gerados pelo método <i>Eigenfaces</i> e função de distância Manhattan e votação com média absoluta.	75
Tabela 6.11. Resultados com descritores gerados pelo método <i>Eigenfaces</i> e função de distância Chi-Squared e votação com média absoluta.	75
Tabela 6.12. Resultados com descritores gerados pelo método <i>EBGM</i> e função de distância Euclidiana e votação com média absoluta.	76
Tabela 6.13. Resultados com descritores gerados pelo método <i>EBGM</i> e função de distância Manhattan e votação com média absoluta.	76

Tabela 6.14. Resultados com descritores gerados pelo método *EBGM* e função de distância Chi-Squared e votação com média absoluta.

77

Listas de Figuras

Figura 2.1. Modelo de Ficha de Identificação Digital utilizada em inquérito policial.	9
Figura 2.2. Esquema de um Globo Ocular.	10
Figura 2.3. Biometria da Íris (Costa <i>et al.</i> , 2006).	11
Figura 2.4. <i>Eigenfaces</i> (Silva, 2013).	12
Figura 2.5. <i>Elastic Bunch Graph Method</i> (Lu, X, 2003).	13
Figura 2.6. Módulo Sensor. Realiza a aquisição da biometria.	14
Figura 2.7. Módulo de Avaliação da Qualidade e Extração de Características.	15
Figura 2.8. Módulo de Correspondência e Tomada de Decisão. Compara os descritores e retorna um valor de similaridade.	15
Figura 2.9. Etapa de registro.	17
Figura 2.10. Reconhecimento por verificação.	18
Figura 2.11. Reconhecimento por identificação.	19
Figura 2.12. Reconhecimento por Verificação com registro de amostra original	21
Figura 2.13. Funções FMR e FNMR de distribuição cumulativa das similaridades de comparações impostoras e genuínas. Pontos de limiar das métricas EER, FMR100, FMR1000 e FMRZero.	23
Figura 2.14. Gráfico de demonstração das taxas de erros EER, FMRs.	24
Figura 2.15. Gráfico da Curva ROC – <i>Receiver Operating Characteristic</i> .	25
Figura 3.1. <i>Eigenfaces</i> .	29
Figura 3.2. (a) Grafo da face; (b) Jet; (c) Image Graph; e (d) Bunch Graph. (Wiskott et al., 1997).	32
Figura 3.3. (a) grafo de imagem de entrada; e (b) Bunch Graph, os <i>Jets</i> em cinza são os que possuem maior semelhança com <i>Jets</i> do grafo de imagem de entrada. (Wiskott et al., 1997).	33

Figura 4.1. Algoritmo de treinamento do Classificador Floresta de Caminhos Ótimos.	41
Figura 4.2. Construção da Árvore Geradora de Custo Mínimo (AGM).	42
Figura 4.3. Relação de adjacência de um Grafo Completo.	42
Figura 4.4. Vértices protótipo dentro da ACM.	43
Figura 4.5. Floresta de Caminhos Ótimos (OPF).	44
Figura 4.6. Algoritmo de Classificação da Floresta de Caminhos Ótimos.	44
Figura 4.7. Processo de Classificação de uma nova amostra Z.	45
Figura 4.8. Algoritmo de Classificação da Floresta de Caminhos Ótimos.	46
Figura 5.1. Algoritmo de treinamento da família de classificadores.	51
Figura 5.2. Algoritmo para dividir a conjunto de descritores de treinamento.	52
Figura 5.3. Algoritmo para calcular o valor de similaridade dos descritores.	53
Figura 5.4. Algoritmo de Classificação do MOPF.	55
Figura 6.1. Resultados com um classificador OPF utilizando descritores do método <i>Eigenfaces</i> e várias funções de distância.	63
Figura 6.2. Resultados com um classificador OPF utilizando descritores do método <i>EBGM</i> e várias funções de distância.	64
Figura 6.3. Resultados com descritores do método <i>Eigenfaces</i> utilizando diferentes métodos de votação com famílias de (a) 04 (b) 05 (c) 10 (d) 15 (e) 20 (f) 25 (g) 50 classificadores.	66
Figura 6.4. Resultados com descritores do método <i>EBGM</i> utilizando diferentes métodos de votação com famílias de (a) 04 (b) 05 (c) 10 (d) 15 (e) 20 (f) 25 (g) 50 classificadores.	67
Figura 6.5. Gráficos com os resultados dos experimentos com o método de votação de média parcial.	69
Figura 6.6. Gráficos com os resultados dos experimentos com o método de votação de média absoluta.	73

Figura 9.1. OPF com abordagem por votação com descritores *Eigenfaces*. 87

Figura 9.2. OPF com abordagem por votação com descritores *EBGM*. 88

1. Introdução

O reconhecimento biométrico tenta responder o problema de reconhecer com precisão uma pessoa, através de uma característica física, ou comportamental, que seja única para essa pessoa. No universo destas características podemos citar a íris, a voz humana, as impressões digitais, a imagem da face humana, dentre outras. Jain *et al.* (2006) e Jain *et al.* (2007).

O estudo de soluções para este problema é do interesse de muitas áreas da sociedade, como na área da segurança civil e segurança pública (Jain, 2008). Por exemplo, na área comercial a garantia do reconhecimento da identidade pode evitar enormes prejuízos nas transações comerciais. Já nos processos jurídicos a veracidade da identidade das pessoas envolvidas no processo é imprescindível, Costa *et al.* (2006) e Bonato *et al.* (2010).

O uso massivo de sistemas biométricos de identificação está associado fundamentalmente a sistemas governamentais. Os bancos de dados biométricos utilizados por estes sistemas têm como características o grande volume de dados armazenados, assim como a quantidade de consultas realizadas sobre estes dados. Em 2010 o Sistema DHS IDENT gerenciado pelo Departamento de Segurança dos Estados Unidos armazenava 110 milhões de registros biométricos e realizava a verificação de 125 000 identidades por dia (Graves, 2010). Outro exemplo deste volume de dados é o Projeto de Identificação Nacional da Índia que prevê o armazenamento e consulta de 1,2 bilhões de registros biométricos (Panchumorthy *et al.*, 2012). Alguns autores tem colocado o problema da consulta a um banco de dados biométrico como um problema de *Big Data* com busca de soluções através de processamento paralelo e distribuído (Panchumorthy *et al.*, 2012) (Kohlwey *et al.*, 2011).

O processo para reconhecer a identidade de uma pessoa com a total confiança é complexo e mesmo para o sofisticado mecanismo natural do ser humano, é impossível garantir o reconhecimento com total certeza.

Dentre as características biométricas mais interessantes temos a face, caracterizada pela facilidade de captura das imagens e pela boa aceitação pelos usuários. Para realizar o reconhecimento facial são comparadas duas imagens da face, verificando assim, se elas pertencem ou não ao mesmo indivíduo.

Algoritmos convencionais de reconhecimento facial iniciam seu processamento extraindo descritores biométricos das imagens faciais através de métodos de extração de características faciais. (Jain *et al.*, 2000). Estes descritores, em forma de vetores de características, devem possuir estruturas e valores que caracterizem a identidade de cada indivíduo. Geralmente, o reconhecimento facial é realizado comparando o descritor obtido de uma imagem de consulta com o descritor de uma imagem alvo, atribuindo-se a esta comparação um valor de similaridade. Este valor de similaridade é comparado com um valor limiar de decisão prefixado para decidir se as imagens pertencem ou não ao mesmo indivíduo. Quando a comparação é realizada entre descritores do mesmo indivíduo temos uma comparação genuína e caso contrário tem uma comparação impostora.

Um algoritmo classificador caracteriza um conjunto de classes. As classes são formadas por conjuntos de amostras que devem possuir características similares. Os algoritmos de classificação supervisionada possuem uma etapa inicial de treinamento que fornece ao classificador amostras, geralmente em forma de vetores de características, de cada uma das classes. Posteriormente, o classificador deve ser capaz de classificar novas amostras entre uma das classes definidas. Alpaydin (2010) e Theodoridis *et al.* (2009).

Algoritmos classificadores supervisionados têm sido utilizados para reconhecimento facial. A abordagem convencional define uma classe para cada indivíduo, Afonso *et al.* (2012) e Papa *et al.* (2009a), ou seja, cada classe é treinada com amostras na forma de descritores biométricos de um mesmo indivíduo. O reconhecimento facial é realizado quando um novo descritor é classificado na classe de um indivíduo.

Outra abordagem para o uso de classificadores foi apresentada inicialmente por Moghaddam *et al.* (2000). Esta abordagem parte da hipótese de que as imagens obtidas como diferença de imagens faciais do mesmo indivíduo possuem padrões comuns e diferentes dos padrões das imagens diferenças de imagens faciais de indivíduos diferentes. Destas imagens diferença podem ser extraídos descritores em forma de vetores aplicando algoritmos de extração de características faciais. Com estes descritores podem ser criadas duas classes: a *Classe Intrapessoal* com descritores extraídos das imagens diferença de um mesmo indivíduo e a *Classe Interpessoal* com descritores extraídos das imagens diferença de indivíduos diferentes.

Desta forma, o problema é transformado em um problema de classificação em duas classes. Para realizar a comparação das imagens de consulta e alvo, primeiramente é obtida uma imagem diferença destas. A imagem diferença é obtida subtraindo a imagem alvo da imagem consulta. Desta imagem diferença é extraído um descritor que é comparado com cada uma das duas classes intrapessoal e interpessoal para estimar um valor de similaridade entre as imagens consulta e alvo. Este valor é comparado com o limiar de decisão para decidir se a comparação é genuína ou impostora. Este trabalho foi posteriormente estendido por outros autores Wang *et al.* (2003) e Chen *et al.* (2012), utilizando o mesmo princípio das duas classes. Para a extração de vetores de características faciais esta abordagem tem utilizado os métodos *Eigenfaces*, Turk *et al.* (1991a) e Turk *et al.* (1991b), e o método *EBGM*, (Wiskott *et al.*, 1997).

Nesta dissertação propomos uma extensão da abordagem de Moghaddam *et al.* (2000) através do uso de uma família de classificadores das classes intrapessoal e interpessoal. Estes classificadores são treinados com subconjuntos disjuntos e aleatórios de descritores de imagens diferença. O conjunto de descritores de treinamento é dividido igualmente entre cada classificador com o objetivo de reduzir o custo computacional do processamento de cada classificador. Para realizar o reconhecimento facial, os resultados de todos os classificadores são combinados através de métodos de votação para obter um valor de similaridade final.

A motivação para o uso de uma família de classificadores é propor uma alternativa para processamento distribuído da classificação aplicada no reconhecimento facial. Uma proposta de solução com uma família de classificadores deve ter como requisitos: (a) reduzir o custo computacional da classificação através da integração dos resultados de subproblemas resolvidos paralelamente por cada um dos classificadores da família. (b) permitir flexibilidade para a adição de novos classificadores na família. Consideramos que esta abordagem por utilizar vários classificadores independentes, pode ser a base de uma implementação distribuída e paralela de reconhecimento facial utilizando técnicas de processamento distribuído como *map-reduce* (Dean *et al.*, 2008).

Neste trabalho utilizamos o classificador supervisionado Floresta de Caminhos Mínimos (OPF), Papa *et al.* (2008) e Papa *et al.* (2009a). O OPF é um classificador baseado em grafos no qual cada amostra de treinamento é associada a um vértice e as arestas estão associadas à

similaridade ou distância entre essas amostras. O método cria uma floresta de forma tal que cada classe pode ser representada por um conjunto de árvores nessa floresta. Esta forma de construção do classificador, permite que possam ser representadas amostras espacialmente dispersas (Alpaydin, 2010). Uma nova amostra é classificada determinando qual a árvore que oferece o menor caminho, do conjunto de vértices representativos de cada classe. O classificador OPF conforme Papa (2008), apresenta resultados computacionais comparáveis no quesito da acurácia, mas resultados superiores na questão do tempo de processamento quando comparados com classificadores tradicionais, como Máquinas de Vetores de Suporte (SVM) e Redes Neurais (NN), Alpaydin (2010), Papa (2009a), Papa (2009b) e Ponti Jr (2011).

O classificador OPF atende os requisitos já que o custo computacional da classificação depende do número de amostras de treinamento. Adicionalmente, a dissertação mostra como os resultados dos classificadores OPF da família podem ser integrados para obter um resultado final de similaridade.

O objetivo desta dissertação é pesquisar o impacto do uso de uma família de classificadores OPF na precisão do reconhecimento facial. A implementação paralela desta proposta está fora do escopo da dissertação.

Para a realização desta pesquisa foram utilizadas imagens faciais do banco de dados FERET, Phillips *et al.* (2000) e Phillips *et al.* (2003), das quais foram extraídos descritores faciais das imagens diferenças utilizando as técnicas de *Eigenfaces*, Turk *et al.* (1991a) e Turk *et al.* (1991b) e Casamento de Grafos Elásticos (*EBGM*) (Wiskott *et al.*, 1997).

Foram obtidos resultados experimentais com descritores biométricos extraídos pelos métodos *Eigenfaces* e Casamento de Grafos Elásticos ou *EBGM* e integrando os resultados com vários métodos de votação. Os resultados experimentais mostram que o uso da família de classificadores OPF melhora em alguns casos a precisão do reconhecimento facial.

Nosso trabalho foi estruturado da maneira descrita a seguir. O Capítulo 2 aborda conceitos básicos de biometria, as etapas e módulos que compõem um sistema biométrico assim como as principais métricas utilizadas para avaliar a acurácia de algoritmos biométricos. A seguir o Capítulo 3 descreve os principais métodos de reconhecimento facial utilizados nesta

dissertação. Na sequência, o Capítulo 4 descreve em detalhe o classificador Floresta de Caminhos Ótimos (OPF) junto com a abordagem aprimorada e a proposta de uso vários classificadores com algoritmos de votação. O Capítulo 5 descreve a proposta de família de classificadores OPF. O Capítulo 6 descreve os experimentos realizados e analisa os resultados da pesquisa e finalmente o Capítulo 7 apresenta as conclusões e trabalhos futuros propostos.

2. Conceitos Básicos de Biometria

Este capítulo apresenta um referencial teórico básico sobre biometria. São apresentados os tipos de biometria, os módulos de um sistema biométrico, as etapas de um sistema biométrico e algumas métricas de avaliação utilizadas na biometria.

No reconhecimento biométrico automatizado de maneira geral, dada uma amostra biométrica de um indivíduo, um algoritmo específico realizará a extração das características biométricas na amostra. Estas características serão utilizadas pelos algoritmos biométricos para decidir se duas amostras correspondem ao mesmo indivíduo.

O principal uso da identificação biométrica é para o reconhecimento de pessoas em ocorrências que exigem considerado nível de segurança. No entanto, já podemos encontrar o emprego desta tecnologia em outras aplicações e ambientes, como por exemplo, para realizar o *login* em aparelhos como notebooks, acesso a smartphones e até em sistemas de travamento de veículos.

As técnicas biométricas são de interesse em muitas áreas da sociedade, como os bancos na área comercial, por exemplo. Também há interesse na área da segurança pública e particular. Nestas áreas é comum que pessoas envolvidas em transações comerciais ou em inquéritos policiais devam ser identificadas com a máxima confiança.

O reconhecimento biométrico busca identificar um indivíduo através de características físicas ou comportamentais. Estas características devem ser únicas e pertencer a todos os indivíduos do grupo. É necessário que estas características garantam a confiança da identificação, por isso preferencialmente são utilizados atributos físicos como as impressões digitais, atributos da face, da íris, da voz utilizando o som (a frequência) e outros (Zhang *et al.*, 2006).

Assim como nos atributos físicos, a biometria também pode utilizar as características comportamentais específicas da pessoa, como a maneira de caminhar, ou reconhecimento de características da voz, por exemplo, a maneira como a pessoa fala (o sotaque) dentre outras.

Uma biometria deve possuir requisitos que buscam garantir a identificação, assim como dificultar a falsificação da identidade. O primeiro requisito que abordaremos é o da

universalidade, que tem como propósito garantir que toda população de usuários do sistema possua o atributo ou, caso não possua, o sistema deve fornecer algum recurso para que o usuário possa ser autenticado. Por exemplo, algum usuário pode não ter a impressão digital legível, assim este usuário pode utilizar um cartão magnético.

Outro requisito é da unicidade, onde a característica biométrica ou atributo biométrico na prática deve ser único para cada usuário do sistema biométrico. Em outras palavras, a possibilidade de pessoas diferentes possuírem o mesmo atributo biométrico deve ser suficientemente desprezível, de maneira que garanta a confiabilidade do sistema biométrico.

O requisito da permanência exige que as características biométricas sejam imutáveis, embora os atributos biométricos na prática sejam suscetíveis à alteração pela influência do tempo, da saúde, do ambiente e outros eventos.

Também tem que ser possível medir o atributo biométrico através de um dispositivo para atender o requisito da coleta. Para este requisito o indivíduo que fornece o atributo biométrico deve aceitar o processo de coleta da característica biométrica. Alguns empecilhos para este requisito podem ser de ordem cultural, de privacidade ou mesmo questões como a higiene (Lucas *et al.*, 2006).

O requisito da confidencialidade define que o sistema biométrico deve assegurar que a informação somente seja acessada por usuários autorizados, no processo da autenticação. Já o requisito da integralidade deve garantir que somente pessoas com autorização possam alterar a informação.

A disponibilidade é o requisito que deve garantir que a informação seja acessível aos usuários que requisitarem e possuam permissão do sistema biométrico.

Tipos de Biometrias

Existem diversos tipos de biometrias que podem ser utilizados para o reconhecimento humano com objetivo de atender a necessidade da identificação (Jain *et al.*, 2007). Neste trabalho focaremos nossas pesquisas nas biometrias baseadas nas características físicas. Estas biometrias são mostradas na Tabela 2.1, onde a primeira coluna descreve alguns tipos de

biometrias físicas; na segunda coluna temos as vantagens de cada biometria; a terceira coluna mostra as desvantagens que elas oferecem e a última coluna apresenta o custo, ou seja, investimento (R\$) em equipamentos para aquisição das características biométricas.

Tabela 2.1. Tipos de Biometria (Bonato *et al.*, 2010).

Tipos	Vantagens	Desvantagens	Custo
Impressão digital	Simples, barato;	Fácil de ser fraudado;	Baixo;
Retina	Precisão;	Caro, difícil de treinar os usuários;	Alto;
Íris	Fácil aquisição, melhor precisão de todos os métodos;	Custo, baixa aceitabilidade;	Médio;
Reconhecimento facial	Similar ao processo humano,	Pode ser difícil de adquirir, baixa precisão;	Baixo a médio;
Reconhecimento por voz	Fácil aquisição;	Fácil de ser fraudado;	Baixo;

As seções a seguir trazem uma breve explanação sobre as biometrias apresentadas na Tabela 2.1.

Biometria da Impressão Digital

A primeira linha na Tabela 2.1 apresenta a biometria da impressão digital (Falguera, 2008). Esta biometria é um método de relativa aceitação no Brasil, devido ao seu histórico de utilização de forma impressa como identificação de indivíduos, para tentar solucionar casos como de homicídios, por exemplo.

A Ciência Forense já utiliza este método e contribuiu com seu conhecimento para os algoritmos atuais. Os órgãos de segurança pública no Brasil utilizam este método para cadastrar pessoas que possuam registro criminal. Também órgãos de identificação de pessoas imprimem no documento de identidade ou RG, esta biometria que possui métodos de coleta das digitais simples, rápidos e não invasivos ao corpo humano. É um método barato em relação a algumas outras biometrias, mas é um processo suscetível de fraude, pois os sistemas utilizados geralmente não identificam se a biometria utilizada para o reconhecimento, é de uma pessoa ou uma prótese produzida para enganar o sistema biométrico.

O método de reconhecimento baseado em impressões digitais consiste em realizar a

leitura das linhas das digitais. Geralmente as linhas desenhadas nas pontas dos dedos, como se pode ver na Figura 2.1. Estas linhas conhecida como minúcias, ou *minutiae* em **latim**, podem definir características únicas para cada indivíduo. Dentre as estruturas destas linhas temos os deltas, que são as estruturas onde três linhas definem a forma de um triângulo; também temos as bifurcações, que descrevem os pontos onde uma linha se divide formando duas novas linhas. Estes são apenas alguns exemplos entre muitos outros definidos pela ciência forense. Com a análise destas estruturas é possível obter um descritor para reconhecer a pessoa a quem pertence o código biométrico (Jain et al., 2000).

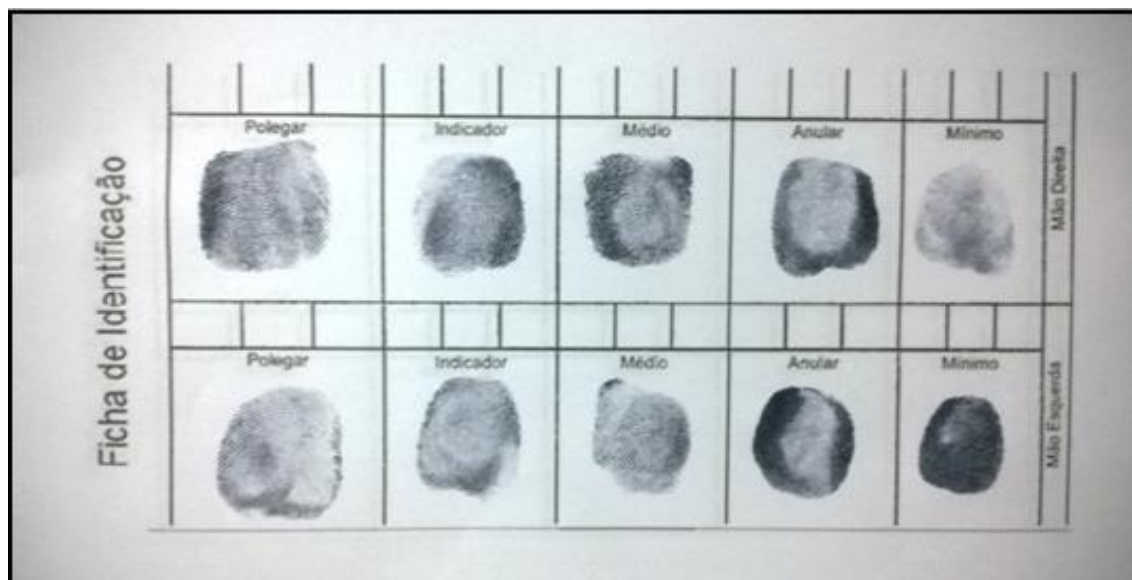


Figura 2.1. Modelo de Ficha de Identificação Digital utilizada em inquérito policial.

Biometria da Retina

Na segunda linha da Tabela 2.1 temos a biometria da retina (Theisen, 2008). A retina é uma das partes do globo ocular que forma o órgão da visão também constituído pela córnea e a esclera, estruturas que protegem o globo ocular.

O cerne desta biometria é a retina, que é o local onde a imagem invertida é apresentada ao sistema de visão e está localizada no fundo do globo ocular. Esta estrutura através das células fotorreceptoras converte os sinais luminosos em impulsos nervosos e estes sinais são levados ao cérebro através do nervo óptico (ver a Figura 2.2).

A Figura 2.2 mostra um esquema em que a luz refletida da imagem entra no globo ocular através da pupila. O cristalino que funciona como uma lente inverte os raios luminosos, fazendo com que uma imagem invertida seja apresentada na retina.

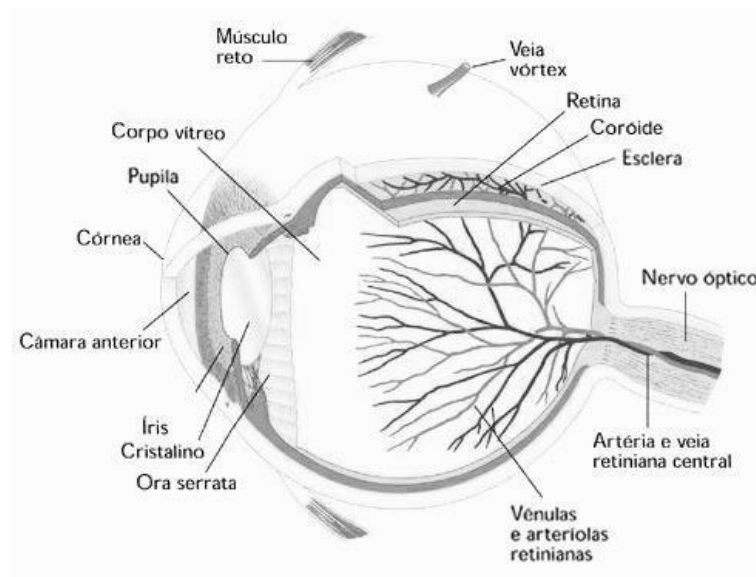


Figura 2.2. Esquema de um Globo Ocular.

A biometria da retina consiste em um escaneamento da sua área. Esta operação tem como objetivo extrair um descritor da estrutura vascular da retina que, assim como as digitais explicadas anteriormente, forma uma estrutura única para cada indivíduo. Este método possui a vantagem de que a estrutura vascular se encontra em local de difícil acessibilidade, o que diminui sensivelmente as possibilidades de fraudes. Este método é tido como uma biometria de ótima precisão, como mostra a Tabela 2.1 (Bonato *et al.*, 2010).

Um dos pontos negativos desta abordagem é a obtenção da imagem biométrica. O processo de escaneamento pode ser incômodo para alguns usuários. A pessoa pode interpretar que o escaneamento causa algum dano ao local da retina. No caso de existir esta interpretação, a propriedade da disponibilidade pode ser prejudicada neste sistema.

Biometria da Íris

Na biometria da íris (Figura 2.3) terceiro exemplo da Tabela 2.1, temos a precisão como principal vantagem. Uma das características utilizadas nesta biometria é o padrão de

pigmentação. A alta complexidade da estrutura da íris atende a condição biométrica de quanto mais aleatório o padrão, melhor para a característica biométrica, Jain *et al.* (2000) e Daugman (2003). Também podem ser utilizadas características da reação da íris a estímulos, o que torna o sistema biométrico capaz de reconhecer uma íris viva e uma prótese (Costa, 2009).

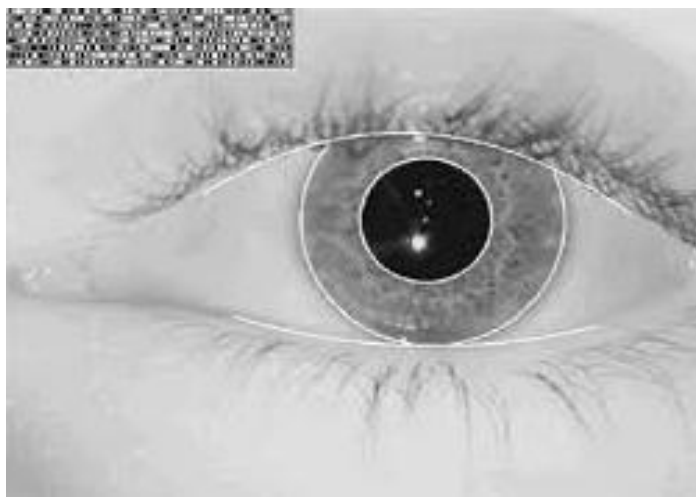


Figura 2.3. Biometria da Íris (Costa *et al.*, 2006).

Como desvantagens esta biometria tem o custo do equipamento de captura, e a rejeição pelos usuários. Há facilidade de fraude, pois o sistema não diferencia a imagem do olho humano de uma fotografia, ou de uma filmagem.

O método da biometria da íris inicia-se com a localização do disco colorido que chamamos de íris, circulado na Figura 2.3. A pupila é o orifício no centro da íris e por não ter estrutura interna, facilita sua localização. Realizada a localização da pupila, o processo poderá identificar a íris e extrair assim uma faixa contínua da qual são extraídos descritores baseados em textura e pigmentação. A alta diversidade das formas encontradas na pigmentação da íris fornece ao método uma ótima precisão (Jain *et al.*, 2000), conforme classificada na Tabela 2.1. Em outras palavras, é um dos melhores métodos quanto ao requisito da unicidade.

Biometria da Face

O próximo exemplo na Tabela 2.1 é o reconhecimento facial (Amaral *et al.*, 2011). Este é um método de boa aceitação pelos usuários, principalmente por ser similar ao reconhecimento

realizado pelo ser humano. Quanto às questões de desvantagem na dificuldade de adquirir e na baixa precisão, comentados na Tabela 2.1, são pontos que a tecnologia vem superando, pois os dispositivos de obtenção de imagens são cada vez mais precisos, fornecendo imagens de boa qualidade para os sistemas biométricos. Com a popularização da fotografia, bons equipamentos para a coleta da imagem da face encontram-se cada vez mais acessíveis atualmente.

No método de reconhecimento facial, uma das abordagens utilizadas é a global. Esta abordagem busca a solução do problema através do processamento de toda a imagem. Uma das abordagens globais é conhecida como *Eigenfaces*, Turk *et al.* (1991a) e Turk *et al.* (1991b); falaremos sobre esta abordagem no Capítulo 3.

Na Figura 2.4 podemos ver alguns exemplos de *Eigenfaces* que foram extraídas de imagens da face.

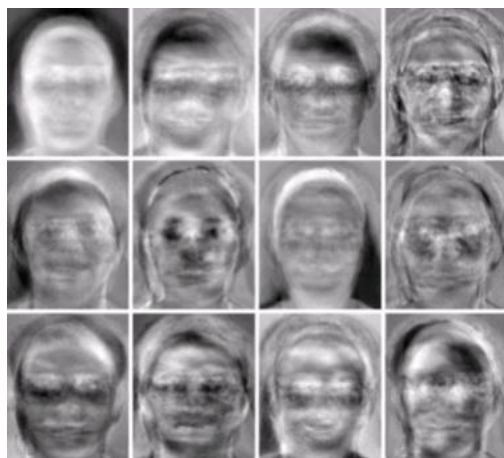


Figura 2.4. *Eigenfaces* (Silva, 2013).

Outra abordagem também utilizada na biometria da face é conhecida como local. Na abordagem local podem-se obter certos pontos da face e destes obter características específicas da sua biometria. Este método utiliza fundamentalmente textura para definir características biométricas como por exemplo, a textura próxima ao nariz, olhos ou outro ponto da face que possa fornecer informações para usar no processo biométrico. Com estas características constroem-se os descritores.

Entre as técnicas locais mais empregadas encontramos o *Elastic Bunch Graph Method* (EBGM), Wiskott *et al.* (1997) e Bonato *et al.* (2010). Na EBGM, mostrada na Figura 2.5, em

um conjunto de pontos da face são extraídos descritores de textura, baseados nos Filtros de Gabor (Lu, 2003).

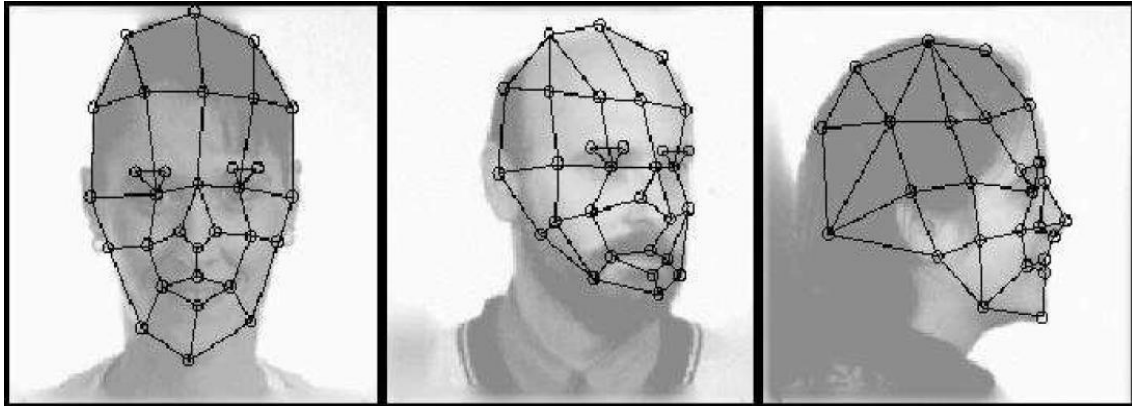


Figura 2.5. *Elastic Bunch Graph Method* (Lu, X, 2003).

Nosso trabalho tem como um dos objetivos testar o classificador Floresta de Caminhos Ótimos (OPF) aplicado a biometria facial.

Biometria da Voz

A Tabela 2.1 continua com o reconhecimento biométrico da voz (Petry *et al.*, 2004). Esta biometria possui várias características especiais que podem ser utilizadas para identificar um indivíduo. A voz humana é criada quando o ar passa pelas cordas vocais e as características físicas das cordas vocais, assim como da forma da cavidade bucal e do nariz interferem no som produzido. Por estas características serem diferentes nos seres humanos é que cada indivíduo possui uma voz peculiar.

Uma técnica utilizada para a biometria da voz é do reconhecimento do som propriamente dito. Nesta técnica a característica utilizada para o reconhecimento biométrico é uma característica física e pode ser usada a frequência do som, por exemplo. Outras técnicas no reconhecimento da voz são dos sistemas que identificam o locutor. Estas técnicas podem utilizar uma característica comportamental como a maneira da pessoa falar, sotaque que o indivíduo usa, e outras características.

Módulos de um Sistema Biométrico

Genericamente, um sistema biométrico possui cinco módulos principais (Jain *et al.*, 2007).

- Módulo Sensor: interface entre o homem e o sistema biométrico;
- Módulo de Avaliação da Qualidade: Aprova a qualidade da amostra;
- Módulo de Extração de Características: Extrai as características da amostra aprovada;
- Módulo de Tomada de Decisão: compara dois descritores e decide o resultado da comparação biométrica;
- Módulo de Banco de Dados: armazena informações sobre o usuário e informações da biometria.

Módulo Sensor

O módulo sensor em um sistema de reconhecimento digital, mostrado como exemplo na Figura 2.6, é o dispositivo onde se coleta da amostra biométrica do usuário do sistema. Dependendo do tipo de biometria podem existir módulos sensores diferentes. No caso da impressão digital é um leitor especializado. No caso da biometria facial, o módulo sensor é geralmente uma câmera que captura a imagem da face do indivíduo.

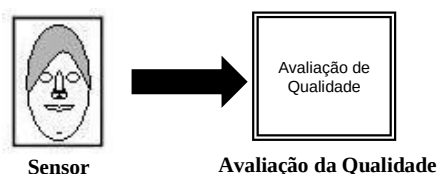


Figura 2.6. Módulo Sensor. Realiza a aquisição da biometria.

Módulo Avaliação da Qualidade

O próximo módulo é o da Avaliação da Qualidade mostrado na Figura 2.7. Este módulo tem a função de avaliar a qualidade da amostra biométrica capturada. Se este módulo avalia que a qualidade da amostra é inferior à exigida pelo sistema biométrico, esta amostra é descartada e se solicita uma nova captura ao indivíduo.

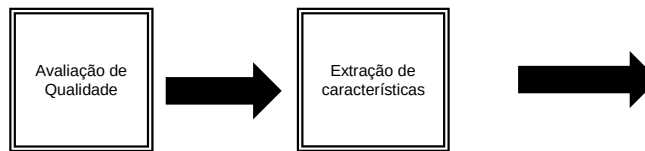


Figura 2.7. Módulo de Avaliação da Qualidade e Extração de Características.

Módulo de Extração de Características

Quando a amostra é considerada adequada, como vemos na Figura 2.7, o Módulo de Extração de Características produz um descritor com as propriedades ou características biométricas extraídas, em geral em forma de vetor de características, geralmente denominado *template*.

Os vetores de características ou descritores biométricos, são as estruturas que o sistema biométrico armazena e utiliza para identificar uma pessoa. Bonato *et al.* (2010) e Jain *et al.* (2007). Dependendo do nível de segurança desejado, este descritor pode ser criptografado.

Módulo de Tomada de Decisão

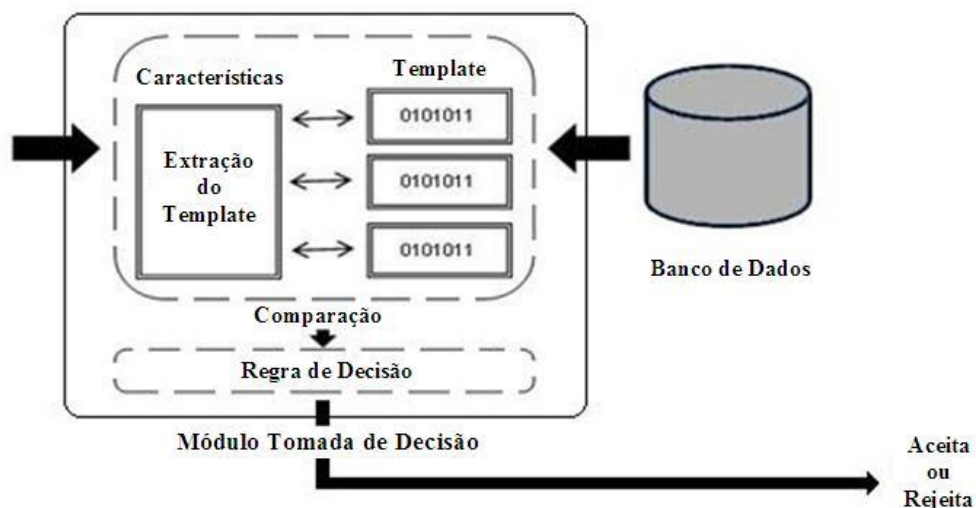


Figura 2.8. Módulo de Correspondência e Tomada de Decisão. Compara os descritores e retorna um valor de similaridade.

O Módulo de Tomada de Decisão tem como função comparar dois descritores biométricos e tomar a decisão se correspondem ou não ao mesmo indivíduo, isto é, se a

comparação entre eles é genuína ou impostora. A Figura 2.8 mostra um típico módulo de decisão que compara dois descritores, um obtido da captura de uma imagem de consulta e outro armazenado em um Banco de Dados.

O Módulo Tomada de Decisão recebe a informação biométrica (*template*) a ser identificada e compara com todos os *templates* do banco de dados. O resultado desta comparação é utilizado para que uma Regra de Decisão conclua se os *templates* são do mesmo indivíduo (Aceita), ou conclui que não pertencem ao mesmo indivíduo (Rejeita).

Geralmente em um sistema biométrico, o algoritmo de comparação gera um valor de similaridade que é comparado com um limiar de decisão. Se o valor de similaridade é maior que o limiar o sistema decide que a comparação é genuína. Em caso contrario decide que a comparação é impostora. Este limiar de decisão pode ser ajustado o que significa que a decisão tomada pelo sistema depende do valor deste limiar.

Módulo de Banco de Dados

O Módulo de Banco de Dados, mostrado na Figura 2.8, tem a função de armazenar as informações dos usuários do sistema biométrico. Em uma abordagem convencional, o módulo de bancos de dados armazena os descritores biométricos extraídos no módulo de extração de características. No entanto, em outras abordagem o banco de dados pode armazenar as amostras biométricas originalmente capturadas .

Além da informação biométrica (descritor ou amostra), este módulo também armazena outros dados que identificam o indivíduo proprietário dessa informação biométrica. Dependendo da abordagem utilizada para a captura das informações, este processo pode ser supervisionado por um ser humano ou não (Jain *et al.*, 2007).

Fases de um Sistema Biométrico

Os sistemas biométricos possuem em geral duas fases ou etapas:

- Registro;

- Reconhecimento (Verificação e Identificação).

A seguir são descritas cada uma destas fases:

Fase de Registro

Na fase de Registro, (*Enrollment*) uma nova amostra biométrica é cadastrada no Banco de Dados. Para esta fase são utilizados o Módulo Sensor, o Módulo de Avaliação da Qualidade e Extração das Características e o Módulo de Bancos de Dados (Jain *et al.*, 2007). Como ilustrado na Figura 2.9, em geral a amostra biométrica é capturada pelo módulo sensor, a seguir é verificada sua qualidade, o descritor é extraído e armazenado no banco de dados. Note-se que em alguns sistemas, a etapa de captura pode armazenar diretamente a amostra biométrica sem extrair o descritor.

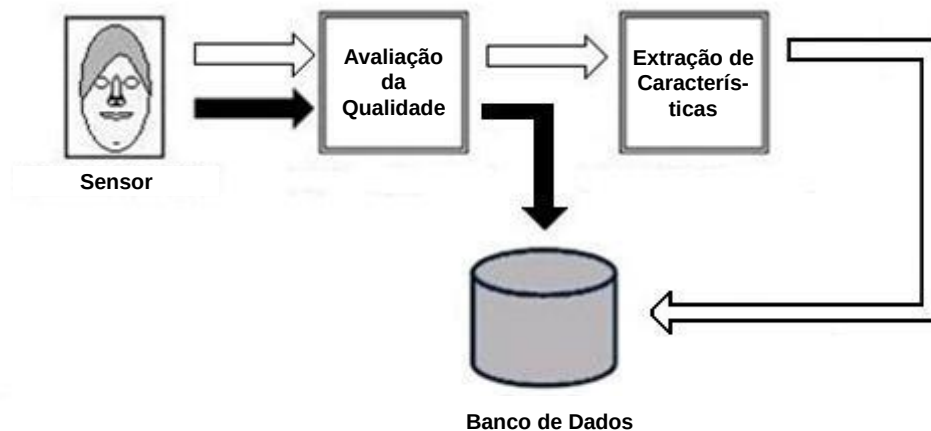


Figura 2.9. Etapa de registro.

Um sistema biométrico pode registrar uma ou várias amostras biométricas associadas ao mesmo indivíduo. Por exemplo, em um sistema de reconhecimento facial podem ser registradas varias imagens da face. O registro de várias amostras associadas ao mesmo indivíduo em geral aumenta a precisão de um sistema biométrico.

Fase de Reconhecimento

Uma vez registrada a biometria, a próxima etapa de um sistema biométrico é o reconhecimento. Existem duas formas em que o reconhecimento se apresenta em um sistema

biométrico: Verificação e Identificação.

Reconhecimento por Verificação

Em um sistema de reconhecimento por verificação o usuário declara uma identidade junto com uma amostra biométrica e o sistema verifica se essa amostra se corresponde com o registro desse usuário no Banco de Dados.

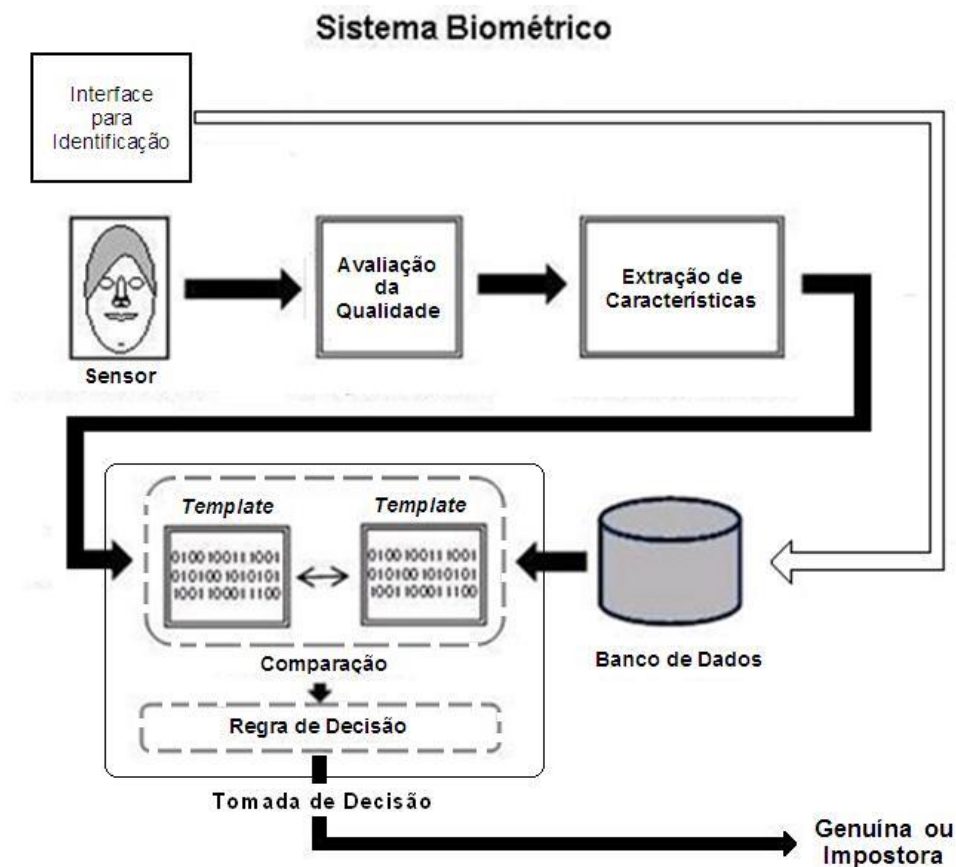


Figura 2.10. Reconhecimento por verificação.

A Figura 2.10 mostra esse processo no qual, inicialmente, a identidade pretendida é fornecida ao sistema. A partir desta informação, o sistema recupera no Banco de Dados o registro biométrico associado com essa identidade e que deve ter sido armazenado previamente na fase de registro. Na sequência, o usuário fornece uma nova amostra biométrica utilizando o Módulo Sensor. Esta amostra é processada pelos Módulos de Avaliação de Qualidade e de Extração de Características se obtendo um descritor biométrico. No próximo passo, no Módulo

de Tomada de Decisão, este descritor biométrico é comparado com o registro recuperado e é gerado um valor de similaridade. De acordo com o limiar de decisão, determina-se se a comparação é genuína ou impostora.

O reconhecimento por verificação é de um para um e dois erros podem ocorrer:

- Uma amostra genuína (amostra que pertence ao indivíduo a ser identificado) pode ser erroneamente classificada como impostora quando comparada com o registro do indivíduo.
- Uma amostra impostora (de um indivíduo diferente àquele a ser identificado) pode ser classificada como genuína quando comparada com o registro do indivíduo.

Reconhecimento por Identificação

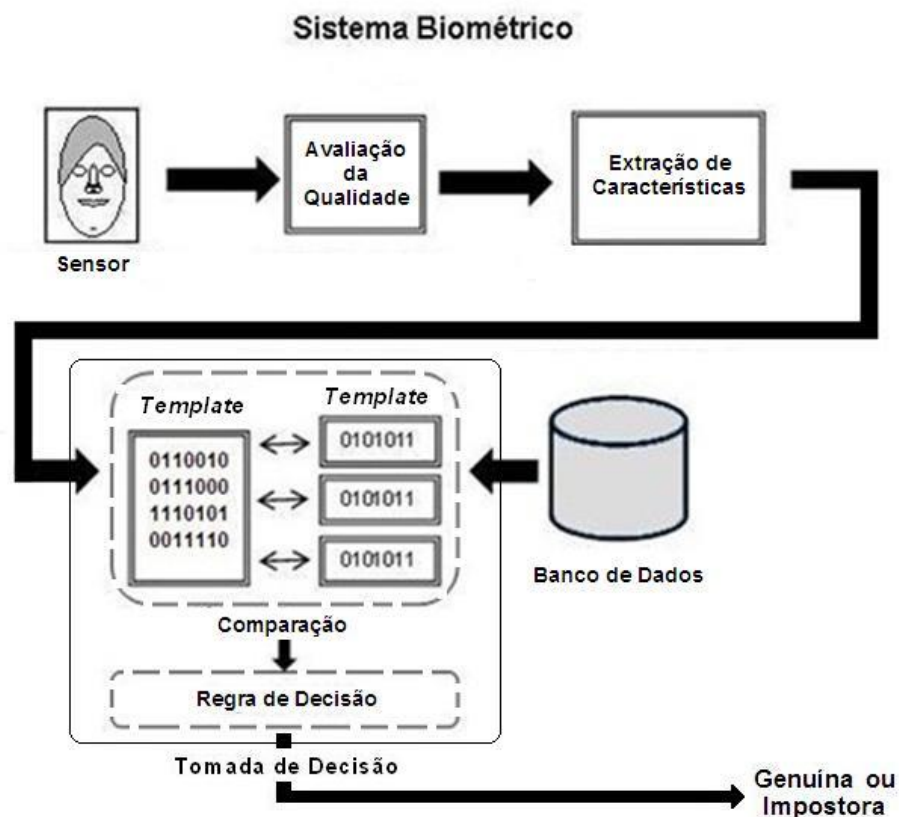


Figura 2.11. Reconhecimento por identificação.

Em um sistema de reconhecimento por Identificação ao sistema é fornecida unicamente uma amostra biométrica. Com esta amostra, o sistema tem que procurar todo o banco de dados

e decidir se existe algum registro de algum indivíduo que se corresponda com a amostra recebida e em caso positivo recuperar a identidade desse indivíduo. A Figura 2.11 mostra o processo de Identificação. Inicialmente a amostra biométrica é obtida no Módulo Sensor embora seja comum fornecer esta amostra sem a presença do indivíduo, isto é, sem a captura da amostra pelo Módulo Sensor. A seguir esta amostra é processada pelos Módulos de Avaliação de Qualidade e de Extração de Características se obtendo um descritor biométrico.

No próximo passo, no Módulo de Tomada de Decisão, este descritor biométrico é comparado com todos os registros no banco de dados. De acordo com o limiar de decisão o sistema decide se alguma das comparações realizadas é genuína, o que corresponde a identificar o indivíduo no banco de dados.

O Reconhecimento por Identificação é de um para muitos. Neste podem ocorrer dois erros:

- uma amostra genuína ser classificada erroneamente como não estando no banco de dados ou;
- uma amostra impostora, ser classificada como existente no banco de dados.

Tanto no Reconhecimento por Verificação quanto por Identificação a tomada de decisão é realizada a partir de um valor de similaridade e um limiar de decisão.

Reconhecimento por Verificação com Registro da Amostra Original

Algumas propostas de Reconhecimento facial, Moghaddam *et al.* (2000), Wang *et al.* (2003) e Chen *et al.*, (2012), consideram na etapa de registro o armazenamento da amostra original, isto é, da imagem facial e não do descritor obtido da amostra original por extração de características, ou seja, o *template*.

A Figura 2.12 mostra esse processo no qual, inicialmente, a identidade pretendida é fornecida ao sistema. A partir desta informação, o sistema recupera no banco de dados a imagem facial associada com essa identidade. Na sequência, o usuário fornece uma nova imagem facial utilizando o Módulo Sensor que é seguidamente processada pelo Módulo de avaliação de qualidade. No próximo passo, no Módulo de Extração de Características, as duas

imagens são combinadas e dessa combinação é extraído um descritor que caracteriza quão próximos são os indivíduos nas duas imagens.

O Módulo de Tomada de Decisão, ver Figura 2.12, recebe esse descritor único dessa combinação e atribui a ele um valor de similaridade. De acordo com o limiar decide-se se a comparação é genuína ou impostora.

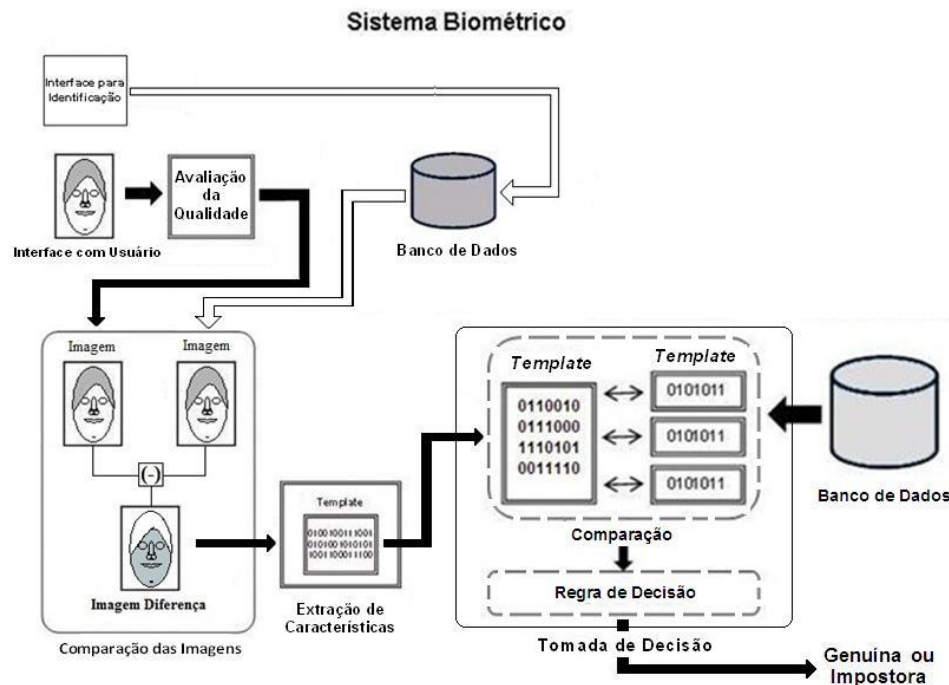


Figura 2.12. Reconhecimento por Verificação com registro de amostra original

No caso das abordagens em Moghaddam *et al.* (2000), Wang *et al.* (2003) e Chen *et al.* (2012) a combinação de imagens que é realizada é a diferença entre as duas imagens. Desta imagem diferença é extraído um descritor utilizando o Módulo de extração de características. Posteriormente, o Módulo de Tomada de Decisão utiliza um classificador que calcula um valor de similaridade deste descritor com duas classes predefinidas: a classe dos descritores de imagens diferença obtidas de imagens do mesmo indivíduo e a classe dos descritores de imagens diferença obtidas de imagens de indivíduos diferentes.

A pesquisa e os resultados desta dissertação estão orientados a um sistema de verificação deste tipo.

Algumas Métricas de Avaliação de Algoritmos Biométricos

Em geral, um algoritmo de reconhecimento por verificação biométrica não responde de forma booleana. Assim, o sistema atribui um valor de similaridade a uma comparação e neste processo é inevitável ocorrerem erros. Com as estimativas destes erros podemos determinar qual o valor de limiar de decisão que melhor descreve as necessidades de nosso sistema. O limiar de decisão é um ponto ou valor determinado, na escala de similaridade, que o algoritmo utiliza para decidir se a comparação é genuína ou impostora.

Os valores de similaridade das comparações entre as amostras ajudam a calcular algumas métricas que permitem a avaliação da qualidade de um algoritmo. Para obtermos as métricas de avaliação, utilizamos funções que estimam as taxas de erros dos algoritmos, denominadas Taxa de Falsa Aceitação (*False Match Rate*- FMR) e a Taxa de Falsa Rejeição (*False Non Match Rate*- FNMR).

Para estimar estas funções (FMR - FNMR) para um algoritmo, é selecionado um conjunto suficientemente grande e representativo de comparações genuínas e impostoras, isto é, de comparações de amostras do mesmo indivíduo e de indivíduos diferentes. A partir destas comparações, são construídas duas funções de distribuição ou histogramas: $Gen(x)$ com os valores de similaridade obtidos de comparações genuínas e $Imp(x)$ com os valores de similaridade obtidos de comparações impostoras. Estas funções descrevem como se distribuem estatisticamente os valores de similaridade de comparações genuínas $Gen(x)$ e impostoras $Imp(x)$. Idealmente os valores em $Gen(x)$ devem estar melhor distribuídos nos valores altos de similaridade e as os valores em $Imp(x)$ nos valores baixos.

A partir destas funções de distribuição $Gen(x)$ e $Imp(x)$ são construídas as funções de distribuição cumulativas que correspondem às curvas FNMR(x) e FMR(x) como podem ver na Figura 2.13.

Note-se que a função FMR(x) estima qual a taxa das comparações impostoras que o algoritmo pode erroneamente considerar como genuínas, se o limiar é fixado no valor x . Já a função FNMR(x) estima a taxa de comparações genuínas que o algoritmo retornará incorretamente como impostoras se o limiar é definido como x . Exemplificando, em um controle de acesso, o FMR(x) será a taxa de indivíduos que erradamente terão o acesso liberado

se o limiar é fixado em x e o $FNMR(x)$ será a taxa de indivíduos que erradamente terão seu acesso rejeitado se o limiar é fixado em x (Jain *et al.*, 2007).

A Figura 2.13 mostra um exemplo das funções FMR e FNMR. Observamos que o valor de FMR descreve a segurança do sistema biométrico, enquanto o FNMR define a conveniência do sistema, ou seja, evitar o incomodo de que uma comparação genuína seja classificada como impostora. Vemos que quando o valor de limiar no eixo x aumenta, o valor de FMR diminui e o valor de FNMR aumenta, isto é, quanto mais seguro o sistema também ele será menos conveniente porque rejeita mais comparações genuínas. Pelo contrario, quando o limiar x diminui, o sistema será menos seguro, mas rejeitara muito menos comparações genuínas. Os objetivos de um sistema biométrico determinam qual o valor que o limiar deve possuir.

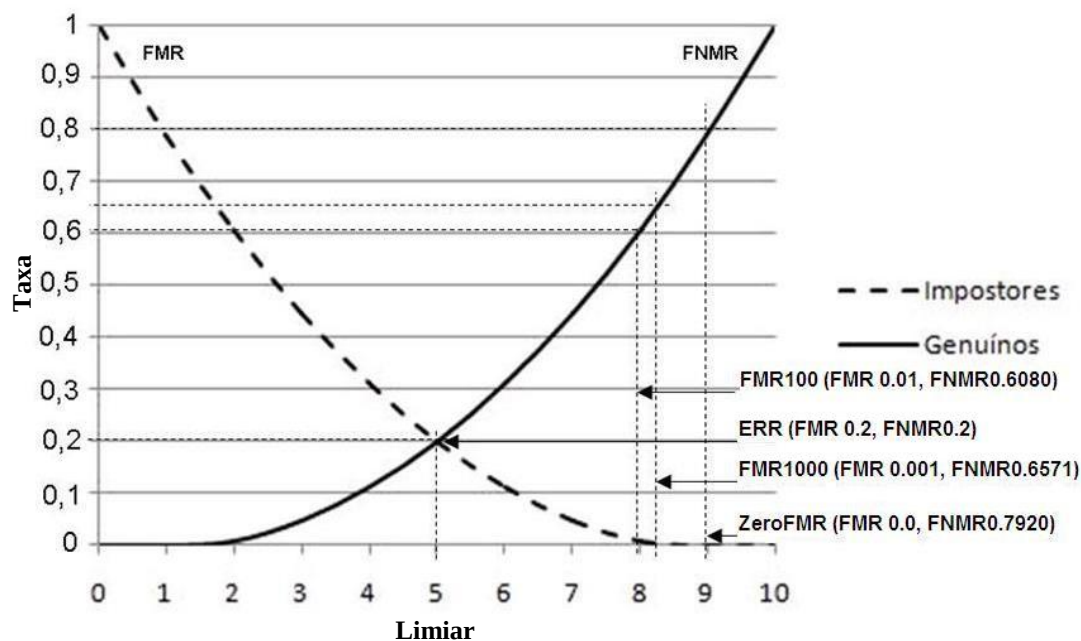


Figura 2.13. Funções FMR e FNMR de distribuição cumulativa das similaridades de comparações impostoras e genuínas. Pontos de limiar das métricas EER, FMR100, FMR1000 e FMRZero.

As principais métricas biométricas descrevem como é o comportamento simultâneo das funções FMR e FNMR. Estas métricas são as seguintes:

- **EER (Equal Error Rate):** Métrica mais utilizada para avaliar sistemas biométricos. É calculada no ponto de limiar onde as duas funções tem o mesmo valor. Descreve qual a taxa de erros esperada se você deseja a mesma taxa de falsa aceitação (FMR) e de falsa

rejeição (FNMR). Em um sistema ideal o EER seria igual à zero.

- O FMR100: Corresponde ao valor da FNMR(x) no ponto de limiar onde o valor de FMR(x) é igual a 0,01. Descreve qual a taxa esperada de falsa rejeição se o limiar é fixado em um ponto em que a taxa de falsa aceitação é 0,01, isto é, admitindo um 1% de insegurança no sistema.
- O FMR1000: Corresponde ao valor da FNMR(x) no ponto de limiar onde o valor de FMR(x) é igual a 0,001. Descreve qual a taxa esperada de falsa rejeição se o limiar é fixado em um ponto em que a taxa de falsa aceitação é 0,001, isto é, admitindo um 0,1% de insegurança no sistema.
- O FMRZero: Corresponde ao valor da FNMR(x) no ponto de limiar onde o valor de FMR(x) é igual a Zero(0). Descreve qual a taxa esperada de falsa rejeição se o limiar é fixado no ponto em que a taxa de falsa aceitação é 0, isto é, o sistema é totalmente seguro.

Em nossa pesquisa utilizamos estas métricas para avaliar o desempenho dos algoritmos biométricos.

A Figura 2.14 apresenta uma forma de representar os valores de EER e dos FMRs graficamente descrevendo unicamente os valores notáveis.

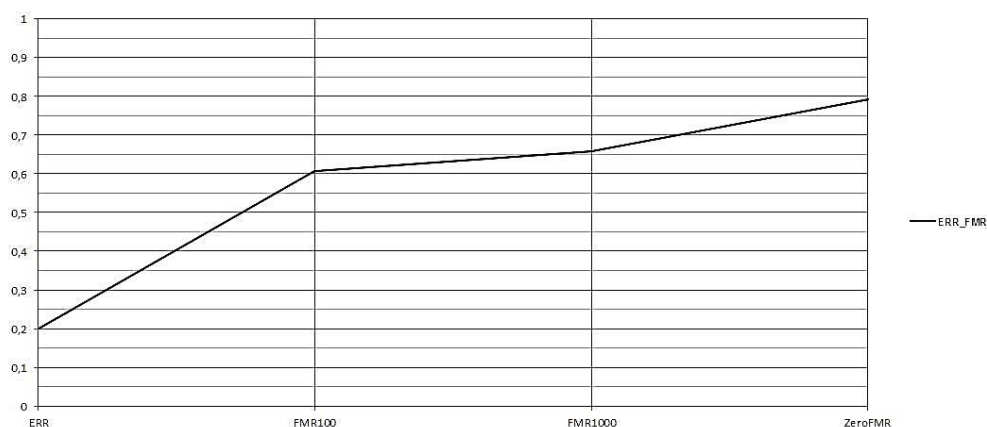


Figura 2.14. Gráfico de demonstração das taxas de erros EER, FMRs.

Existem outros nomes utilizados pela comunidade biométrica para descrever as funções de distribuição apresentadas. O FMR, ou Taxa de Falsa Aceitação, também pode ser denominado FAR, abreviação de *False Accept Rate*. Assim como taxa *False Reject Rate*, ou

Taxa de Falsa Rejeição é utilizada no lugar do FNMR.

Uma forma utilizada pela comunidade biométrica para as funções FMR e FNMR, assim como as métricas de avaliação, é através da Curva ROC (*Receiver Operating Characteristic*) (Lucas *et al.*, 2006) representada na Figura 2.15.

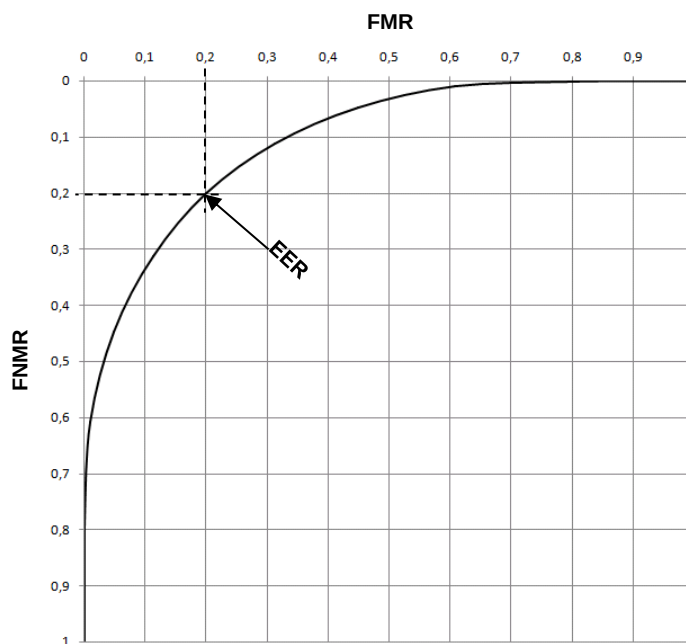


Figura 2.15. Gráfico da Curva ROC – Receiver Operating Characteristic.

A Curva ROC é formada pela pelos pares de valores formados pela combinação de valores de FMR e FNMR em cada ponto do limiar. Neste tipo de representação, a combinação de valores das funções FMR e FNMR pode ser verificada sem necessidade do valor de limiar. Por outro lado, podemos identificar claramente as métricas EER e FMRs. Este tipo de representação é muito adequado para determinar qual a combinação de taxas de erros é desejada para um sistema biométrico. Por esta razão, ela é a mais distribuída pelas empresas que vendem tecnologia biométrica para seus usuários possam personalizar seus sistemas.

Considerações Finais

Este capítulo apresentou conceitos básicos de biometria. Foram abordados os módulos de um sistema biométrico e seu uso nas fases do sistema. Adicionalmente foram caracterizados

os sistemas de reconhecimento por verificação e por identificação. Em particular, foi descrito um sistema de verificação com registro das amostras capturadas que será o tipo de sistema para o qual está orientada a pesquisa neste trabalho. Terminando o capítulo, descrevemos as métricas de avaliação que utilizaremos em nossa pesquisa, assim como suas formas de representação.

O capítulo a seguir descreve os métodos de reconhecimento facial relacionados com nossa pesquisa.

3. Métodos de Reconhecimento Facial

Este capítulo apresenta os métodos biométricos de Reconhecimento Facial. Em particular serão apresentados os métodos de tipo global *Eigenfaces*, Turk *et al.* (1991a) e Turk *et al.* (1991b) e o método de tipo local *Elastic Bunch Graph Method* ou *EBGM* (Wiskott *et al.*, 1997). O capítulo também apresenta o uso de classificadores e de classes no reconhecimento facial, e, adicionalmente, será apresentada a abordagem Bayesiana (Moghaddam *et al.*, 2000).

Reconhecimento Facial

Atualmente, não se tem conhecimento de uma técnica de reconhecimento facial automatizada que tenha total precisão no reconhecimento de todas as faces (Zhao *et al.*, 2003). O sistema de reconhecimento facial humano possui a capacidade de lembrar muitos rostos durante a vida e os identificar mesmo depois da ação dos anos ou com oclusões tais como óculos, barba ou cabelos. Apesar deste sofisticado sistema de reconhecimento facial conseguir bons resultados, não está isento de erros na identificação de um indivíduo (Tan *et al.*, 2005).

As técnicas de reconhecimento facial surgiram por volta dos anos setenta, com as técnicas iniciais de processamento das dimensões do rosto, Jain *et al.* (2007) e Zhao *et al.* (2003). Logo foram superadas por técnicas denominadas como globais e locais, Prodossimo *et al.*, (2012) e Tan *et al.*, (2006).

A abordagem global, também chamada de abordagem holística, propõe a solução do problema através do processamento de toda a imagem, extraindo informações que possibilitem o processo de identificação. Esta abordagem tenta oferecer as vantagens de ao utilizar toda a imagem como entrada e garantir um grande número de informações para o processo do reconhecimento. Podemos considerar como desvantagem a grande dimensionalidade dos vetores de características (descritores) a serem processados, o que necessariamente exige considerável capacidade computacional.

A abordagem global incentivou o estudo de métodos que propunham a redução da dimensionalidade dos vetores. Uma proposta para alcançar este objetivo está nas *Eigenfaces*, introduzidas por Turk *et al.* (1991a). Nesta proposta temos a representação da imagem em um

vetor k-dimensional obtido da Análise dos Componentes Principais (PCA) da imagem da face (Moghaddam, 1999), definindo uma base de autovetores a ser utilizada na redução da dimensionalidade.

Na busca de métodos para a redução das informações utilizadas na identificação biométrica, precisamos abordar a técnica de reconhecimento local que busca definir pontos específicos na imagem facial. Destes pontos podem ser extraídas as características relevantes. Estas características serão utilizadas para construir as informações da identificação biométrica.

Assim, neste capítulo também comentaremos sobre a Elastic Bunch Graph Method (*EBGM*) que é uma técnica local bastante utilizada. Esta técnica oferece a vantagem de em alguns casos poder reduzir o tamanho do vetor de característica e, principalmente, possuir a condição de identificar uma imagem processando apenas algumas partes das informações contidas na face.

Em nosso trabalho utilizaremos o método global *Eigenfaces* e como método local o *EBGM* descritos a seguir.

Eigenfaces

O método *Eigenfaces* é uma abordagem global (Turk *et al.*, 1991a). Esta abordagem utiliza o conceito de Análise de Componentes Principais (PCA). O PCA é uma técnica estatística que realiza uma transformação linear em um conjunto de dados multidimensionais. Deste modo, no novo sistema de coordenadas, os componentes mais relevantes estão nas primeiras dimensões ou eixos, isto é, a maior variância dos dados projetados fica ao longo da primeira coordenada. A segunda maior variância fica ao longo da segunda coordenada, e assim por diante. Estas primeiras dimensões são chamadas de componentes principais.

A matriz de transformação utilizada tem como colunas os autovetores da matriz de covariância estimada de um conjunto de dados de treinamento. Esta técnica é utilizada para reduzir a dimensionalidade dos dados, já que uma vez que os dados são transformados pela matriz, garante-se que as primeiras dimensões desses vetores possuam a maior parte da informação relevante para a discriminação dos dados. Desta forma, a análise dos dados pode

ser realizada unicamente nas dimensões principais, podendo descartar grande parte das outras dimensões.

A abordagem de *Eigenfaces* considera cada imagem como um vetor, isto é, as linhas da matriz $n \times m$ de pixels da imagem são unidas em um vetor de $n \times m$ coordenadas, o que nos fornece um espaço de faces R^{nm} . Desta forma, um conjunto de imagens de treinamento é considerado como um conjunto de vetores e sobre este conjunto de vetores é aplicado o método de Componentes Principais. Como resultado, quando a matriz de covariância é calculada para estes vetores (imagens), os autovetores obtidos nas colunas são realmente um conjunto de imagens chamadas de *Eigenfaces*. A Figura 3.1 contém um exemplo de 15 *Eigenfaces* dos 15 componentes principais para um determinado conjunto de imagens faciais de treinamento. Este método permite uma importante redução de dimensionalidade, por exemplo, uma imagem de 64×64 pixels corresponderia a um vetor de dimensão 4096, mas se pode representar a informação mais relevante de todas elas em algumas dezenas de dimensões.

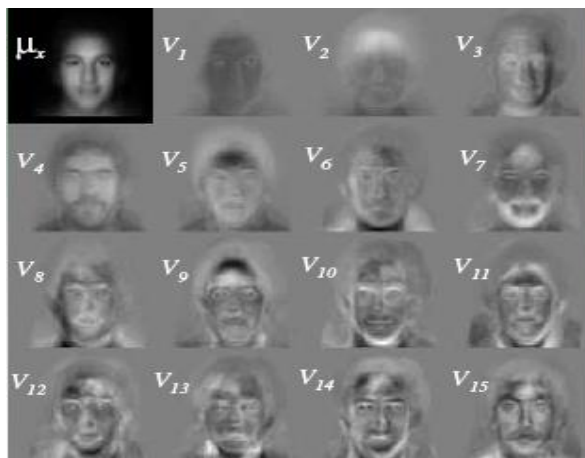


Figura 3.1. Eigenfaces.

Com a matriz de transformação construída, uma nova imagem (vetor) pode ser transformada e projetada nas dimensões principais. Os coeficientes da projeção da imagem transformada em cada um dos componentes principais são organizados em um vetor. Este vetor corresponde ao descritor ou vetor de características para essa imagem. Note-se que a combinação linear destes coeficientes com os autovetores dos componentes principais, gera uma imagem aproximada da imagem original (Moghaddam, 1999).

A hipótese desta técnica é que imagens da face de um mesmo indivíduo devem ter coeficientes similares quando transformadas e, portanto, a distância entre seus vetores de coeficientes deve ser menor que entre vetores de indivíduos diferentes. A similaridade de duas imagens é calculada através da distância euclidiana destes vetores.

Para exemplificar o processo do cálculo das *Eigenfaces*, consideraremos um conjunto com uma quantidade x de imagens de pessoas diferentes e um número y de imagens para cada indivíduo, totalizando, assim, x vezes y imagens. Em seguida calculamos a média das faces, através da equação 1 (Silva, 2013):

$$\Psi = \frac{1}{M} \sum_{i=1}^M \Gamma_i \quad (1)$$

Na equação 1, o valor de M é a quantidade de imagens (x vezes y) e Γ é o vetor resultante da matriz $n \times m$ (pixels) de cada imagem. O resultado da função é expresso por Ψ , que representa a face média. A equação 2, a seguir, calcula o valor de Φ_i . Esta função resulta na diferença entre a face média (Ψ) e os vetores Γ_i .

$$\Phi_i = \Gamma_i - \Psi \quad (2)$$

Com os vetores das diferenças Φ_i podemos construir a matriz A , em que cada coluna contém os vetores Φ_i . A matriz de covariância é definida da seguinte forma:

$$C = AA^T \quad (3)$$

Onde C é a matriz de covariância e A é a matriz resultante dos valores de Φ . O resultado de AA^T torna-se inviável para computar, assim para contornar este problema utilizaremos $A^T A$, reduzindo seus autovetores v , para um número N de dimensões que venha a viabilizar o processo computacional, como mostrado na equação 4.

$$A^T A v_i = \lambda v_i \quad (4)$$

Multiplicando por A , obteremos as *Eigenfaces* $\lambda A v_i$ associadas aos autovetores $C = AA^T$, conforme pode ser visto em Silva (Silva, 2013) e também expresso na equação 5, a seguir.

$$AA^T A v_i = \lambda A v_i \quad (5)$$

Na equação 5 temos $A v_i$, que são os autovetores da matriz de covariância AA^T , associados as N maiores autovalores, onde N é o número de dimensões que foram reduzidas.

Calculados os autovetores, é preciso construir o “espaço de faces” utilizando as *Eigenfaces* calculadas. Isto pode ser feito com a equação 6.

$$\omega_k = u_k^T \Phi \quad (6)$$

Com a equação 6 formamos o vetor $[\omega_1, \omega_2, \dots, \omega_O]$, sendo O a quantidade de *Eigenfaces*.

$$Y^T = [\omega_1, \omega_2, \dots, \omega_O]$$

Para reconhecimento de uma nova imagem, calcula-se o vetor Φ (diferença) da nova imagem e, obtemos a projeção Y no “espaço de faces” com o vetor Φ da nova imagem. Depois se calcula a menor distância (d) entre o vetor de diferença da nova imagem com a média (Ψ), e os vetores ω do “espaço de faces”, como mostrado na equação 7.

$$d = \min_k ||Y - Y_O|| \quad (7)$$

Se a menor distância encontrada for menor que o limiar pré-estabelecido, então a nova imagem pertence a imagem qual foi encontrada a distância d . Caso contrário a nova imagem é desconhecida.

Elastic Bunch Graph Method (EBGM)

Dentre as técnicas locais mais utilizadas encontramos o *Elastic Bunch Graph Method* (EBGM). Nesta técnica, de um conjunto de pontos da face são extraído descritores de textura baseados nos Filtros de Gabor (Wang *et al.*, 2003).

A hipótese deste método é que em faces do mesmo indivíduo a descrição da textura deve ser similar nos mesmos pontos da face (Wiskott *et al.*, 1997).

O método *EBGM* compreende etapas de treinamento e reconhecimento. Na etapa do treinamento é criada uma estrutura utilizada para localizar os pontos de interesse. Estes pontos, em nosso exemplo da Figura 3.2, são regiões características da face e são predeterminados.

Alguns dos pontos que podem ser utilizados são: centro dos olhos, centro da boca, topo da cabeça, como podemos ver alguns exemplos na Figura 3.2a.

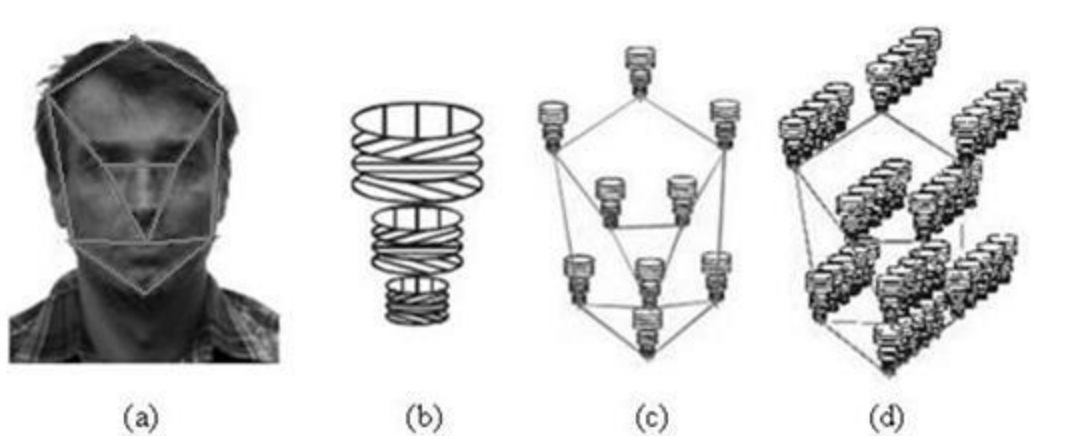


Figura 3.2. (a) Grafo da face; (b) Jet; (c) Image Graph; e (d) Bunch Graph. (Wiskott et al., 1997).

O método é treinado utilizando um conjunto de faces representativas da população a ser reconhecida. Para cada face, os pontos de interesse são marcados manualmente. Estes pontos são conectados por um conjunto de arestas predefinidas, formando um grafo, como ilustrado na Figura 3.2a. Cada um destes pontos é convoluído com um banco de Filtros de Gabor, com diferentes frequências e direções, como demonstrado na Figura 3.2b, criando assim um descritor de cada ponto. O conjunto de coeficientes obtidos da convolução em cada ponto é denominado *Jet*. O grafo de uma imagem visto na Figura 3.2c é dado pelo conjunto dos *Jets* de cada vértice e pelo tamanho de cada aresta. A união dos *Jets* em cada vértice e o tamanho médio de cada aresta do conjunto de grafos de uma imagem das faces utilizadas forma o *Bunch Graph* (Figura 3.2d).

O reconhecimento pelo método *EBGM* é dividido em duas atividades: cadastramento e casamento. Em ambas as atividades são obtidas as coordenadas dos olhos da imagem facial capturada. Faz-se uma transformação, a fim de alinhar a face, nivelando os olhos nas coordenadas predefinidas no *Bunch Graph* de treinamento.

Posteriormente, ocorre o processo de localização dos pontos de interesse da nova imagem. Partindo da posição relativa dada pelas arestas do *Bunch Graph*, a vizinhança de cada vértice da nova imagem é explorada, procurando o ponto de maior semelhança com um dos *Jets* do vértice correspondente no *Bunch Graph*, mostrado na Figura 3.3.



Figura 3.3. (a) grafo de imagem de entrada; e (b) Bunch Graph, os *Jets* em cinza são os que possuem maior semelhança com *Jets* do grafo de imagem de entrada. (Wiskott et al., 1997).

Após a localização dos pontos de interesse, é construído o grafo da nova imagem. No cadastramento o grafo da imagem é armazenado junto com uma identidade fornecida pelo indivíduo.

Um descritor típico deste método é formado pelos *Jets* de cada imagem. Para calcular a semelhança entre as faces são utilizadas a textura dada pelos coeficientes dos vértices, além da morfologia dada pelos ângulos e comprimentos das arestas. Geralmente, como função de similaridade é calculada a distância euclidiana entre os coeficientes dos *Jets*.

Até este ponto do capítulo, comentamos sobre o Reconhecimento Facial, que é a biometria utilizada em nossas pesquisas. Também falamos sobre os métodos de extração de características biométricas como o *Eigenfaces*. O segundo método de extração que comentamos e usaremos em nosso trabalho é o Elastic Bunch Graph Method (*EBGM*).

Em nosso trabalho de pesquisa utilizaremos os métodos de extração comentados, para obter material para testes. Nas seções seguintes deste capítulo, abordaremos assuntos referentes aos Classificadores de Padrões, a quantidade de classes utilizadas pelo classificador e sobre a Abordagem Bayesiana na verificação e identificação de imagens.

Reconhecimento Facial Utilizando Classificadores

Algoritmos classificadores de padrões recebem como entrada um conjunto de vetores de

características. Uma das funções do classificador é organizar estes vetores em um tipo de estrutura que seja útil à classificação. Com os vetores organizados, os algoritmos podem classificar novos vetores de entrada.

Os algoritmos classificadores podem trabalhar com dois tipos de abordagens: a supervisionada e a não supervisionada. Na abordagem não supervisionada, não há conhecimento prévio de informações referentes às classes. O algoritmo classificador recebe apenas o conjunto de vetores de características, sem nenhuma informação de como estes vetores serão organizados em grupos, ou seja, em classes. Na abordagem não supervisionada, as informações de como dividir os vetores de características em classes, fazem parte do método. O algoritmo definirá as classes conforme as características que considerar mais adequadas para a divisão do conjunto de vetores (Alpadin, 2010).

Outra abordagem é a supervisionada. Nesta abordagem são fornecidas informações prévias a respeito das classes. Com base nas definições destas classes é que o algoritmo classificador realiza o processo de aprendizado. Em outras palavras, o algoritmo classificador recebe o conjunto de vetores de características, juntamente com as informações dos vetores que compõem as classes.

Um algoritmo de classificação de padrão supervisionado pode ser utilizado para o reconhecimento facial. Assim, o classificador recebe como entrada um conjunto de amostras biométricas e as organiza em grupos ou classes, conforme as informações recebidas com o conjunto de descritores. Na fase de aprendizado do algoritmo, que chamaremos fase de treinamento, as amostras biométricas são avaliadas e organizadas em conjuntos. Com esta organização são extraídas informações. Com estas informações é possível formar uma estrutura, que chamamos de OPF (*Optimum-Path Forest* ou Floresta de Caminhos Ótimos). Esta estrutura será utilizada para classificar novas amostras biométricas.

Em nosso trabalho utilizaremos uma fase de classificação, onde o classificador recebe um segundo conjunto de descritores, denominado conjunto de classificação. Preferencialmente o conjunto de classificação possui amostras diferentes das utilizadas como entrada na fase de treinamento. Comparando as amostras de classificação com as classes do conjunto de treinamento, o sistema retorna um valor de distância. Este valor representa a proximidade da

amostra de classificação com as classes do conjunto de treinamento. Conhecendo os valores de distância da amostra de classificação com as classes de treinamento, o algoritmo classificador pode classificar as amostras biométricas do conjunto de classificação.

Existem vários classificadores de padrões atualmente que podem ser utilizados na biometria como, por exemplo, Máquina de Vetores de Suporte (SVM), Redes Neurais Artificiais (ANN-MLP) ou Floresta de Caminhos Ótimos (OPF), Papa (2008), Papa *et al.* (2009a), Papa *et al.* (2009b), Papa *et al.* (2012) e Marana *et al.*, (2009). Destes classificadores falaremos apenas do OPF em nosso trabalho, pois é o algoritmo que queremos analisar em nossa pesquisa.

Quantidades de Classes do Classificador

A abordagem convencional do uso de classificadores em reconhecimento biométrico define uma classe para cada indivíduo, Afonso *et al.* (2012) e Papa *et al.* (2009a). Dado um novo vetor de característica, o classificador deve determinar a que classe ele pertence, isto é, a que indivíduo pertence o vetor de características.

Esta abordagem tem duas desvantagens:

- Para definir cada classe precisamos um de conjunto grande de amostras de cada indivíduo.
- Quando um novo indivíduo é cadastrado, o classificador precisa ser recalculado.

No método em que o sistema utiliza diversas classes para solucionar o problema do reconhecimento facial, tem a necessidade de que a cada nova amostra inserida no conjunto de treinamento uma nova fase de treinamento seja processada.

Com este processo de um novo treinamento o conjunto das classes é recalculado, para que todas as classes sejam atualizadas.

Abordagem Bayesiana

A abordagem bayesiana está direcionada a sistemas de verificação e identificação que

armazenam a imagem facial propriamente, e não descritores. Nesta abordagem as imagens são integradas através da sua imagem diferença.

Parte-se da hipótese de que imagens obtidas como diferença de duas imagens faciais possuem padrões diferentes se pertencem ao mesmo indivíduo ou a indivíduos diferentes. Seguindo esta hipótese, das imagens diferença são extraídos descritores utilizando métodos de extração de características como *Eigenfaces* e *EBGM* descritos anteriormente.

O uso de classificadores (Li *et al.*, 2004) nesta abordagem está direcionado a organizar o conjunto de treinamento em duas classes de descritores. Uma primeira classe com os descritores obtidos de imagens diferença das imagens do mesmo indivíduo que é denominada de *Classe Intrapessoal*. A segunda classe é formada pelos descritores obtidos de imagens diferença de indivíduos diferentes, denominada de *Classe Interpessoal*. Essas classes servem para treinar um Classificador Bayesiano que dá o nome ao método.

O método parte da criação de uma imagem diferença. A equação 6 mostra essa diferença subtraindo as imagens I_1 e I_2 . Esta operação é realizada com a subtração de pixel a pixel das imagens.

$$\Delta = I_1 - I_2, \quad (6)$$

Na proposta de Moghaddam et al. (2000), o Classificador Bayesiano para o reconhecimento facial utiliza uma medida de similaridade probabilística, que é baseada na intensidade da imagem. Podemos observar que na equação 6 temos Δ , que é a variação da intensidade da imagem, I_1 é a representação da primeira imagem e I_2 a representação da segunda imagem.

A equação 7 mostra um valor de probabilidade de um descritor diferença, ou seja, o Δ calculado na equação 6 estar para a classe intrapessoal (Ω_I), e estar para a classe interpessoal (Ω_E).

$$P(\Delta|\Omega_I) \text{ e } P(\Delta|\Omega_E), \quad (7)$$

Assim, a medida da similaridade pode ser definida como indicado na equação 8.

$$S(I_1, I_2) = P(\Delta \in \Omega_I) = P(\Omega_I | \Delta) \quad (8)$$

Significando $S(I_1, I_2)$ o valor da similaridade entre a imagem I_1 e a imagem I_2 . Já $P(\Omega_I | \Delta)$ caracteriza a verossimilhança, ou seja, a probabilidade da variação Δ pertencer à classe intrapessoal Ω_I . Já $P(\Omega_E | \Delta)$ caracteriza a probabilidade da variação Δ pertencer à classe interpessoal Ω_E .

Utilizando a regra de *Bayes* calculamos a similaridade entre as imagens I_1 e I_2 , classificando a verossimilhança da imagem:

$$S(I_1, I_2) = \frac{P(\Delta | \Omega_I)}{P(\Delta | \Omega_I) + P(\Delta | \Omega_E)} \quad (9)$$

A probabilidade pode ser calculada conhecendo o número de amostras de cada classe do problema e o total de amostras da base de dados. Considerando um problema de reconhecimento de padrões com duas classes e que o número de amostras das duas classes seja igual, saberemos que $P(\Omega_I) = P(\Omega_E) = 1/2$. Logo, para decidir se a imagem pertence à classe intrapessoal, classe das imagens dos mesmos indivíduos, ou a interpessoal, classe de imagens de indivíduos diferentes, o algoritmo utiliza a regra da máxima a priori. Esta regra determina que duas imagens sejam do mesmo indivíduo se o resultado da equação 10 for verdadeiro.

$$S(I_1, I_2) > 1/2. \quad (10)$$

A proposta de Moghaddam calcula a função de probabilidade a partir dos descritores do método *Eigenfaces*, aplicados sobre a imagem diferença.

O uso de um Classificador Bayesiano em vetores de alta dimensionalidade, como é o caso de descritores biométricos, impõe um alto custo computacional, tanto no treinamento do classificador como no processo de classificação, Moghaddam *et al.* (1997) e Theodoridis *et al.* (2009). Para diminuir este custo os autores, Moghaddam *et al.* (1997), propõem um método de cálculo aproximado.

Adicionalmente, um Classificador Bayesiano tem um bom comportamento se a classe intrapessoal, descritores genuínos, e classe interpessoal, descritores impostores, se distribuem

de forma normal e linearmente separáveis (Moghaddam *et al.*, 2000).

Propostas similares utilizando como extração e característica o método do *EBGM* foram apresentadas em Wang *et al.* (2003) e Li *et al.* (2004). Recentemente em Chen *et al.* (2012) novas reformulações das abordagens têm sido realizadas obtendo resultados mais precisos.

Considerações Finais

O Capítulo 3 mostrou os métodos de Reconhecimento Facial que são de interesse desta pesquisa. Para os testes que serão realizados em nosso trabalho utilizaremos descritores obtidos pelos métodos de reconhecimento facial global, *Eigenfaces*, e de reconhecimento local *Elastic Bunch Graph Method (EBGM)*. Foi apresentado, ainda, o método de Reconhecimento Fácil Bayesiano por ter apresentado pela primeira vez a proposta com o uso de duas classes, que comentamos anteriormente como Classe Intrapessoal e Classe Interpessoal (Moghaddan *et al.*, 2000).

O próximo capítulo apresenta o classificador supervisionado Floresta de Caminhos Ótimos ou OPF que foi utilizado em nossa pesquisa.

4. Classificador Floresta de Caminhos Ótimos

Este capítulo apresenta o classificador Floresta de Caminhos Ótimos (OPF) em sua versão original assim como sua abordagem melhorada com o objetivo de aprimorar seu desempenho computacional. Adicionalmente será apresentada uma extensão do classificador para seu uso com algoritmos de votação.

Classificador Floresta de Caminhos Ótimos

O Classificador Floresta de Caminhos Ótimos (OPF) apresenta uma proposta para a classificação supervisionada de padrões utilizando conceitos da teoria dos grafos. Esta abordagem surgiu como uma generalização da Transformada de Florestas da Imagem (IDT) (Falcão *et al.*, 2004).

O método Floresta de Caminhos Ótimos (OPF), Papa *et al.* (2008) e Papa *et al.* (2009a), é um algoritmo classificador supervisionado baseado em grafos que vem produzindo bons resultados em aplicações para outros problemas, Alpaydin (2010) e Afonso *et al.* (2012). Neste classificador é criada uma floresta onde os nós são descritores. Para estes nós é considerada uma relação de adjacência completa, e utilizada uma função de distância para definir os relacionamentos dos nós. Na floresta são classificadas árvores de forma tal, que cada árvore está associada a uma classe. Uma mesma classe pode ser representada por mais de uma árvore.

Esta forma de construção do classificador permite que possam ser representadas classes não linearmente separáveis, com amostras espacialmente dispersas (Alpaydin, 2010). O algoritmo OPF foi utilizado também em Rocha *et al.* (2009) e Cappabianco *et al.* (2012) como base de um algoritmo classificador não supervisionado.

Todo classificador supervisionado é construído a partir de uma fase inicial de treinamento, na qual é fornecido um conjunto de amostras representativas de cada uma das classes que se pretende caracterizar. Posteriormente, o classificador deve ser capaz de determinar a qual classe pertence uma nova amostra apresentada.

No caso do classificador OPF, o conjunto de amostras de treinamento é modelado através de um grafo rotulado no qual cada amostra é representada por um vértice. Estes vértices

são organizados em uma floresta, isto é, um conjunto de árvores. Cada árvore da floresta contém unicamente vértices de uma mesma classe, sendo que para uma mesma classe podem existir varias árvores. Desta forma, cada classe é caracterizada por conjunto de árvores dentro da floresta construída.

Para gerar o conjunto de árvores representativas de cada classe, este método aplica um processo de maximização do mapa de conectividade entre amostras da mesma classe. Neste processo, cada vértice recebe um valor que caracteriza o custo da sua conectividade com seu grupo. Este valor está associado ao caminho de menor custo desse vértice com algum vértice protótipo da sua árvore. Os vértices protótipos são aqueles dentro da árvore que são os mais próximos aos vértices de outra classe.

Uma das propriedades importantes deste classificador é que uma classe é caracterizada por estruturas sem restrições espaciais. O conjunto de árvores de uma classe pode estar espacialmente separado, isto é, não formando um único agrupamento espacial. Este conjunto de árvores de uma classe também pode ter formatos diferentes, ou seja, as diferentes classes não precisam, por exemplo, ser linearmente separáveis.

Para classificar uma nova amostra, procura-se em toda a floresta o vértice para o qual o custo de conectividade com a nova amostra seja menor. A classe desse vértice é atribuído à nova amostra.

A seguir apresentamos em detalhe os algoritmo de treinamento e o algoritmo classificação do OPF.

Algoritmo de Treinamento

A Figura 4.1 mostra o algoritmo *OPF_Treinamento()*, que utilizamos para construir o classificador Floresta de Caminhos Ótimos (OPF) apresentado em Papa *et al.* (2012). O algoritmo recebe como entrada um conjunto de amostras Z_1 , cada uma delas associada a um conjunto de classes λ e um par de vetores de características (v, d) para computar a distância.

O primeiro passo do *OPF_Treinamento()* é construir um grafo completo com cada descritor associado a um vértice e uma aresta entre dois vértices quaisquer como mostrado na

Figura 4.3. Cada aresta recebe um valor com a distância entre as amostras dos vértices adjacentes a ela segundo a função de distância *dist*.

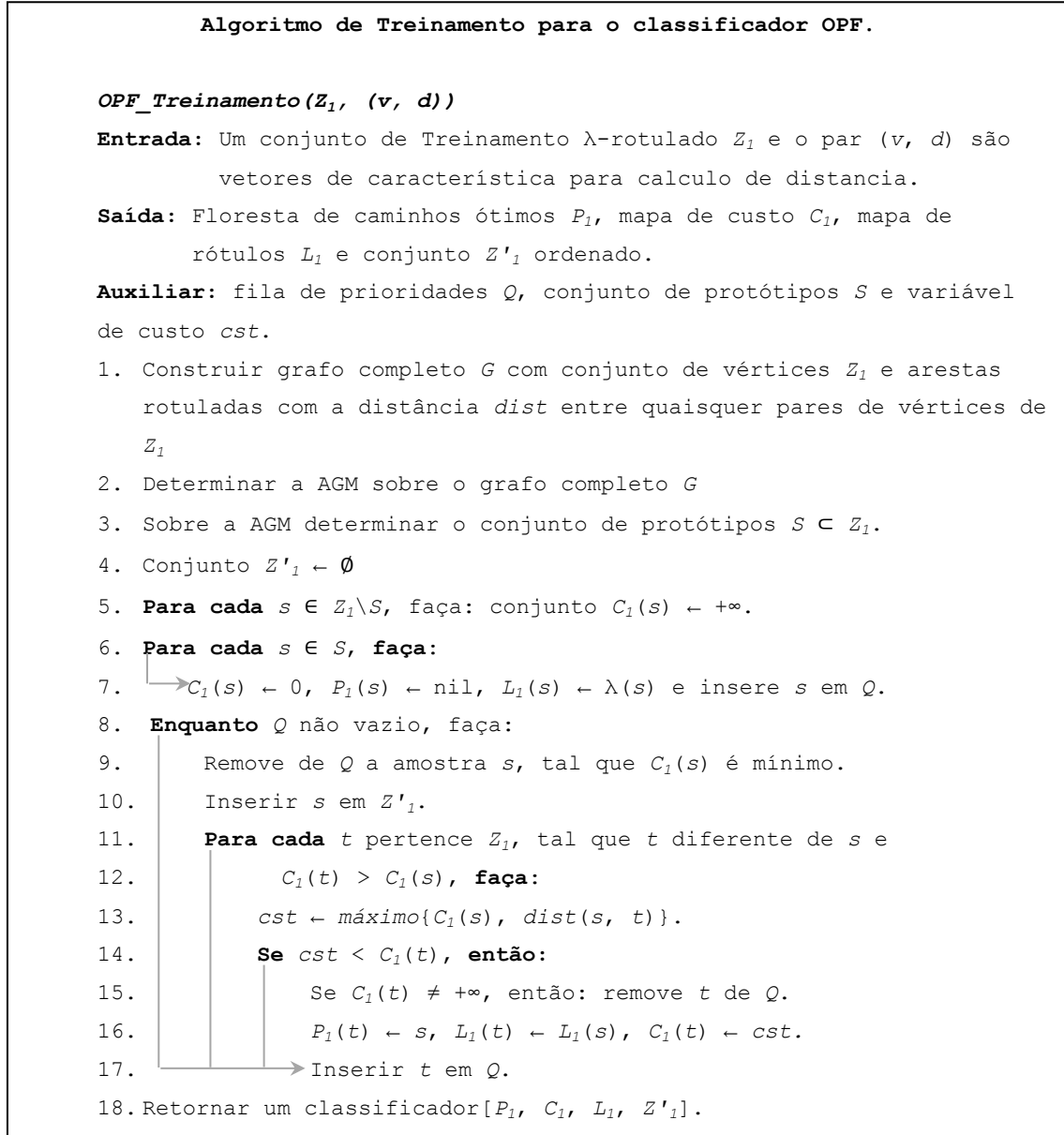


Figura 4.1. Algoritmo de treinamento do Classificador Floresta de Caminhos Ótimos.

Note-se que no grafo completo da Figura 4.3, cada amostra se relaciona com todas as outras amostras, independentemente da classe. Na figura, os círculos brancos representam uma classe e círculos pretos outra.

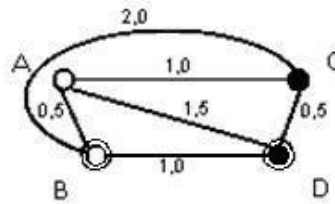


Figura 4.3. Relação de adjacência de um Grafo Completo.

A partir do grafo completo, é construída a Árvore Geradora de Custo Mínimo (AGM) (Cormen *et al.*, 2001). Como ilustrado na Figura 4.2, a AGM determina um grafo conexo tal que a soma das distâncias das arestas selecionadas é mínima. Esta árvore representa um mapa de conectividade mínima entre todas as amostras, independentemente da classe associada a estas amostras.

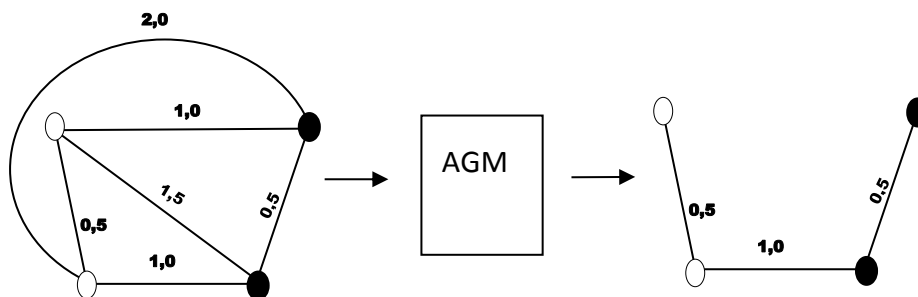


Figura 4.2. Construção da Árvore Geradora de Custo Mínimo (AGM).

Sobre a AGM são identificados os vértices protótipos, este processo é executado pelo algoritmo da linha 1 a 3. Os protótipos são representados na Figura 4.3, Figura 4.4 e Figura 4.5 pelos círculos dentro das circunferências tracejadas, e, são os vértices de uma classe no grafo completo, adjacentes aos vértices de outra classe. Na proposta deste classificador esses vértices são considerados os mais representativos em cada classe.

Na próxima etapa do algoritmo *OPF_Treinamento()* executado da linha 8 a 17, que os autores do método chamam de conquista, será determinado o menor caminho entre os vértices no grafo completo. A ideia é que cada protótipo conquiste os vértices não protótipos, oferecendo um caminho menor.

No algoritmo *OPF_Treinamento()* para a conquista dos vértices, é associado a cada vértice t (linha 11) um valor $C_1(t)$ (linha 12), que corresponde ao custo total do menor caminho

entre ele e um vértice protótipo. Para representar o caminho ótimo, cada vértice t possui uma referencia $P_1(t)$ ao próximo vértice no seu caminho até um vértice protótipo. Adicionalmente, para cada vértice t é definido um rótulo $L_1(t)$ que indica qual o vértice protótipo que o conquistou. As linhas 4 a 7 do algoritmo mostram a inicialização destas estruturas onde o custo C_1 de cada vértice protótipo é inicializado em zero (0) e na linha 5, cada vértice não protótipo é inicializado com infinito (∞).

O *OPF_Treinamento()* utiliza uma fila com prioridades Q onde são armazenados os vértices para os quais já foi determinado o caminho ótimo, mas que ainda podem oferecer um melhor caminho aos outros vértices. Inicialmente, todos os vértices protótipos são inseridos na fila (linha 7). Em cada passo, mostrado nas linhas 11 a 16 do algoritmo, um vértice s é extraído da fila e oferece aos outros vértices um caminho de menor custo através dele. Para isto, é determinado o valor máximo $\{C_1(s), d(s, t)\}$ na linha 13. Neste processo, cada vez que um vértice melhora seu caminho ótimo, seu custo, referência e rótulo são atualizados, este processo é executado da linha 14 a 17.

Na Figura 4.4 podemos observar um exemplo de um conjunto com relação de adjacência completa e as definições dos nós protótipos (B, D). A tabela ao lado do grafo mostra os respectivos pesos de cada nó com valor zero para os protótipos e infinito para os nós não protótipos. Os custos das arestas do grafo foram definidos pela função de distância.

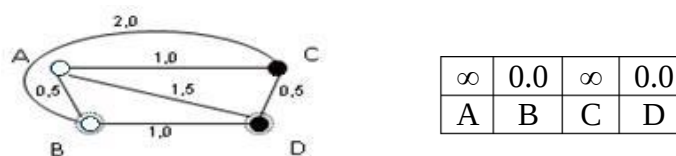


Figura 4.4. Vértices protótipo dentro da ACM.

Após percorrer todos os nós do grafo o *OPF_Treinamento()* retorna como resultado: o mapa de caminhos ótimos (P_1) que corresponde a uma floresta de caminhos ótimos; o vetor de custo ótimo de cada vértice (C_1); o vetor de rótulos de cada vértice e finalmente o conjunto de vértices de treinamento em ordem não decrescente do seu custo ótimo, o retorno destas estruturas é executado na linha 18 do algoritmo.

No exemplo mostrado na Figura 4.5, tínhamos inicialmente um grafo completo com os

nós (amostras) A, B, C e D de um conjunto OPF. Após a AGM encontrar os protótipos B e D, no processo de conquista o protótipo B conquistou A e o protótipo D conquistou C.

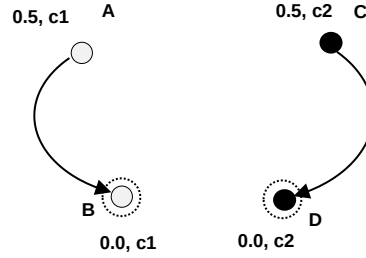


Figura 4.5. Floresta de Caminhos Ótimos (OPF).

Algoritmo de Classificação

O algoritmo de classificação (Papa *et al.*, 2012) é apresentado na Figura 4.6. Chamaremos o algoritmo de *OPF_Classificacao()*, e, ele recebe como entrada a Floresta de Caminhos Ótimos $[P_1, C_1, L_1]$ gerada na fase de treinamento (*OPF_Treinamento()*) e um conjunto Z_2 de amostras para serem classificadas.

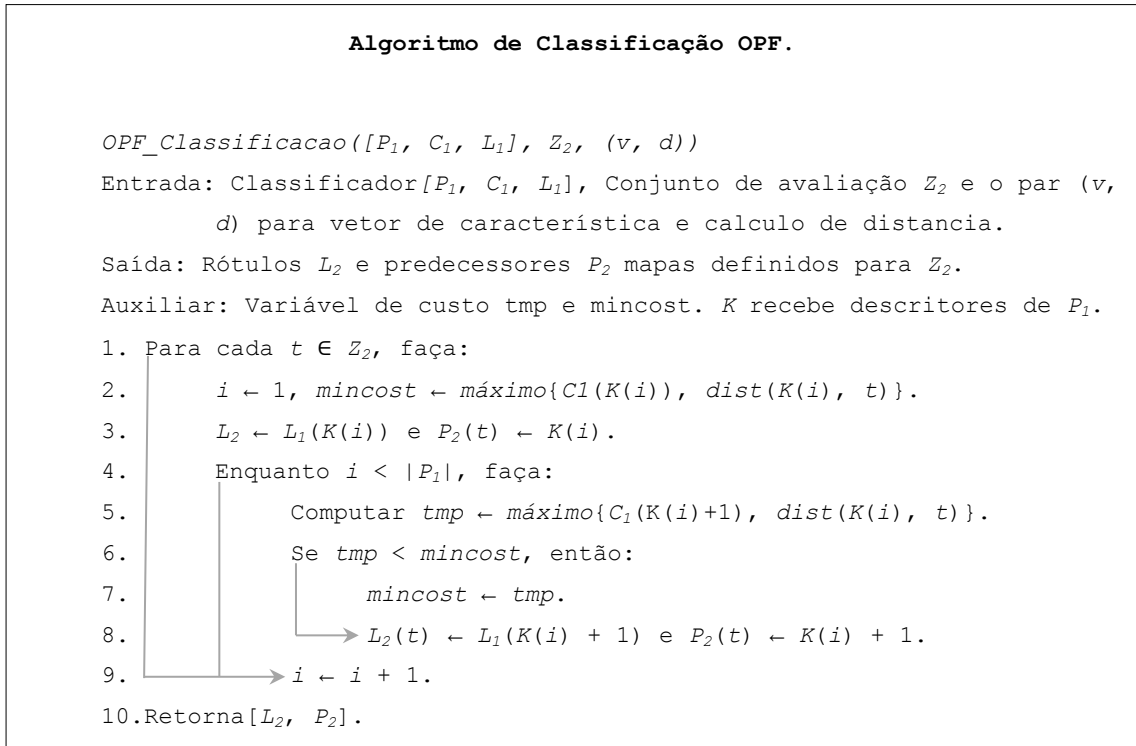


Figura 4.6. Algoritmo de Classificação da Floresta de Caminhos Ótimos.

No processo de classificação uma amostra t é considerada um vértice do conjunto de classificação (Z_2) que deve ser conquistado (linha 1). Neste processo K recebe as amostras do OPF, ou seja, de P_1 . A seguir, é utilizada a distância ($dist(K(i), t)$) entre t e cada um dos vértices de treinamento (OPF). Este cálculo mostrado na equação 11, é utilizado na linha 2 do algoritmo *OPF_Classificao()*.

$$C_2(t) = \text{mínimo}\{ \text{máximo}\{ C_1(s), d(s, t) \} \}, \forall s \in Z'_1. \quad (11)$$

A classe atribuída à amostra t (classificação) será aquela associada ao vértice $K(i)$ para o qual se obteve o valor mínimo. A amostra de classificação também receberá como peso o valor mínimo resultante da equação 11 ($\text{mínimo}(\text{máximo}())$).

A Figura 4.7 mostra um exemplo do resultado deste processo. O classificador tem o OPF com os descritores protótipos B e D e os descritores não protótipos A e C. Cada descritor do OPF possui um valor e o rótulo da classe, definido no treinamento. Para classificar uma nova amostra, fase de classificação, amostra Z no exemplo, aplica o processo de conquista. No exemplo da Figura 4.7 a amostra foi conquistada pelo nó D que pertence a classe 2.

Note-se que os vértices do mapa de caminhos ótimos são percorridos até o último vértice da estrutura P_1 obtida na fase de treinamento, como podemos ver na condição na linha 4 e, neste *loop* (linha 4 a 9) é determinada a classe da amostra de classificação, lembrando que K contém os vértices de P_1 .

O processo de conquista mostrado na Figura 4.7 com a amostra Z é executado para todas as amostras do conjunto de classificação $\{Z, X, Y, W \text{ e } V\}$.

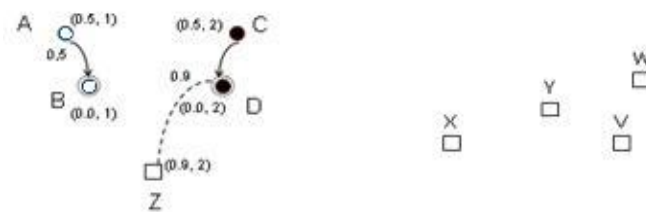


Figura 4.7. Processo de Classificação de uma nova amostra Z.

Como resultado, o algoritmo retorna o conjunto de amostras Z_2 com seus respectivos rótulos em L_2 (linha 10). Os rótulos em L_2 indicam a classe à qual foi associada cada uma

das amostras do conjunto de classificação (Z_2).

O laço principal do algoritmo da linha 1 a 9 percorre todos os elementos do conjunto Z_2 de amostras para serem classificadas. Para cada amostra de classificação (Z_2), são percorridos todos os vértices de treinamento da floresta em P_1 , ver linha 4 a 9.

Classificador Floresta de Caminhos Ótimos Aprimorado

Em Papa *et al.* (2012) apresenta-se o algoritmo classificador Floresta de Caminhos Ótimos Aprimorado (*Enhanced Optimum-Path Forest - EOPF*). Este algoritmo (*EOPF_Classificacao()*, Figura 4.8) basicamente utiliza o mesmo processo de conquista de vértices do seu predecessor (*OPF_Classificacao()*), mas busca uma redução do custo computacional do processamento através de algumas alterações.

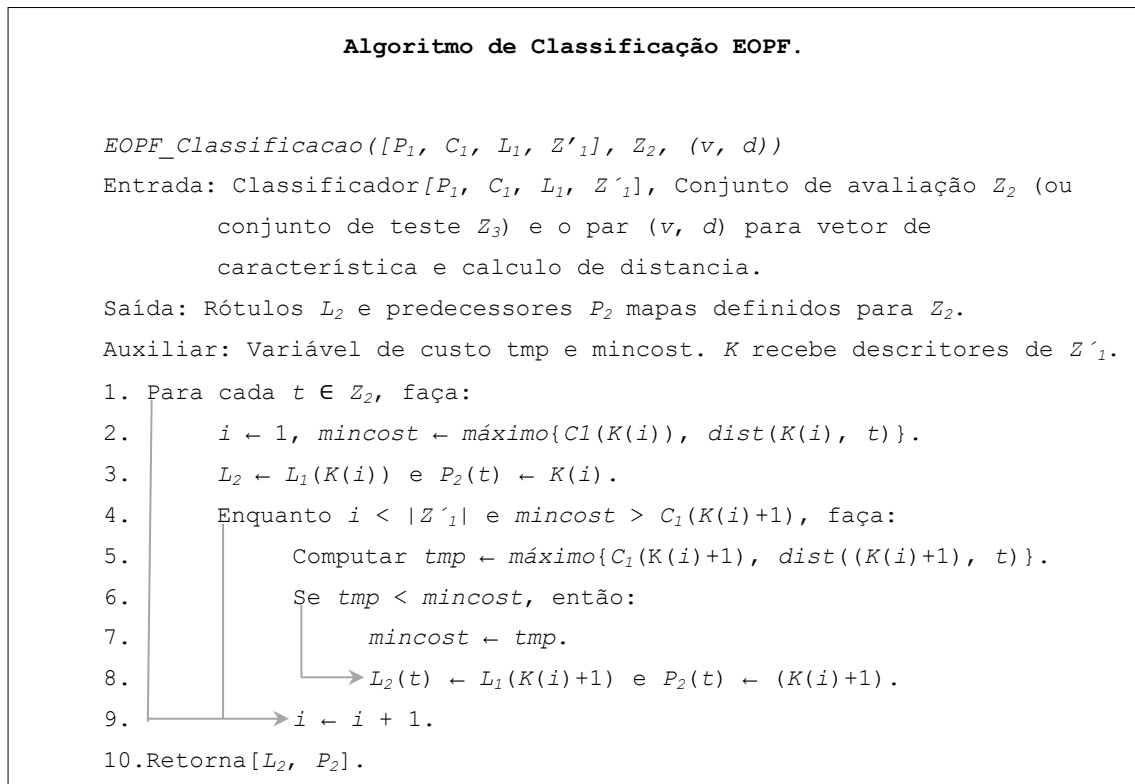


Figura 4.8. Algoritmo de Classificação da Floresta de Caminhos Ótimos.

A primeira diferença está na divisão do grupo de treinamento em dois conjuntos. O algoritmo utiliza o primeiro conjunto para realizar a fase de treinamento seguindo os passos

descritos no algoritmo da Figura 4.1. Um segundo conjunto de treinamento é utilizado para treinar novamente o resultado do primeiro. Com este procedimento, esperamos que o tamanho do conjunto de treinamento seja reduzido, e com isso deve ser reduzido o tempo de execução para a classificação de novas amostras do conjunto de classificação.

A partir desta redução, o processo busca um aprimoramento no treinamento, com o objetivo conseguir um conjunto de classes mais apuradas. Com este conjunto de classes, se espera uma precisão maior do algoritmo, ao processar os descritores na fase de classificação.

Para melhorar o tempo de execução do *EOPF_Classificacao()*, é realizado no algoritmo um processo denominado Poda. Neste processo o classificador ordena o conjunto dos descritores e com este conjunto ordenado pode considerar determinadas condições nas distâncias calculadas pelo OPF. Por exemplo, se o algoritmo conhece a ordem dos valores das arestas (relação), pode decidir se o resultado atual é o melhor ou não. Com esta decisão, o classificador evita que o algoritmo continue o processo de conquista até o fim do grafo. Ao parar o processo, o algoritmo evita o trabalho de um processamento desnecessário das amostras.

A técnica da Poda pretende reduzir o conjunto de descritores processados. Utilizando o conjunto de amostras mais relevantes, ainda assim, garantindo a precisão do algoritmo com um tempo de execução compatível com o processamento. Como o algoritmo OPF tem complexidade $O(N^2)$, sendo N o número de amostras, a redução do tamanho do conjunto proporciona uma melhora no tempo de execução da classificação. Com esta proposta, o novo algoritmo do OPF apresentou resultados interessantes nas pesquisas realizadas Papa *et al.* (2012).

O algoritmo *EOPF_Classificacao()* na Figura 4.8, possui uma estrutura similar ao *OPF_Classificacao()* (Figura 4.1), diferenciando nos seguintes pontos: parâmetros recebido; a estrutura Z'_1 , que é o conjunto de treinamento (OPF) ordenado pelo custo das amostras de maneira não decrescente; e na comparação da linha 4; onde a condição $mincost > C_1(K(i)+1)$ verifica se o custo atual da amostra de treinamento é maior que o custo atual da amostra de treinamento posterior. A condição é satisfeita, e o algoritmo continua processando, até que não se encontre custo menor no conjunto de amostras de treinamento (Z'_1).

Aprimoramento com Vários Classificadores OPFs e Métodos de Votação

Outra abordagem visando um melhor desempenho do método é o uso de vários classificadores OPF para classificar grandes conjuntos de vetores de características. Esta técnica (Ponti Jr *et al.*, 2011) consiste na divisão do conjunto de treinamento em n subconjuntos de igual tamanho. Estes subconjuntos podem ser processados de forma independente por vários classificadores OPF. Para uma nova amostra, cada classificador decide a que classe ela pertence e, finalmente, através de um algoritmo de votação, se decide a classe da amostra.

Existem vários algoritmos de votação que podem ser utilizados (Naldi *et al.*, 2009). Como por exemplo, o algoritmo da média absoluta, onde é considerado o valor da média das notas (valor de similaridade obtido por cada OPF). O cálculo da média também pode variar para o algoritmo que chamamos de Média Parcial. Neste é calculado o valor da média dos resultados, desconsiderando-se o maior e o menor valor. Ainda temos os métodos de votação do máximo e do mínimo em que se consideram a nota maior dos resultados dos OPFs para o máximo, e o menor resultado para o mínimo.

Os resultados apresentados em Ponti Jr *et al.* (2011) descrevem uma melhora com relação ao OPF original. Além desta vantagem, esta abordagem possibilita o uso do processamento distribuído entre vários processadores ou sistemas, o que torna a proposta interessante para o estudo do tempo de processamento de grandes conjuntos de classificação.

Considerações Finais

No Capítulo 4 apresentamos o algoritmo classificador Floresta de Caminhos Ótimos (OPF). Também apresentamos duas variações do classificador: o classificador Floresta de Caminhos Ótimos Aprimorado e algoritmo do OPF com a abordagem por votação. Foi descrito o processo de melhora do classificador OPF a partir da poda de vértices da floresta, assim como a abordagem de múltiplos classificadores integrados por algoritmo de votação. O classificador OPF será o centro da pesquisa desenvolvida nesta dissertação e descrita no próximo capítulo.

5. Família de Classificadores OPF para Reconhecimento Facial

Este capítulo apresenta a proposta desta dissertação para o uso de uma família de classificadores OPF para reconhecimento facial.

Como apresentado na Capítulo 3, uma das abordagens de um sistema de verificação considera a extração de um descritor a partir da integração de duas imagens. A partir deste descritor é obtido um valor de similaridade que estima se as duas imagens comparadas são do mesmo indivíduo ou não.

Também no Capítulo 3 foi apresentada a abordagem Bayesiana que é direcionada a esse tipo de sistema de verificação. Nesta abordagem as imagens são integradas através da sua imagem diferença. O método cria um Classificador Bayesiano treinado com duas classes de descritores: a Classe Intrapessoal, treinada com descritores obtidos de imagens do mesmo indivíduo e a Classe Interpessoal, treinada com descritores obtidos de imagens indivíduos diferentes. Quando duas imagens são comparadas, um descritor é extraído da imagem diferença e para este descritor é atribuído um valor de similaridade que estima sua proximidade com cada uma das duas Classes Intrapessoal e Interpessoal. Este valor de similaridade é utilizado para o reconhecimento por verificação.

Nesta dissertação propomos uma abordagem para reconhecimento facial baseada na classificação das Classes Intrapessoal e Interpessoal. Nossa abordagem apresenta características de ser facilmente adaptável ao processamento distribuído e paralelo pois, a proposta consiste em substituir o uso de um único classificador por uma família de classificadores (vários classificadores) mais leves e integrar seus resultados através de métodos de votação. Cada um destes classificadores da família poderia ser executado de maneira paralela o que diminuiria o tempo total de processamento. Para a criação da família, o conjunto de descritores de treinamento deve ser dividido igualmente e de forma aleatória entre os classificadores da família.

Neste trabalho realizamos a pesquisa utilizando uma família de classificadores OPF apresentado no capítulo anterior. A decisão de estudar o comportamento do classificador OPF baseia-se em duas características deste algoritmo apresentadas a seguir:

- **Classificadores da família mais leves**

O custo da classificação no OPF depende da quantidade de descritores utilizados para o treinamento. Como na nossa abordagem dividimos os descritores de treinamento entre todos os classificadores, quanto maior o número de classificadores na família menor será o custo computacional da classificação de cada classificador OPF.

- **Atribuição de um valor de similaridade na classificação**

Um classificador OPF convencional recebe um vetor de características como entrada e retorna como saída qual a classe a que esse vetor pertence. Para o uso de um classificador OPF na nossa abordagem, ele tem que ser adaptado para retornar como saída um valor de similaridade.

No caso do classificador OPF é possível estimar o valor de similaridade de um descritor a partir de sua relação de distância com cada uma das classes Intrapessoal e Interpessoal. Esta função de similaridade, apresentada na equação 12, deve ser normalizada no intervalo entre zero e um (0,1) e deve atribuir valor um (1) à total similaridade e zero (0) a total dissimilaridade.

Dado um descritor $F_{n,m}$ obtido da diferença entre as imagens I_n e I_m , e dadas as classes X (intrapessoal) e Y (interpessoal) definimos como função de similaridade entre I_n e I_m .

$$sim(I_n, I_m) = 1 - \frac{d(F_{n,m}, X)}{d(F_{n,m}, X) + d(F_{n,m}, Y)} \quad (12)$$

em que:

- $d(F_{n,m}, X)$ é o valor da distância do descritor $F_{n,m}$ com a Classe Intrapessoal.
- $d(F_{n,m}, Y)$ é o valor da distância do descritor $F_{n,m}$ com a Classe Interpessoal.

Nas seções a seguir apresentamos nossas propostas de alteração nas funções do algoritmo OPF. Estas alterações se justificam na necessidade de que o algoritmo trabalhe como um classificador biométrico, assim como classificar os descritores com base nos valores de

similaridade.

Etapas do Processo

Na proposta deste trabalho temos duas etapas para o reconhecimento facial:

1. Treinamento de uma família de classificadores OPF $C_1, C_2, \dots, C_n$ a partir de um conjunto de descritores rotulados nas classes Intrapessoal X e Interpessoal Y .
2. Dado um conjunto de descritores W de teste, determinar o valor de similaridade de cada um deles utilizando a família de classificadores OPF $C_1, C_2, \dots, C_n$ e um método de votação V .

Etapa 1: Treinamento da Família de Classificadores

Foi definido e implementado o algoritmo *Gerar_Familias()*, utilizado para o treinamento da família de classificadores. Este algoritmo descrito na Figura 5.1, recebe como entrada todo o conjunto de descritores de treinamento rotulados nas classes Intrapessoal (X) e Interpessoal (Y), e um valor N que indica a quantidade desejada de classificadores na família.

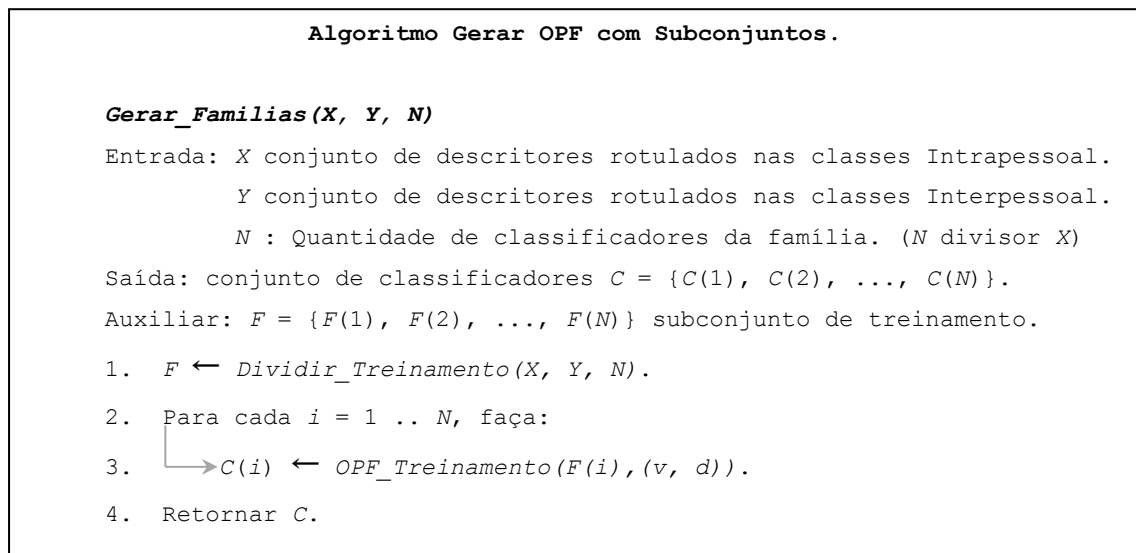


Figura 5.1. Algoritmo de treinamento da família de classificadores.

Na linha 1 o conjunto de descritores é dividido, pelo algoritmo *Dividir_Treinamento()*, em N subconjuntos (F), todos os subconjuntos gerados devem conter a mesma quantidade de

descritores de cada uma das classes. Nas linhas 2 e 3 cada um destes subconjuntos F é utilizado para treinamento de um classificador OPF (C), utilizando o algoritmo original de treinamento do OPF, o $OPF_Treinamento()$ definido na Figura 4.1 no capítulo anterior. Como resultado o algoritmo retorna N classificadores na linha 4, sendo $C = \{C(1), C(2), \dots, C(N)\}$.

Algoritmo de para dividir o conjunto de treinamento

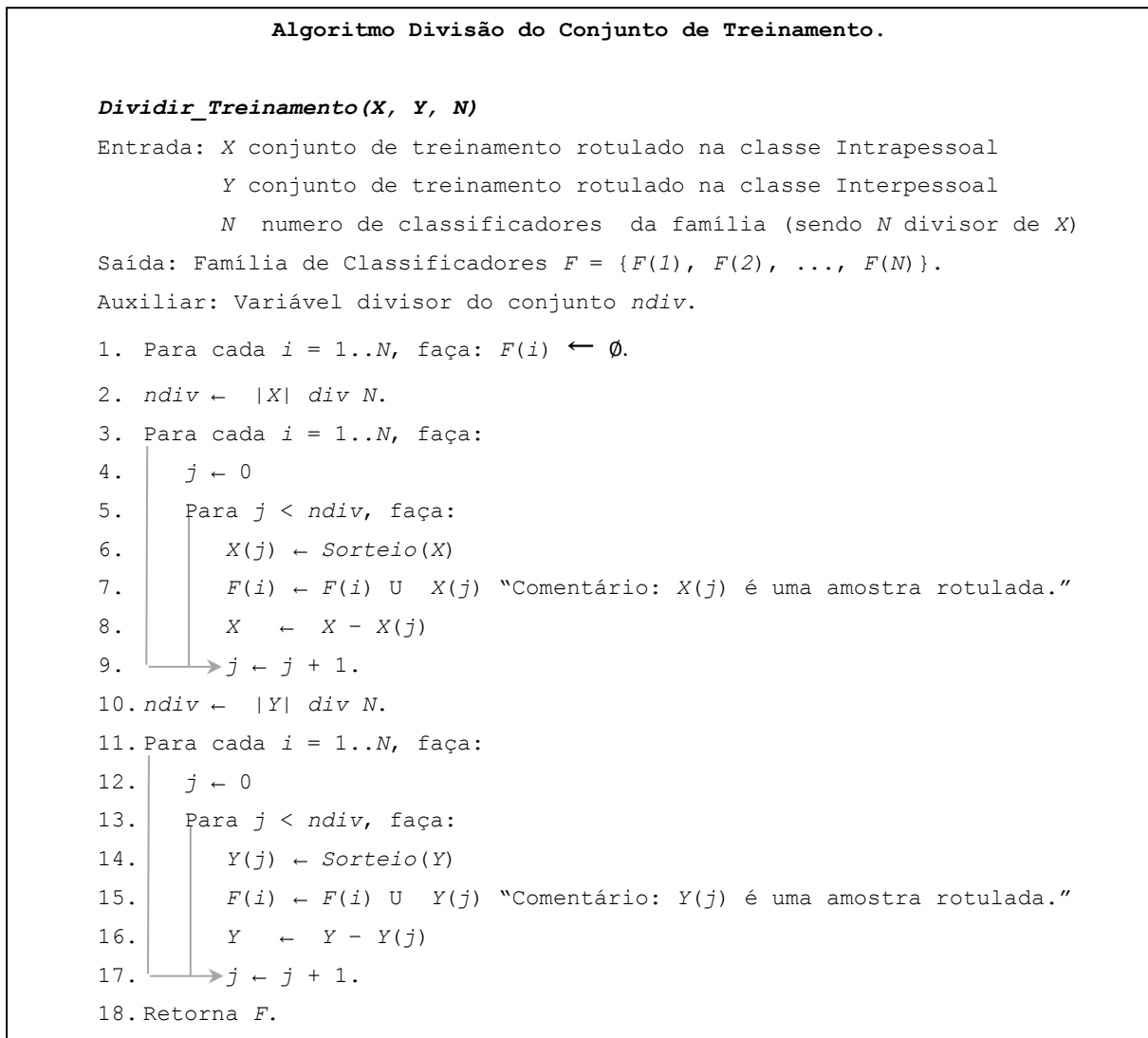


Figura 5.2. Algoritmo para dividir a conjunto de descritores de treinamento.

Para realizar a divisão do conjunto de descritores de treinamento foi definido o algoritmo $Dividir_Treinamento()$, descrito na Figura 5.2. Este algoritmo distribui igual e

aleatoriamente os descritores entre N conjuntos, garantindo a mesma quantidade de descritores das classes Intrapessoal e Intrapessoal em cada conjunto.

O algoritmo recebe um valor N , que informa em quantos subconjuntos deve dividir o conjunto original de treinamento. Com este valor calcula, nas linhas 2 e 10, quantas amostras da classe intrapessoal e quantas da classe interpessoal caberão para cada subconjunto. Os *loops* das linhas 5 a 9, para classe intrapessoal, e nas linhas 13 a 17 para a classe interpessoal, inserem as amostras aleatoriamente nos subconjuntos $F(i)$ nas linhas 7 e 15.

No final teremos N subconjuntos, sendo N um número inteiro ≥ 2 e $N \leq \lfloor \text{tamanho do conjunto de treinamento} / 2 \rfloor$. Cada subconjunto é um OPF votante da família de classificadores.

Etapa 2: Cálculo do valor de similaridade de um conjunto de descritores de teste

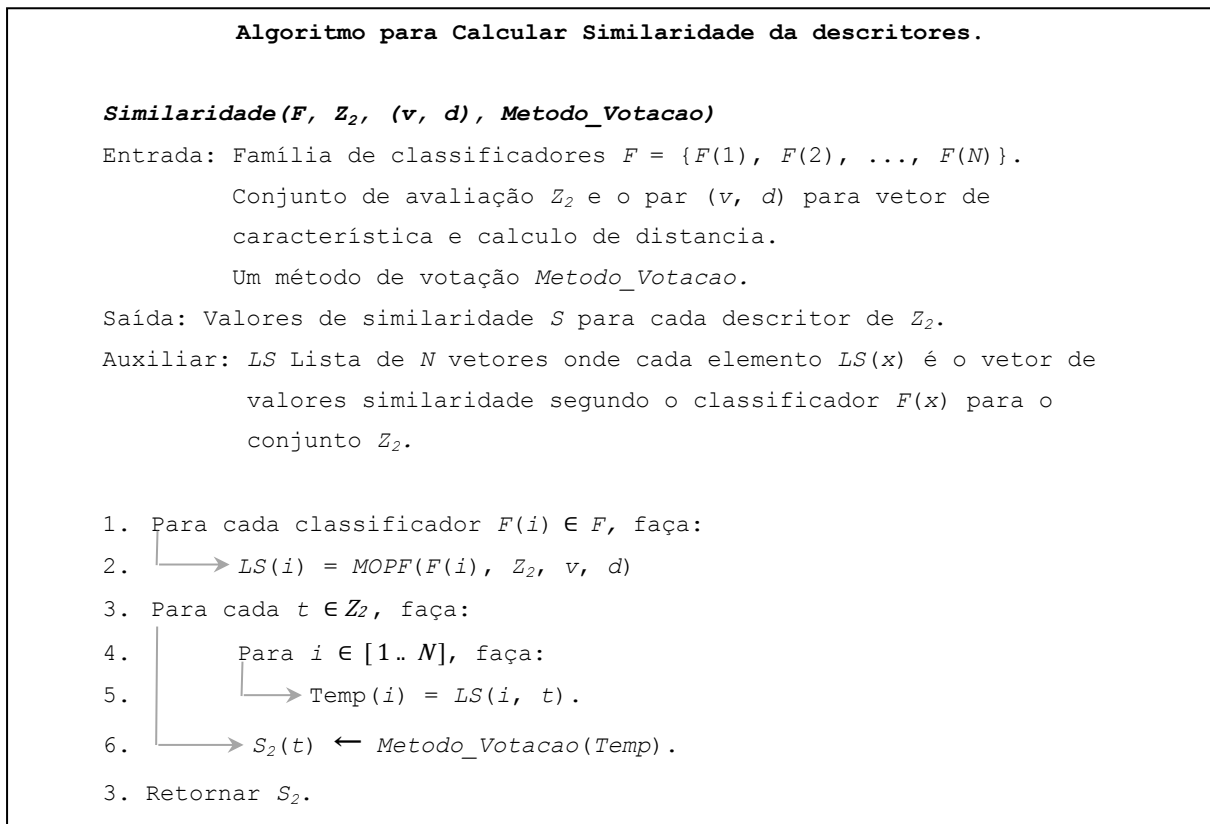


Figura 5.3. Algoritmo para calcular o valor de similaridade dos descritores.

Esta etapa utiliza a família (F) de N classificadores construídos na etapa 1 para atribuir um valor de similaridade a um conjunto de novos descritores a serem classificados (conjunto de

classificação), recebidos como entrada (Z_2). O algoritmo *Similaridade()*, definido na Figura 5.3, recebe a família de classificadores F e conjunto de descritores de teste Z_2 . Para cada classificador $F(i)$ da família, o algoritmo *MOPF()*, que comentaremos a seguir nesta capítulo (Figura 5.4), é chamado e retorna o vetor de valores de similaridades calculadas para os descritores em Z_2 em cada classificador $F(i)$.

Na sequência, para cada descritor t de Z_2 é construído o vetor *Temp* com cada um dos valores de similaridade calculados para t em cada classificador. A seguir, a função *Metodo_Votacao()* de votação integra todos esses valores de similaridade em um valor final de similaridade de acordo com o método de votação adotado. O resultado final do algoritmo é a lista de similaridades S_2 para cada descritor de Z_2 .

A Figura 5.4 mostra o algoritmo *MOPF()* que calcula os valores de similaridade para um conjunto de descritores Z_2 utilizando um classificador. Para implementar este algoritmo foram realizadas modificações ao algoritmo classificador OPF que foi apresentado na Figura 4.6.

Para classificar uma amostra, o algoritmo OPF original percorre todos os vértices do classificador mantendo sempre qual o custo do menor caminho e o rótulo da classe que oferece esse caminho. No nosso caso, temos unicamente as duas classes Intrapessoal e Interpessoal e precisamos manter o custo do menor caminho a cada uma dessas duas classes para poder calcular a similaridade segundo a equação 12.

Como mostra a Figura 5.4 o algoritmo utiliza as variáveis *costX* e *costY* para manter em cada passo o custo de menor caminho para cada uma das classes Intrapessoal e Interpessoal.

No laço nas linhas 1 a 16 cada descritor $t \in Z_2$ é analisado percorrendo cada um dos vértices do classificador. Para cada vértice calcula-se o custo do descritor para esse vértice na linha 6 e verifica-se se o custo que ele oferece é menor que o custo já armazenado na variável associada a sua classe. Se o custo é menor, então a variável é atualizada com o novo valor.

No final do processamento, na linha 15, o valor de similaridade para o descritor é calculado utilizando a equação 12.

Algoritmo de Classificação MOPF.

MOPF($[P_1, C_1, L_1, Z'_1], Z_2, (v, d)$)

Entrada: Classificador $[P_1, C_1, L_1, Z'_1]$, Conjunto de avaliação Z_2 e o par (v, d) para vetor de característica e calculo de distancia.

Saída: Rótulos L_2 e predecessores P_2 mapas definidos para Z_2 .

Auxiliar: Variável de custo tmp e $mincost$, distâncias das classes intrapessoal D_1 e interpessoal D_2 .

```

1. For cada  $t \in Z_2$ , faça:
2.    $i \leftarrow 0$ .
3.    $costX \leftarrow +\infty$ 
4.    $costY \leftarrow +\infty$ 
5.   Enquanto  $i < |Z'_1|$  e  $(\text{máximo}(costX, costY) > C_1(k_{(i+1)}))$ , faça:
6.      $tmp \leftarrow \text{máximo}\{C_1(k_{(i+1)}), \text{dist}(k_{(i+1)}, t)\}$ .
7.     Se  $L_1(k_{(i+1)}) = 0$ , então:
8.       Se  $tmp < costX$ , então:
9.          $costX \leftarrow tmp$ 
10.         $L_2(t) \leftarrow L_1(k_{(i+1)})$  e  $P_2(t) \leftarrow k_{(i+1)}$ .
11.      Senão
12.        Se  $tmp < costY$ , então:
13.           $costY \leftarrow tmp$ 
14.           $L_2(t) \leftarrow L_1(k_{(i+1)})$  e  $P_2(t) \leftarrow k_{(i+1)}$ .
15.         $L_2(t) \leftarrow (1 - (costX / (costX + costY)))$ .
16.       $i \leftarrow i + 1$ .
17. Retorna  $[L_2, P_2]$ .

```

Figura 5.4. Algoritmo de Classificação do MOPF.

Considerações Finais

Neste capítulo foi apresentada a proposta do uso da família de classificadores OPF para serem aplicados no reconhecimento facial. Apresentamos as adaptações deste classificador, visando sua adequação a este tipo de sistema, assim como alguns aspectos da base de imagens considerada.

O próximo capítulo descreve os experimentos realizados neste trabalho.

6. Experimentos e Análise dos Resultados

Este capítulo apresenta os resultados experimentais para avaliar a família de classificadores OPF para reconhecimento facial.

6.1. Metodologia

Para a avaliação da proposta foram criados dois conjuntos de descritores de treinamento e de teste utilizando imagens faciais do banco de imagens FERET, Phillips *et al.* (2000) e Phillips *et al.* (2003). As imagens são de pessoas distintas e de ambos os sexos

O conjunto de treinamento foi criado com 65 indivíduos com 6 imagens para cada um deles. Este conjunto permitiu gerar 975 imagens diferença do mesmo indivíduo, a partir de todas as combinações de duas imagens possíveis. Foram geradas também 2080 imagens diferença de indivíduos diferentes selecionadas aleatoriamente do conjunto das 74.880 combinações possíveis.

O conjunto de classificação foi criado com 50 indivíduos com 4 imagens para cada um deles. Este conjunto permitiu gerar 300 imagens diferença do mesmo indivíduo, a partir de todas as combinações de duas imagens possíveis. Foram geradas também 1264 imagens diferença de indivíduos diferentes selecionadas aleatoriamente do conjunto das 19600 combinações possíveis.

Sobre estas imagens diferença foram aplicados os métodos *Eigenfaces* e *EBGM* para gerar os descritores. Para cada um destes métodos foram gerados dois conjuntos de descritores das Classes Intrapessoal e Interpessoal para o treinamento e os dois conjuntos destas mesmas classes para o teste da família de classificadores OPF. Para gerar estes descritores, foi utilizada a biblioteca *CSU Face Evaluation System* da Universidade de Colorado (Beveridge et. al., 2005).

Parâmetros pesquisados

Os experimentos realizados incluíram a pesquisa da precisão com variação dos seguintes parâmetros:

- Número de classificadores da família
- Métodos de votação.
- Funções de distância.

Número de classificadores da família

Foram criadas famílias de quatro (04), cinco (05), dez (10), quinze (15), vinte (20), vinte e cinco (25) e cinquenta (50) classificadores. Em todos estes casos o conjunto de treinamento foi aleatoriamente dividido em partes iguais entre todos os classificadores da família.

Métodos de votação pesquisados

Como foi apresentado na proposta, o valor de similaridade final associado ao descritor de uma imagem diferença é calculado a partir da integração pelos métodos de votação. Cada método de votação realiza a integração dos valores de similaridade calculados por cada um dos classificadores da família, conforme a equação definida para cada método. Em nossa pesquisa consideramos quatro métodos de votação:

- Mínimo: o valor menor de todos os valores de similaridade dos classificadores da família;
- Máximo: o valor maior de todos os valores de similaridade dos classificadores da família;
- Média Absoluta: o valor da média dos valores de similaridade dos classificadores da família;
- Média Parcial: o valor da média dos valores de similaridade dos classificadores da família, excluindo o maior e menor valor de similaridade.

Funções de distância que foram consideradas

Em todos os métodos de reconhecimento facial utilizados assim como para a construção dos classificadores, a função de similaridade é estimada a partir de uma função de distância entre os descritores das imagens diferença. Em nossa pesquisa utilizamos várias alternativas de

funções de distância entre os descritores que são listadas a seguir:

- Bray-Curtis;

A função Bray-Curtis (Bray *et al.*, 1957) calcula a dissimilaridade entre os descritores representada com um valor entre zero e um, através a equação 13:

$$distBC = \sum_i^N \frac{(dt_i - dc_i)}{somD_i} \quad (13)$$

onde:

O dt_i é o descritor do OPF (treinamento); dc_i é o descritor de classificação e $somD_i$ representa a soma dos descritores ($dt_i + dc_i$). A equação retorna em $distBC$ a somatória dos valores de dissimilaridade das características do descritor que vão de i até N .

- Canberra;

A função Canberra (Lance *et al.*, 1966) é mostrado na equação 14, onde $distC$ recebe a distância entre dois descritores em um espaço vetorial. As variáveis dt_i e dc_i , assim com na dissimilaridade Bray-Curtis, representam os respectivos descritores.

$$distC = \sum_i^N \frac{|dt_i - dc_i|}{|dt_i| + |dc_i|} \quad (14)$$

- Euclidiana Log;

A função Euclidiana Log, calculada na equação 15 e retorna o valor em $distLE$, utiliza o Teorema de Pitágoras para calcular a distância entre os vetores. Sobre o resultado aplica a função logarítmica.

$$distLE = \log \sum_i^N \sqrt{(dt_i - dc_i)^2 + (dt_i - dc_i)^2} \quad (15)$$

- Euclidiana;

A função Euclidiana calcula a distância com base no Teorema de Pitágoras. A equação 16 mostra o cálculo que guarda o resultado em $distE$. As variáveis dt_i e dc_i contêm os valores de características dos descritores de treinamento e classificação.

$$distE = \sum_i^N \sqrt{(dt_i - dc_i)^2 + (dt_i - dc_i)^2} \quad (16)$$

- Manhattan;

A equação 17 mostra a distância Manhattan (Barroso, 1998), onde $distM$ armazena a soma das distâncias entre os descritores.

$$distC = \sum_i^N |dt_i - dc_i| + |dt_i - dc_i| \quad (17)$$

- Chi-Squared;

A função Chi-Squared é mostrada na equação 18, onde $distCS$ recebe o resultado da equação que utiliza as variáveis dt_i e dc_i com os valores dos descritores de treinamento e classificação respectivamente. As variáveis $somDt$ e $somDc$ contêm a soma dos valores das características do descritor do conjunto de treinamento (OPF), e do conjunto de classificação respectivamente.

$$distCS = \sum_i^N \frac{1}{dt_i + dc_i} * \left(\frac{dt_i}{somDt} - \frac{dc_i}{somDc} \right)^2 \quad (18)$$

- Squared Chord.

A função Squared Chord é mostrada na equação 19. A equação é válida quando os

valores de características dt_i (descritor de treinamento) e dc_i (descritor de classificação) são igual ou maiores que zero. A variável $distSC$ guarda o resultado da equação.

$$distSC = \sum_i^N (\sqrt{dt_i} - \sqrt{dc_i})^2 \quad (19)$$

Métricas de avaliação

Para avaliar os resultados da precisão utilizaremos as métricas biométricas mais importantes e descritas no Capítulo 2. Os resultados serão apresentados com tabelas contendo os valores de EER, FMR100, FMR1000 e FMRZero. Os resultados serão apresentados também através das curvas ROC.

Implementações realizadas

Para realizar os experimentos foram realizadas as seguintes implementações:

1. *Programa classificador estendido OPF.* Este programa implementou a extensão do algoritmo de classificação OPF apresentada no capítulo anterior.
2. *Programa para criação da família de n classificadores.* Este programa divide o conjunto de treinamento em n subconjuntos e realiza o treinamento da família de n classificadores OPF.
3. *Programa para a classificação.* Este programa determina o valor de similaridade de cada descritor do conjunto de teste em cada um dos classificadores da família previamente treinados. Calcula o valor final de similaridade de acordo com o método de votação selecionado. Gera um arquivo com os valores de similaridade calculados.
4. *Programa para cálculo das métricas.* Este programa recebe dois arquivos com valores de similaridades de descritores obtidos de comparações genuínas e impostoras e gera as funções FMR e FNMR e as métricas de avaliação EER, FMR100, FMR1000, FMRZero e curvas ROC.

Para esta implementação foi utilizada a biblioteca libopf (Papa *et al.*, 2014) que implementa o classificador OPF na linguagem C como descrito no Capítulo 5. O código desta biblioteca

foi utilizada e alguns casos modificada para obter os programas necessários.

Experimentos realizados

Os experimentos realizados tiveram como objetivo verificar em que casos o uso de uma família de classificadores melhora a precisão do reconhecimento facial. Neste sentido, foi necessário estudar quais funções de distância e quais métodos de votação produzem os melhores resultados. Em todos os casos foram utilizados descritores obtidos dos métodos *Eigenfaces* e *EBGM*.

Dada a diversidade das variáveis sendo consideradas na pesquisa (Funções de distância, métodos de votação e tamanho das famílias) foi necessário selecionar as funções de distância e métodos de votação. Com este objetivo foram realizados os seguintes experimentos:

- *Experimentos com diferentes funções de distância.* O objetivo destes experimentos foi comparar de forma preliminar os resultados com diferentes funções de distância.
- *Experimentos com diferentes métodos de votação.* O objetivo destes experimentos foi selecionar os métodos de votação com melhores resultados.
- *Experimentos com famílias de classificadores OPF.* Estes experimentos utilizam famílias de classificadores de diferente tamanho combinando as funções de distância e métodos de votação selecionados. Tem como objetivo verificar a hipótese de que os resultados com uma família de classificadores são melhores que com um único classificador.

Na primeira serie de experimentos foi verificado que as funções de distância Manhattan, Euclidiana e Chi-Squared são as que produzem os melhores resultados. Estas funções foram escolhidas para realizar a pesquisa com a família de classificadores.

Na segunda serie de experimentos foram introduzidas famílias de classificadores e foi utilizada a função de distância Chi-Squared para estudar o comportamento de diferentes métodos de votação. O resultado mostrou que os métodos de votação Média Parcial e Média absoluta produzem os melhores resultados.

Finalmente, a terceira serie de experimentos combinou as funções de distância e métodos de votação escolhidos para estudar a precisão quando estes são utilizados com famílias de classificadores.

Nas seções seguintes apresentamos resultados dos experimentos e as análises dos resultados.

6.2. Experimentos com Diferentes Funções de Distância

Como primeiro passo da pesquisa, utilizamos um único classificador OPF, com os descritores criados pelo método da *Eigenfaces* e foram comparados os resultados utilizando as diferentes funções de distância consideradas.

A Figura 6.1 mostra as curvas ROC com a precisão do classificador OPF para as funções de distância consideradas em nosso trabalho.

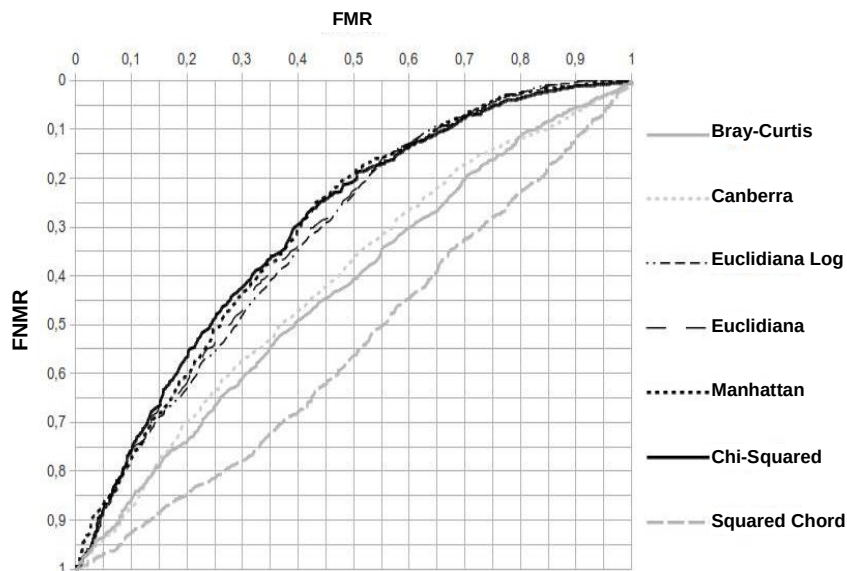


Figura 6.1. Resultados com um classificador OPF utilizando descritores do método *Eigenfaces* e várias funções de distância.

A Tabela 6.1 mostra um resumo dos valores das principais métricas. Na primeira coluna temos as funções de distâncias consideradas; na segunda a métrica EER, na terceira, quarta e quinta colunas temos as métricas FMR100, FMR1000 e FMRZero respectivamente.

Tabela 6.1. Métricas de Avaliação obtidas com classificador OPF com amostras *Eigenfaces*. Comparação dos resultados para cada função de distância.

Funções de Distância	EER	FMR100	FMR1000	FMRZero
Bray-Curtis	0,445247	0,983221	0,998881	1,000000
Canberra	0,435853	0,985459	0,997763	0,997763
Euclidiana Log	0,396657	0,845638	0,911633	0,917226
Euclidiana	0,369750	0,850112	0,911633	0,911633
Manhattan	0,363684	0,891499	0,945190	0,946309
Chi-Squared	0,356029	0,892617	0,994407	0,994407
Squared Chord	0,525196	0,995521	0,998880	1,000000

Podemos ver que, para todas as métricas avaliadas, os melhores resultados são obtidos com as funções de distância Euclidiana, Manhattan e Chi-Squared.

Para o próximo experimento foram utilizados descritores obtidos do método *EBGM* e seus resultados, para as funções de distâncias utilizadas, são apresentados na Figura 6.2.

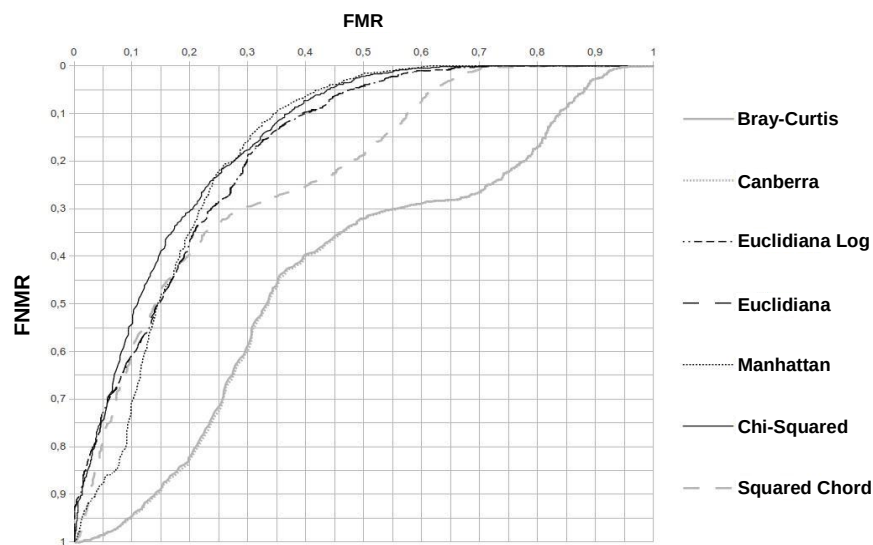


Figura 6.2. Resultados com um classificador OPF utilizando descritores do método *EBGM* e várias funções de distância.

Na Tabela 6.2 temos os resultados obtidos com as funções de distância descritas na primeira coluna, as métricas de avaliação EER na segunda coluna, FMR100 na terceira coluna, FMR1000 na quarta e a métrica FMRZero com os valores na quinta coluna.

Observando os resultados da tabela concluímos que as funções de distância com melhores resultados são a Euclidiana, Manhattan e Chi-Squared.

Tabela 6.2. Métricas de Avaliação obtidas com classificador OPF e descritores do *EBGM*. Comparação dos resultados das funções de distâncias.

Funções de Distâncias	EER	FMR100	FMR1000	FMRZero
Bray-Curtis	0,424878	0,933333	0,951795	0,958974
Canberra	0,396873	0,922998	0,947639	0,954825
Euclidiana Log	0,271094	0,631795	0,728205	0,741538
Euclidiana	0,268718	0,595897	0,694359	0,714872
Manhattan	0,241635	0,544615	0,604103	0,612308
Chi-Squared	0,240013	0,554872	0,678974	0,702564
Squared Chord	0,296941	0,691282	0,743590	0,779487

Note-se que, em geral, os resultados de precisão do reconhecimento são melhores com descritores extraídos do método *EBGM* que do método *Eigenfaces*. Estas métricas se correspondem com os resultados reportados para estes métodos quando utilizados de forma convencional. Considerando que durante a pesquisa as funções Euclidiana, Manhattan e Chi-Squared mantiveram os melhores resultados nos experimentos, a seguir apresentaremos unicamente os resultados com estas funções.

6.3. Experimentos com Diferentes Métodos de Votação

A segunda serie de experimentos consistiu em avaliar quais os melhores métodos de votação. Para estes experimentos foi escolhida a função de distância Chi-Squared que mostrou os melhores resultados nos experimentos realizados com um único classificador.

Foram selecionadas famílias de 04, 05, 10, 15, 20, 25 e 50 classificadores com combinação de seus valores de similaridade utilizando os diferentes métodos de votação propostos.

A Figura 6.3 mostra os resultados dos experimentos com descritores obtidos pelo método *Eigenfaces* e a Figura 6.4 os resultados com descritores do método *EBGM*. Nestas figuras é apresentado o comportamento para cada família.

Todos estes experimentos mostram que os métodos de votação com melhores resultados foram os da Média Absoluta e Média Parcial. Este comportamento foi igualmente comprovado utilizando as outras funções de distância selecionadas: Euclidiana e Manhattan.

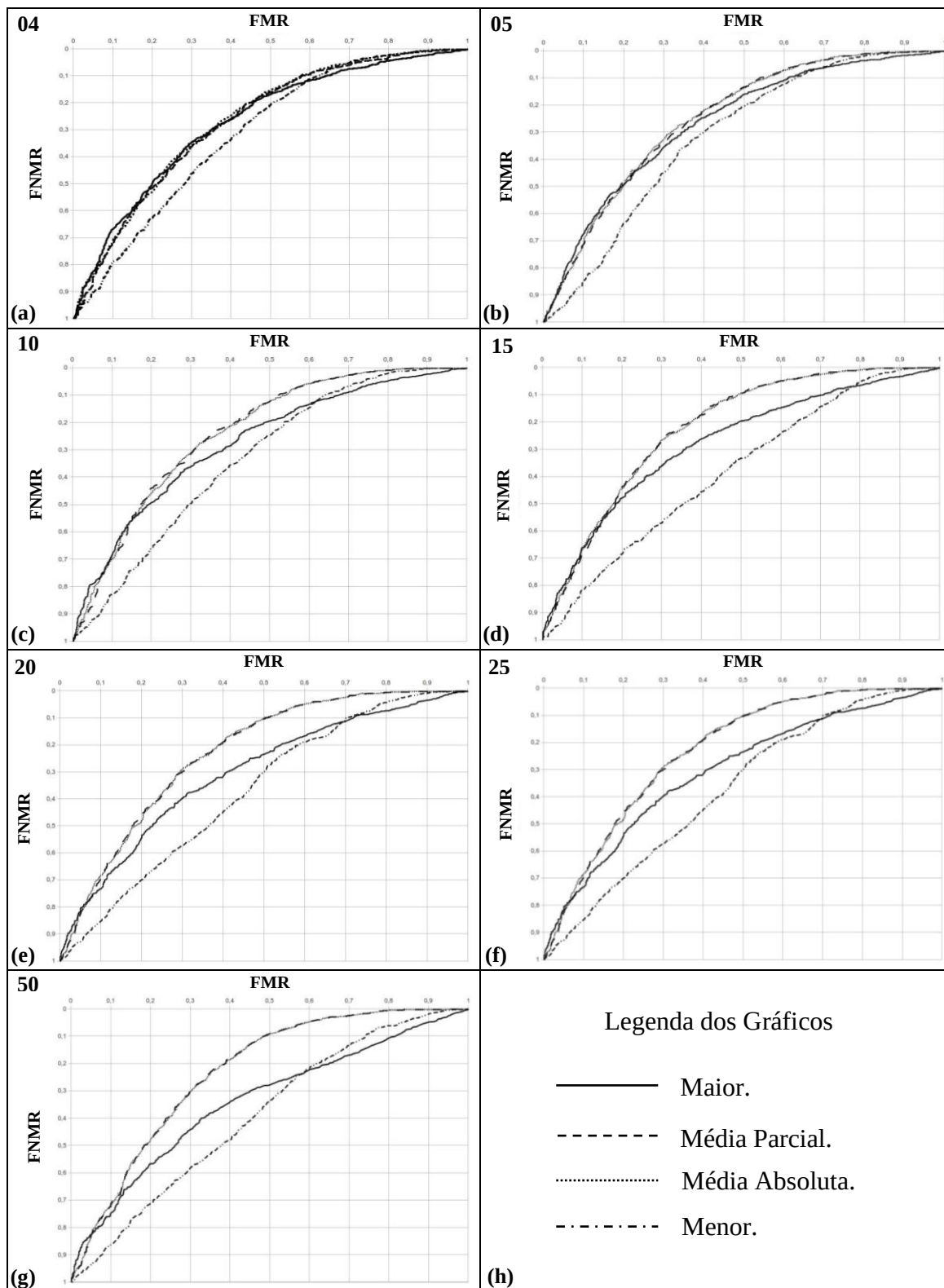


Figura 6.3. Resultados com descritores do método *Eigenfaces* utilizando diferentes métodos de votação com famílias de (a) 04 (b) 05 (c) 10 (d) 15 (e) 20 (f) 25 (g) 50 classificadores.

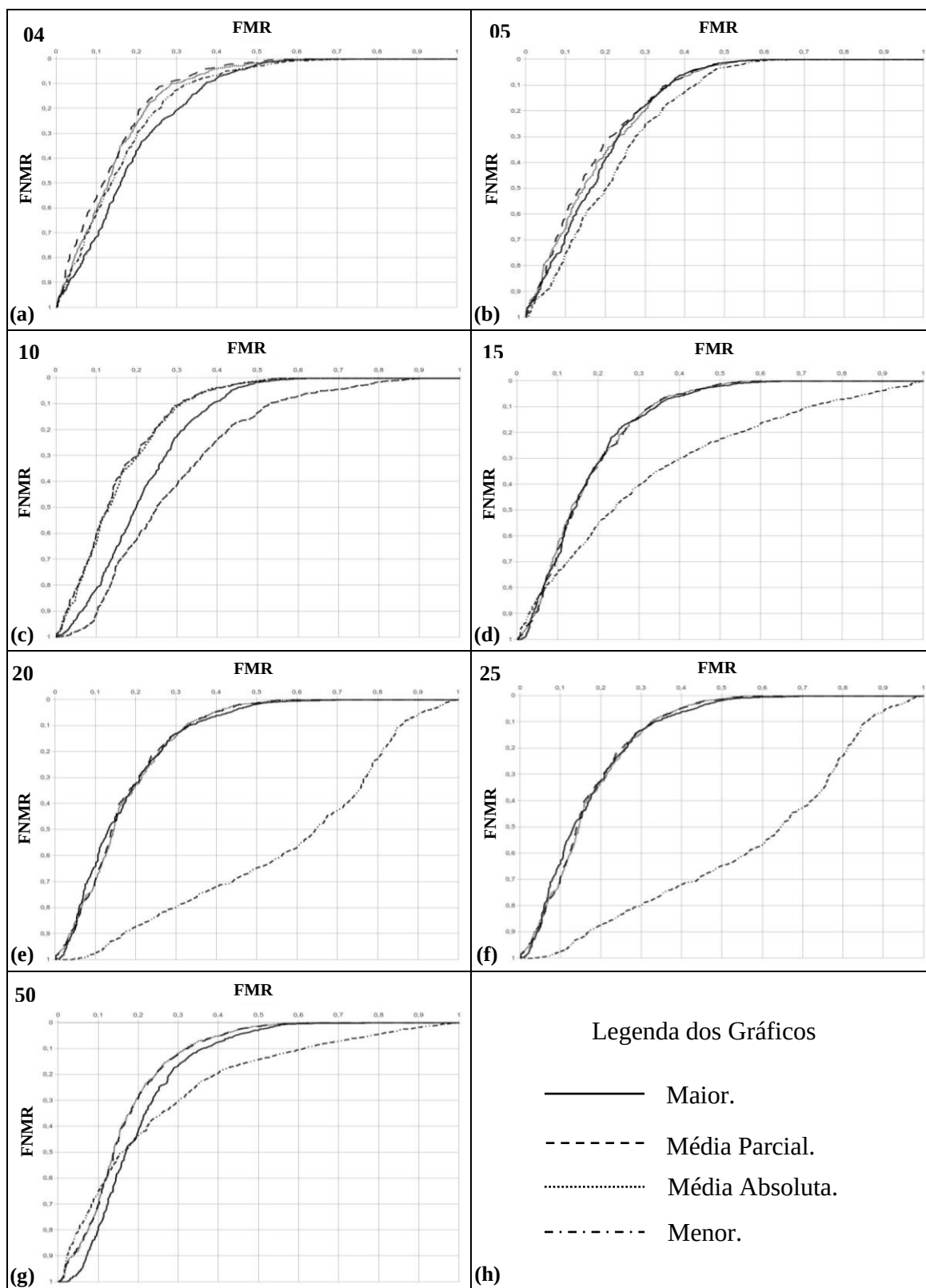


Figura 6.4. Resultados com descritores do método *EBGM* utilizando diferentes métodos de votação com famílias de (a) 04 (b) 05 (c) 10 (d) 15 (e) 20 (f) 25 (g) 50 classificadores.

Os objetivos destes primeiros passos da pesquisa foram selecionar o conjunto de funções de distância e de métodos de votação com melhores resultados de precisão. Nas próximas seções é apresentado o comportamento das diferentes famílias de classificadores utilizando as funções de distância Chi-Squared, Euclidiana e Manhattan e os métodos de votação Média Absoluta e Média Parcial.

Em todos os outros experimentos realizados foram avaliadas todas as funções de distância, mas de maneira consistente os melhores resultados bem diferenciados foram sempre obtidos com essas três funções de distância. Por esta razão, para simplificar os próximos experimentos, a partir de agora todos os resultados serão apresentados unicamente com estas três funções. No anexo A apresentamos outros resultados não relevantes considerando outras funções de distância.

6.4. Experimentos com Famílias de Classificadores

A seguir apresentamos os resultados dos experimentos comparando o comportamento de diferentes tamanhos da família de classificadores. Estes experimentos tem como objetivo verificar a hipótese de que com uma família de classificadores é possível obter melhores resultados de precisão que com um único classificador. Para estes experimentos foram exploradas todas as funções de distância consideradas no trabalho, assim como os métodos de votação. Foram utilizados descritores dos métodos *Eigenfaces* e *EBGM*.

6.4.1. Experimentos com o Método de Votação Média Parcial

A Figura 6.5 mostra as curvas dos resultados dos experimentos considerando o método de votação por média parcial. A coluna da esquerda mostra os resultados para descritores obtidos com o método *Eigenfaces* e a coluna da direita com descritores do método *EBGM*. Cada linha corresponde aos resultados segundo uma função de distância diferente: Euclidiana, Manhattan e Chi-Squared.

Os gráficos (a), (c) e (e) na Figura 6.5, mostram os resultados de precisão com descritores gerados a partir do método *Eigenfaces* e utilizando respectivamente as distâncias

Euclidiana, Manhattan e Chi-Squared.

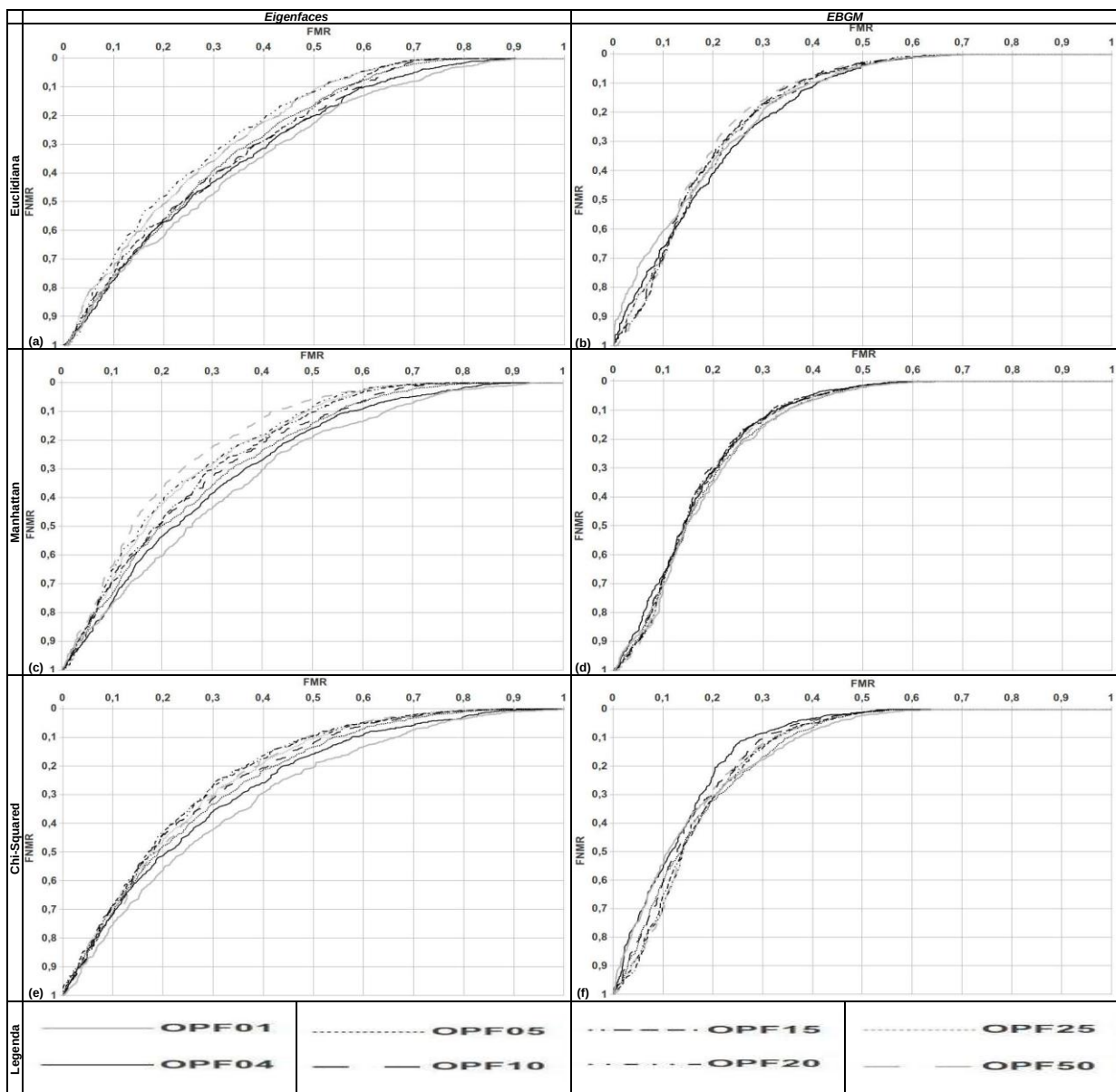


Figura 6.5. Gráficos com os resultados dos experimentos com o método de votação de média parcial. A coluna da esquerda mostra resultados com descritores do método *Eigenfaces* e a coluna da direita com descritores do método *EBGM*. Cada linha da figura mostra os resultados para uma função de distância diferente: Euclidiana, Manhattan e Chi-Squared. Cada gráfico mostra as curvas ROC para famílias com 01, 04, 05, 10, 15, 20, 25 e 50 classificadores.

Quando utilizada a distância *Euclidiana* (Figura 6.5a) e considerando a métrica EER na Tabela 6.3, verificamos que o melhor resultado para as amostras *Eigenfaces* foi obtido com 20 classificadores. Observa-se como o aumento do número de classificadores neste caso melhora

de forma consistente a precisão em todas as métricas.

No caso dos resultados com a distância Manhattan com a média parcial (Figura 6.5c), observa-se também uma melhora da precisão com o aumento do número de classificadores. Para esta função de distância os melhores resultados observados a métrica EER da Tabela 6.4, foram com 50 classificadores para *Eigenfaces*.

Tabela 6.3. Resultados com descritores gerados pelo método *Eigenfaces* e função de distância Euclidiana e votação com média parcial.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,36975	0,850112	0,911633	0,911633
04	0,358217	0,818589	0,924972	0,936170
05	0,335966	0,741321	0,810750	0,837626
10	0,346173	0,699888	0,786114	0,818589
15	0,351248	0,704367	0,770437	0,797312
20	0,311639	0,690929	0,790594	0,793953
25	0,320308	0,699888	0,792833	0,818589
50	0,318176	0,675252	0,745801	0,769317

Tabela 6.4. Resultados com descritores gerados pelo método *Eigenfaces* e função de distância Manhattan e votação com média parcial.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,363684	0,891499	0,945190	0,946309
04	0,337156	0,832027	0,894737	0,896976
05	0,321759	0,755879	0,895857	0,921613
10	0,312097	0,713326	0,839866	0,856663
15	0,301446	0,712206	0,795073	0,836506
20	0,288823	0,701008	0,774916	0,815230
25	0,292776	0,699888	0,806271	0,828667
50	0,265296	0,679731	0,818589	0,826428

Tabela 6.5. Resultados com descritores gerados pelo método *Eigenfaces* e função de distância Chi-Squared e votação com média parcial.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,356029	0,892617	0,994407	0,994407
04	0,330715	0,860022	0,977604	0,977604
05	0,314113	0,799552	0,941769	0,964166
10	0,304577	0,765957	0,875700	0,929451
15	0,288512	0,774916	0,873460	0,954087
20	0,290034	0,761478	0,932811	0,932811
25	0,298775	0,736842	0,917133	0,969765
50	0,301164	0,765957	0,911534	0,946249

Com o uso da distância Chi-Squared (gráfico (e) na Figura 6.5) vemos também de

forma consistente um aumento da precisão. Para os melhores resultados na coluna EER da Tabela 6.5 são obtidos quando utilizados a família de 15 classificadores para amostras.

De modo geral, os resultados para todas as famílias de classificadores foram melhores que utilizando um único classificador quando utilizamos descritores gerados pelo método *Eigenfaces*. Em especial a família de 50 classificadores com amostras *Eigenfaces* está entre as quatro melhores acurácias das oito consideradas na Tabela 6.3, Tabela 6.4 e Tabela 6.5.

No caso de descritores gerados pelo método *EBGM* as Figura 6.5b, Figura 6.5d e Figura 6.5e, mostram os resultados de precisão utilizando respectivamente as distâncias Euclidiana, Manhattan e Chi-Squared.

Quando utilizada a distância Euclidiana o melhor resultado para amostras *EBGM* (Figura 6.5b) é com a família de 50 classificadores. Cada uma das métricas melhora seus resultados quando utilizados vários classificadores, mas em nenhum caso existe uma melhora consistente em todas as métricas ou todos os limiares (Tabela 6.6).

Tabela 6.6. Resultados com descritores gerados pelo método *EBGM* e função de distância Euclidiana e votação com média parcial.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,268718	0,595897	0,694359	0,714872
04	0,266580	0,605749	0,716632	0,732033
05	0,274136	0,612936	0,723819	0,737166
10	0,249755	0,596509	0,704312	0,741273
15	0,249442	0,572895	0,691992	0,715606
20	0,256674	0,600616	0,707392	0,726899
25	0,256674	0,600616	0,707392	0,726899
50	0,236989	0,612936	0,722793	0,738193

Os resultados com a distância Manhattan e amostras *EBGM*, mostrados na Figura 6.5d, não mostram melhores resultados com o uso de vários classificadores. A Tabela 6.7 mostra a família com 15 classificadores com o melhor resultado para a métrica EER. Na maioria dos casos, nas métricas da Tabela 6.7 as abordagens com família de classificadores mostram em resultados melhores que a abordagem com um classificador.

Tabela 6.7. Resultados com descritores gerados pelo método *EBGM* e função de distância Manhattan e votação com média parcial.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,241635	0,544615	0,604103	0,612308
04	0,232775	0,650924	0,650924	0,650924
05	0,243396	0,606776	0,606776	0,644764
10	0,229925	0,628337	0,628337	0,646817
15	0,228743	0,599589	0,599589	0,625257
20	0,232664	0,591376	0,591376	0,641684
25	0,232664	0,591376	0,591376	0,641684
50	0,232396	0,574949	0,590349	0,637577

Com o uso da distância Chi-Squared, a Figura 6.5f mostra que existe uma melhor resultado quando utilizados a família de 04 classificadores. A Tabela 6.8 mostra nas métricas EER que a família com 04 classificadores obtém os melhores resultados. Também verificamos neste experimento; em quase todas as métricas a abordagem com família obteve uma melhora em relação a um único classificador.

Tabela 6.8. Resultados com descritores gerados pelo método *EBGM* e função de distância Chi-Squared e votação com média parcial.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,240013	0,554872	0,678974	0,702564
04	0,204312	0,574949	0,574949	0,596509
05	0,246330	0,606776	0,606776	0,618070
10	0,242553	0,617043	0,617043	0,625257
15	0,234086	0,630390	0,630390	0,633470
20	0,234086	0,630390	0,630390	0,633470
25	0,230457	0,634497	0,634497	0,634497
50	0,226664	0,650924	0,650924	0,650924

Os resultados obtidos nos experimentos com as amostras *EBGM*, também mostram que o uso de vários classificadores melhora em todos os casos os resultados de precisão se comparados com o uso de um único classificador. De modo geral, os resultados para a família de 50 classificadores se apresentam melhor que o OPF com um classificador, na métrica EER para todas as funções de distância.

6.4.2. Experimentos com o Método de Votação Média Absoluta

A Figura 6.6 mostra as curvas dos resultados dos experimentos considerando o método de votação por média absoluta.

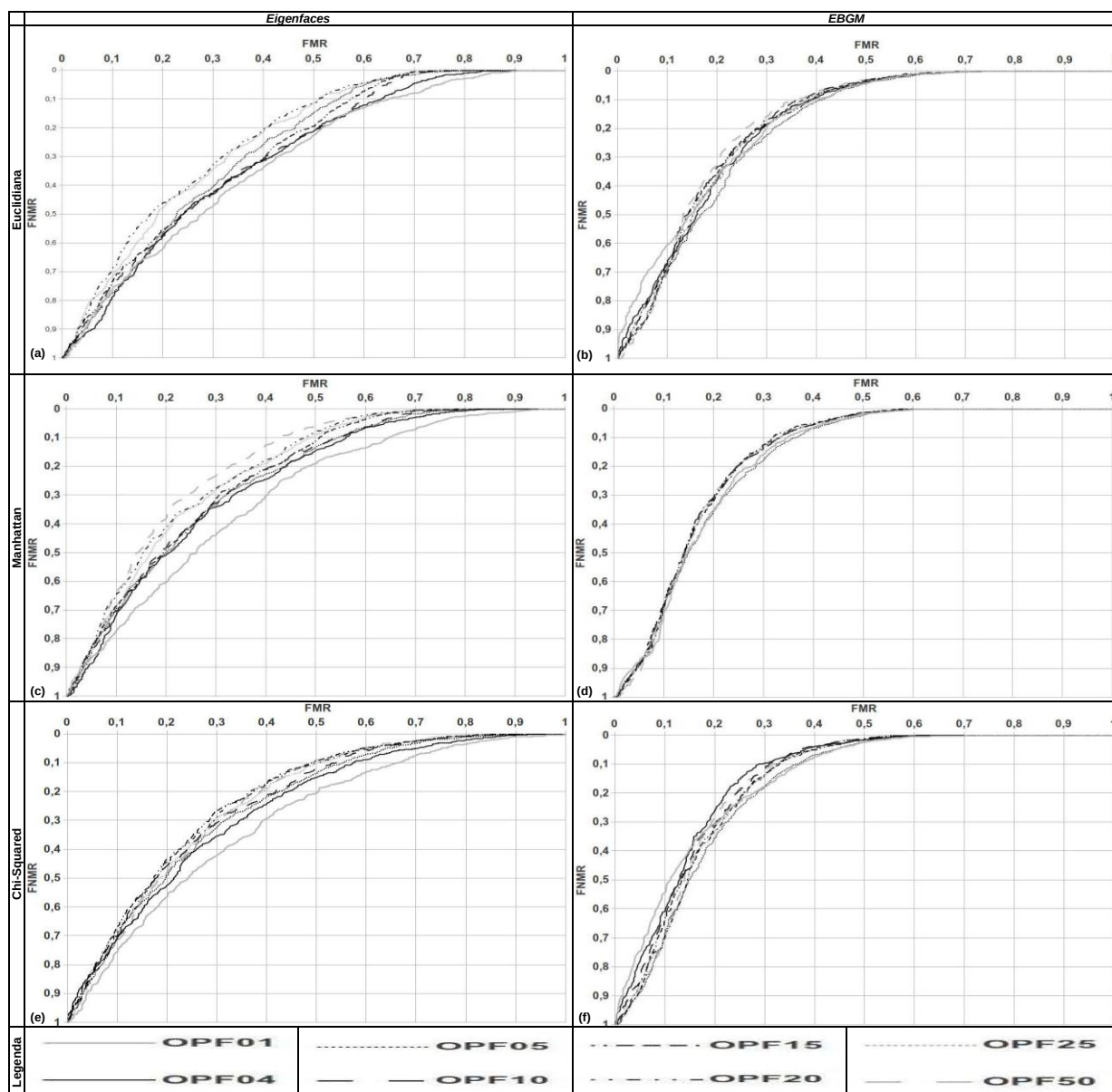


Figura 6.6. Gráficos com os resultados dos experimentos com o método de votação de média absoluta. A coluna da esquerda mostra resultados com descritores do método *Eigenfaces* e a coluna da direita com descritores do método *EBGM*. Cada linha da figura mostra os resultados para uma função de distância diferente: Euclidiana, Manhattan e Chi-Squared. Cada gráfico mostra as curvas ROC para famílias com 01, 04, 05, 10, 15, 20, 25 e 50 classificadores.

Como no caso dos resultados com o método da média parcial, a coluna da esquerda mostra os resultados para descritores obtidos com o método *Eigenfaces* e a coluna da direita com descritores do método *EBGM*.

As Figura 6.6a, Figura 6.6c e Figura 6.6e, assim como as Tabela 6.9, Tabela 6.10 e Tabela 6.11, mostram os resultados da precisão dos classificadores utilizando respectivamente as distâncias Euclidiana, Manhattan e Chi-Squared. Nestes experimentos foram utilizados descritores gerados pelo método *Eigenfaces*, com o método de votação da média absoluta.

Para a distância Euclidiana com as amostras *Eigenfaces* (Figura 6.6a) os melhores resultados são com a família de 20 classificadores na métrica EER. A Tabela 6.9 confirma a família de 50 classificadores nesta métrica. Contudo as famílias de 20 e 50 classificadores apresenta resultado próximos da família com 20 classificadores. Também quando usamos a família de classificadores foi confirmada a melhora da precisão da classificação.

Tabela 6.9. Resultados com descritores gerados pelo método *Eigenfaces* e função de distância Euclidiana e votação com média absoluta.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,369750	0,850112	0,911633	0,911633
04	0,356047	0,798432	0,888018	0,893617
05	0,338040	0,712206	0,810750	0,829787
10	0,350899	0,702128	0,790594	0,829787
15	0,358225	0,706607	0,759239	0,797312
20	0,310190	0,685330	0,772676	0,792833
25	0,319054	0,697648	0,786114	0,818589
50	0,317845	0,674132	0,743561	0,769317

Os resultados na Figura 6.6c mostram que com a distância Manhattan com amostras *Eigenfaces*, a acurácia do algoritmo com 50 classificadores continua sendo a melhor.

A Tabela 6.10 na métrica EER confirma, para este experimento, a família de 50 classificadores como melhor resultado. Todos os classificadores com abordagem por votação, ou seja, a família de classificadores obtiveram melhores resultados que o OPF com um classificador.

Tabela 6.10. Resultados com descritores gerados pelo método *Eigenfaces* e função de distância Manhattan e votação com média absoluta.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,363684	0,891499	0,94519	0,946309
04	0,322739	0,768197	0,832027	0,870101
05	0,310688	0,727884	0,877940	0,905935
10	0,310298	0,706607	0,840985	0,868981
15	0,303970	0,712206	0,811870	0,834267
20	0,287259	0,696529	0,768197	0,820829
25	0,292273	0,701008	0,804031	0,829787
50	0,265851	0,679731	0,816349	0,827548

As curvas da Figura 6.6e, os resultados para a distância Chi-Squared mostram a família com 15 classificadores com melhor acurácia. Na métrica EER da Tabela 6.11, também a família com 15 classificadores apresenta o melhor resultado. A abordagem com a família de classificadores também demonstra melhor acurácia.

Tabela 6.11. Resultados com descritores gerados pelo método *Eigenfaces* e função de distância Chi-Squared e votação com média absoluta.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,356029	0,892617	0,994407	0,994407
04	0,325356	0,860022	0,921613	0,951848
05	0,307751	0,806271	0,93953	0,972004
10	0,303435	0,772676	0,855543	0,93729
15	0,288740	0,759239	0,880179	0,950728
20	0,295227	0,760358	0,909294	0,928331
25	0,295633	0,734602	0,923852	0,973124
50	0,300951	0,765957	0,913774	0,947368

Assim como em 6.4.1 podemos com estes experimentos (6.4.2) ver uma melhora na classificação com o uso da família de classificadores, quando comparamos com o algoritmo OPF com um único classificador.

As Figura 6.6b, Figura 6.6d e Figura 6.6f, e as Tabela 6.12, Tabela 6.13, e Tabela 6.14 mostram os resultados dos classificadores OPFs utilizando a função de distância Euclidiana, Manhattan e Chi-Squared, com o método de votação da média absoluta. A diferença em relação à seção anterior, é que os descritores foram obtidos com o método *EBGM*.

Para a distância Euclidiana (Figura 6.6b) os resultados mostram a família com 50 classificadores com melhor acurácia. Na Tabela 6.12 a coluna EER confirma o resultado do gráfico para a família de 50 classificadores. Os outros classificadores com a proposta de família, também mostram uma melhora na acurácia em relação ao OPF, para esta métrica.

Tabela 6.12. Resultados com descritores gerados pelo método *EBGM* e função de distância Euclidiana e votação com média absoluta.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,268718	0,595897	0,694359	0,714872
04	0,257194	0,611910	0,704312	0,724846
05	0,266121	0,613963	0,721766	0,733060
10	0,248310	0,600616	0,705339	0,735113
15	0,247433	0,577002	0,703285	0,715606
20	0,251507	0,602669	0,708419	0,725873
25	0,251507	0,602669	0,708419	0,725873
50	0,236140	0,619097	0,722793	0,737166

Para a distância Manhattan (Figura 6.6d), a família com 15 classificadores apresenta o melhor resultado, embora as famílias com 20, 25 e 50 classificadores apresentem resultados, que podem ser considerados equivalentes aos de 15 classificadores, para a métrica de avaliação EER. Podemos ver na Tabela 6.13 a confirmação da família com 15 classificadores como melhor resultado na métrica de avaliação EER. Os classificadores com a proposta de família em geral apresentaram melhores resultados.

Tabela 6.13. Resultados com descritores gerados pelo método *EBGM* e função de distância Manhattan e votação com média absoluta.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,241635	0,544615	0,604103	0,612308
04	0,241032	0,635524	0,635524	0,635524
05	0,249193	0,626283	0,626283	0,641684
10	0,231621	0,619097	0,619097	0,643737
15	0,229862	0,588296	0,588296	0,628337
20	0,234535	0,585216	0,589322	0,640657
25	0,234535	0,585216	0,589322	0,640657
50	0,230184	0,578029	0,591376	0,640657

Nos resultados da Figura 6.6f com a distância Chi-Squared, amostras do *EBGM* e método de votação da média absoluta; a família com 04 classificadores obteve o melhor resultado, embora que a família com 50 classificadores, que é um algoritmo teoricamente melhor na questão da programação paralela, obteve resultado praticamente equivalente nesta métrica. A Tabela 6.14 confirma o resultado para a família com 04 classificadores. Os classificadores com a proposta de famílias neste experimento, também ficaram melhores, ou com resultados próximos ao OPF.

Tabela 6.14. Resultados com descritores gerados pelo método *EBGM* e função de distância Chi-Squared e votação com média absoluta.

Classificadores	EER	FMR100	FMR1000	FMRZero
01	0,240013	0,554872	0,678974	0,702564
04	0,216558	0,554415	0,594456	0,618070
05	0,261858	0,573922	0,592402	0,612936
10	0,233847	0,641684	0,641684	0,641684
15	0,242300	0,588296	0,588296	0,635524
20	0,241499	0,626283	0,626283	0,638604
25	0,241499	0,626283	0,626283	0,638604
50	0,224783	0,650924	0,650924	0,650924

Os resultados desta seção, também nos mostram uma melhora na acurácia do classificador OPF quando usamos a proposta com família de classificadores. Com o uso do método de votação média absoluta vemos a família de classificadores com 50, 15 e 20 apresentando resultados consideráveis nos experimentos com várias funções de distância. Estes resultados vêm reforçar os resultados da família de classificadores com o uso da média parcial na seção anterior.

6.5 Análise dos Resultados

Os resultados mostram que o uso de famílias de classificadores melhoram os resultados de precisão em quase todos os experimentos, quando comparados com os resultados com um único classificador.

Concluimos que as funções de distância com melhores resultados para ambos os

métodos de extração de características utilizados, *EBGM* e *Eigenfaces*, são a Euclidiana, Manhattan e Chi-Squared, sendo esta última, com a família de 04 classificadores, a que reporta melhor resultado para a métrica EER.

No caso dos métodos de votação escolhidos a média absoluta e a média parcial reportaram os melhores resultados de precisão, sendo esta última a que se mostrou mais consistente nos resultados.

O método de votação por média parcial com descritores do método *Eigenfaces* melhora em todos os casos a precisão dos algoritmos de forma consistente.

Em experimentos com descritores obtidos pelo método *EBGM* a melhora da precisão não se apresenta de maneira consistente em todas as métricas. No entanto os resultados com 50 classificadores apresenta significantes resultados na quase totalidade dos experimentos, o que é de interesse para a pesquisa, pelo número de votantes utilizados.

Em todos os casos alguma família de classificadores sempre melhora o resultado do EER, que é considerada a principal métrica em reconhecimento biométrico.

Estes resultados mostram que o uso de famílias de classificadores pode potencializar a métrica de avaliação do EER que é a métrica mais observada nos algoritmos biométricos.

A abordagem proposta pode ser utilizada com outros métodos de extração de características faciais, pois o processo de classificação independe da forma como os descritores das imagens diferença são obtidos.

7. Conclusões e Trabalhos Futuros

Esta dissertação apresentou os resultados da pesquisa do uso de uma família de classificadores Floresta de Caminhos Ótimos (OPF) aplicados aos métodos de reconhecimento facial.

Para a pesquisa foi utilizada a abordagem de classificação baseada em duas classes: Intrapessoal e Interpessoal. Estas classes são definidas de descritores extraídos da diferença de imagens de iguais e diferentes indivíduos respectivamente. Como métodos de extração de características foram utilizados os métodos *Eigenfaces* e *EBGM* para os quais existem bibliotecas disponíveis de fácil acesso.

A proposta original do algoritmo classificador OPF foi modificada para converter resultados de classificação em valores de similaridade como outras propostas da literatura.

Foram pesquisadas várias funções de distância para os descritores dos métodos de extração de características e vários métodos de votação para integrar os resultados obtidos por cada classificador da família. A combinação da função de distância Chi-Squared com o método de votação média parcial mostraram os melhores resultados.

Os experimentos mostraram que os resultados do classificador OPF podem ser melhorados se utilizada uma família no lugar de um único classificador. Particularmente a métrica EER mostra melhoras em todos os testes apresentados.

Estes resultados mostram que a proposta pode ser uma alternativa a uma solução paralela de reconhecimento facial. Cada classificador da família pode ser executado de forma independente e seu custo de processamento é menor que se utilizado um único classificador.

Concentramos nossos objetivos neste trabalho nos experimentos de comparação entre o classificador OPF e as famílias de Classificadores OPF, ou seja, não objetivamos aqui, comparar os resultados obtidos com os resultados de precisão do estado da arte na área.

Outro ponto importante em relação aos métodos de extração de características utilizados: *Eigenfaces* e *EBGM*, é que, embora sejam métodos bastante utilizados para comparar a acurácia de classificadores, também não foi nosso objetivo verificar a qualidade dos descritores produzidos, em relação a outros métodos de extração de características. No entanto,

a dissertação apresenta uma abordagem que pode ser facilmente aplicável a outros métodos de extração de características.

Trabalhos Futuros

Como trabalhos futuros associados a esta dissertação propomos estender a proposta considerando métodos de votação ponderada, isto é, atribuindo na votação pesos diferenciados aos diferentes classificadores da família. Os métodos para definição dos pesos dos classificadores serão assunto do trabalho, mas podemos sugerir que seja utilizado, por exemplo, uma avaliação de acurácia por alguma das métricas de avaliação e, com base nesta avaliação pode ser definido os pesos.

Também propomos pesquisar a precisão da abordagem utilizando outros métodos de extração de características faciais, para verificar a acurácia do método de família com dados que possam apresentar outras características biométricas. Ainda propomos pesquisar a precisão da proposta atual, utilizando métodos que caracterizam quantitativamente a diferença entre duas imagens faciais.

Outra proposta de nosso interesse é pesquisar o resultado com outras funções de distância e métodos de votação, que na pesquisa atual não houve condições de serem pesquisadas.

8. Referências

- Afonso, L. C., Papa, J. P., Marana, A. N., Poursaberi, A., Yanushkevich, S. N. & Gavrilova, M. 2012. Optimum-Path Forest Classifier for Large Scale Biometric Applications, Emerging Security Technologies (EST), Third International Conference, 58-61.
- Alpaydin, E. 2010. Introduction to Machine, Massachusetts – USA, The Massachusetts Institute of Technology Press.
- Amaral, V. & Thomaz, C. E. 2011. Extração e Comparação de Características Locais e Globais para Reconhecimento Automático de Imagens de Faces, São Bernardo do Campo - SP – Brasil, XII Workshop de Visão Computacional (WVC).
- Beveridge, J. R., Bolme, D., Draper, B. A. & Teixeira, M. 2005. The CSU Face Identification Evaluation System, Machine Vision and Applications, 16(2), pp 128–138.
- Bonato, C. S. & Neto, R. M. F. 2010. Um Breve Estudo Sobre Biometria, Campus Catalão – Go - Brasil, Departamento de Ciência da Computação da Universidade Federal de Goiás (UFG).
- Bray, J. R., Curtis, J. T. 1957. An Ordination of the Upland forest Communities of Southern Wisconsin, Ecological Monographs, 27(4), 325-349.
- Cappabianco, F. A. M., Falcão, A. X., Yasuda, C. L. & Udupa, J. K. 2012. Brain Tissue MR-Image Segmentation Via Optimum-Path Forest Clustering, Computer Vision and Image Understanding, 116(10), 1047–1059.
- Chen, D., Cao, X., Wang, L., Wen, F. & Sun, J. 2012. Bayesian Face Revisited: A Joint Formulation. European Conference on Computer Vision (ECCV'12), 7574, 566-579.
- Costa, L. R., Obelheiro, R. R. & Fraga, J. S. 2006. Introdução à Biometria, Livro textos dos Minicursos do VI Simpósio Brasileiro de Segurança da Informação e de Sistemas Computacionais, 1, 103-151.
- Costa, R. M. 2009. Nova Abordagem para Reconhecimento Biométrico Baseado em Características Dinâmicas da Iris, São Carlos - SP – Brasil, Escola de Engenharia de São Carlos da Universidade de São Paulo.

- Cormen, T. H., Leiserson, C. E., Rivest, R. L. & Stein, C. 2001. Introduction to Algorithms (vol 6), Cambridge - USA, MIT press.
- Daugman, J. 2003. The Importance of Being Random: Statical Principles of Iris Recognition. Pattern Recognition, 36(2), 279-291.
- Dean J., & Ghemawat, S. 2008. Map Reduce: Simplified Data Processing on Large Cluster. Communications of the ACM, 51(11), 107-113.
- Falcão, A. X., Stolfi, J. & Lotufo, R. 2004. The Image Foresting Transform: Theory, Algorithms, and Applications. IEEE Transactions on Pattern Analysis and Machine Intelligence, 26(1), 19-29.
- Falguera, F. P. S. 2008. Fusão de Métodos Baseados em Minúcias e em Cristas para Reconhecimento de Impressões Digitais, Presidente Prudente – SP - Brasil, Universidade Estadual Paulista Julio de Mesquita Filho.
- Graves, W. 2010. The World's Largest Biometric Application, Gaithersburg - USA, International Biometric Performance Testing Conference - IBPT'10.
- Jain, A. K., Hong, L. & Pankanti, S. 2000. Biometrics Identification, Communications of the ACM, 43(2), 90-98.
- Jain, A. K., Bolle, R. & Pankanti, S. 2006. Biometrics Personal Identification in Networked Society, New York - USA, Springer Science & Business Media.
- Jain, A. K., Flynn, P. & Ross, A. A. 2007. Handbook of Biometrics, New York – USA, Springer Science Business Media.
- Jain, A. K., Nandakumar, K. & Nagar, A. 2008. Biometric Template Security, EURASIP Journal on Advances in Signal Processing, 113.
- Kohlwey, E., Sussman, A., Trost, J. & Maurer, A. 2011. Leveraging the Cloud for Big Data Biometrics: Meeting the Performance Requirements of the Next Generation Biometric Systems, IEEE World Congress on Services, 597-601.
- Lance, G. N., Williams, W. T. 1966. Os programas de computador para a classificação politética hierárquica (análise de similaridade). Journal Computer, 9(1), 60-64.

- Li, Z. & Tang, X. 2004. Bayesian Face Recognition Using Support Vector Machine and Face Clustering, IEEE Computer Society Conference, 2, 374-380.
- Lu, X. 2003. Image Analysis for Face Recognition, Michigan State - USA, Dept. of Computer Science & Engineering Michigan State University.
- Marana, A. N., Papa, J. P. & Chiachia, G. 2009. Análise de Desempenho de Classificadores Baseados em Redes Neurais, Máquinas de Vetores de Suporte e Florestas de Caminhos Ótimos para o Reconhecimento de Dígitos Manuscritos, Bauru - SP - Brasil, UNESP – Faculdade de Ciências – Departamento de Computação.
- Moghaddam, B. & Pentland, A. 1997. Probabilistic Visual Learning for Object Representation, IEEE Transactions on Pattern Analysis and Machine Intelligence, 19(7), 696–710.
- Moghaddam, B. 1999. Principal Manifolds and Bayesian Subspaces for Visual Recognition. The Proceedings of the Seventh IEEE International Conference, 2, 1131-1136.
- Moghaddam, B., Jebara, T. & Pentland, A. 2000. Bayesian Face Recognition. Pattern Recognition, 33(11), 1771-1782.
- Naldi, M. C., Faceli, K. & Carvalho, A. 2009. Uma Revisão Sobre Combinação de Agrupamentos. Revista de Informática Teórica e Aplicada, 16(2), 25-52.
- Panchumorthy, R. , Subramanian, R. & Sarkar, S. 2012. Biometric Evaluation on the Cloud: A Case Study with Humanid Gait Challenge, IEEE Fifth International Conference on Utility and Cloud Computing (UCC'12), 219-222.
- Papa, J. P. 2008. Classificação Supervisionada de Padrões utilizando Floresta de Caminhos Ótimos, Campinas - SP - Brasil, Universidade Estadual de Campinas - Instituto de Computação.
- Papa, J. P., Falcão, A. X. & Suzuki, C. T. N. 2009a. Supervised Pattern Classification Based on Optimum-Path Forest. International Journal of Imaging Systems and Technology. 19(2), 120-131.
- Papa, J. P., Falcão, A., Levada, L., Corrêa D., Salvadeo, D. H. & Mascarenhas, N. D. 2009b. Fast and Accurate Holistic Face Recognition using Optimum-Path Forest, 16th International

- Conference on Digital Signal Processing, 1-6.
- Papa, J. P., Cappabianco, F. A. M. & Falcão, A. X. 2010. Optimizing Optimum-Path Forest Classification for Huge Datasets, Pattern Recognition (ICPR) - 20th International Conference on Pattern Recognition, 4162-4165.
- Papa, J. P., Falcão, A. X., Albuquerque, V. H. C. & Tavares, J. M. R. 2012. Efficient Supervised Optimum-Path Forest Classification for Large Datasets, Pattern Recognition, 45(1), 512–520.
- Papa, J. P., Falcão, A. X. & Suzuki, C. T. N. 2014. LibOPF: A library for the Design of Optimum-path forest Classifiers, disponível em <http://www.ic.unicamp.br/~afalcao/LibOPF>.
- Petry, A., Soares, S. S. & Barone, D. A. C. 2004. Sistema para Autenticação de Usuários por Voz em Redes de Computadores. Workshop em Segurança de Sistemas Computacionais, 4, 263-272.
- Phillips, P. J., Moon, H., Rizvi, S. A. & Rauss, P. J. 2000. The FERET Evaluation Methodology for Face Recognition Algorithms, IEEE Transactions on Pattern Analysis and Machine Intelligence, 22(10), 1090-1104.
- Phillips, P. J., Grother, P., Micheals, R. J., Blackburn D. M., Tabassi, E. & Bone, M. 2003. Face Recognition Vendor Test 2002, Analysis and Modeling of Faces and Gestures (AMFG) – IEEE International Workshop, 44.
- Ponti Jr, M. P. & Papa, J. P. 2011. Improving Accuracy and Speed of Optimum-Path Forest Classifier using Combination of Disjoint Training Subsets. International Workshop on Multiple Classifier Systems, 237-248.
- Prodossimo, F. C., Chidambaram, C., Hastreiter, R. L. F. & Lopes, H. S. 2012. Proposta de uma Metodologia para a Construção de um Banco de Imagens Faciais Normalizadas, PIE, 5, 41368.
- Rocha, L. M., Cappabianco, F. A. M. & Falcão A. X. 2009. Data Clustering as an Optimum-Path Forest Problem with Applications in Image Analysis, International Journal of Imaging Systems and Technology, 19(2), 50-68.

- Silva, A. B. N. 2013. Reconhecimento Facial utilizando *Eigenfaces*, Rio de Janeiro - RJ – Brasil, Instituto Alberto Luiz Coimbra de Pós-Graduação e Pesquisa de Engenharia.
- Tan, X., Chen, S., Zhou, A. H. & Zhang, F. 2005. Recognizing Partially Occluded Expression Variant Faces from Single Training Image per Person with SOM and Soft k-NN Ensemble, *IEEE Transactions on Neural Networks*, 16(4), 875-886.
- Tan, X., Chen, S., Zhou, Z. H. & Zhang, F. 2006. Face Recognition from a Single Image per Person: A survey, *Pattern Recognition*, 39(9), 1725-1745.
- Theisen, E. M. K. 2008. Confiabilidade no Processo de Rastreabilidade por meio de Análise de Imagens Retinais, RS - RS - Brasil, Universidade Santa Cruz do Sul.
- Theodoridis, S. & Koutroumbas, K. 2009. Feature Selection, *Pattern Recognition*, 213-262.
- Turk, M. A. & Pentland, A. S. 1991a. Eigenfaces for Recognition. *Journal of Cognitive Neuroscience*, 3(1), 71-86.
- Turk, M. A. & Pentland, A. S. 1991b. Face Recognition using Eigenfaces, *Proc. Computer Vision and Pattern Recognition (CVPR'91) - IEEE Computer Society Conference*, 586-591.
- Wang, X. & Tang, X. 2003. Bayesian Face Recognition using Gabor Features, *Proceedings of the 2003 ACM SIGMM workshop on Biometrics Methods and Applications (WBMA'03)*, 70-73.
- Wiskott, L., Fellous, J. M., Kuiger, N. & Von Der Malsburg, C. 1997. Face Recognition by Elastic Bunch Graph Matching. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 19, 775-779.
- Zhang, D., Jing, X. & Yang, J. 2006. *Biometric Image Discrimination Technologies*, Hershey – Pensilvânia - EUA, Idea Group Publishing.
- Zhao, W., Chellappa, R., Phillips, P. J. & Rosenfeld, A. 2003. Face Recognition: A Literature Survey, *ACM Computing Surveys (CSUR)*, 35(4), 399-458.

9. Apêndices A. Gráficos para a comparação do OPF com abordagem por votação utilizando várias funções de distância

O apêndice A mostra os resultados com vários tamanhos de famílias com funções de distância não apresentadas na dissertação. A Figura 9.1 mostra os resultados para descritores do método *Eigenfaces* e a Figura 9.2 para descritores do método *EBGM*. Note-se nas figuras que os resultados de precisão para as distâncias Cambera , Bray-Curtis , Squared Chord estão muito distantes das consideradas no texto.

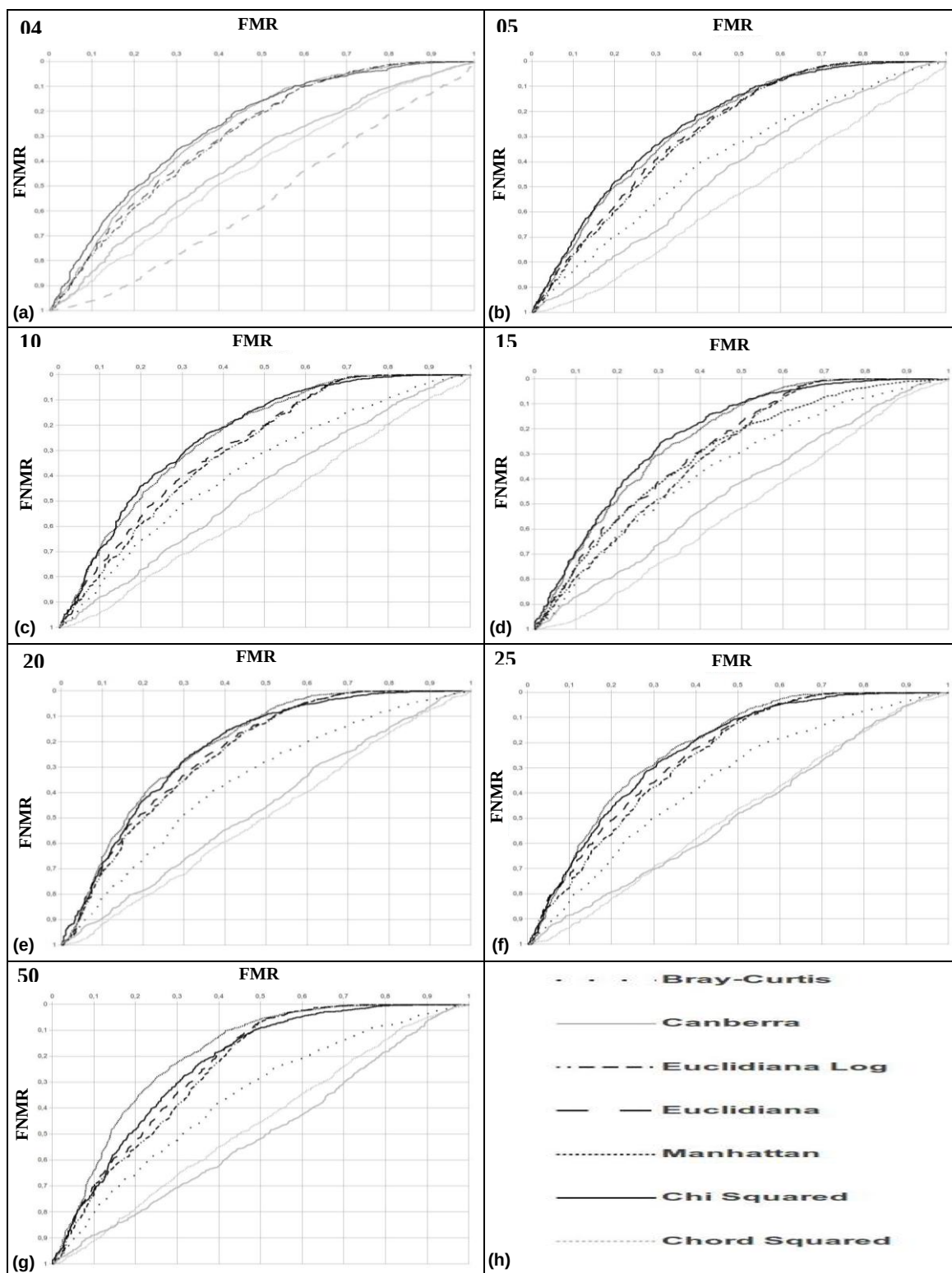


Figura 9.1. OPF com abordagem por votação com descritores *Eigenfaces*.

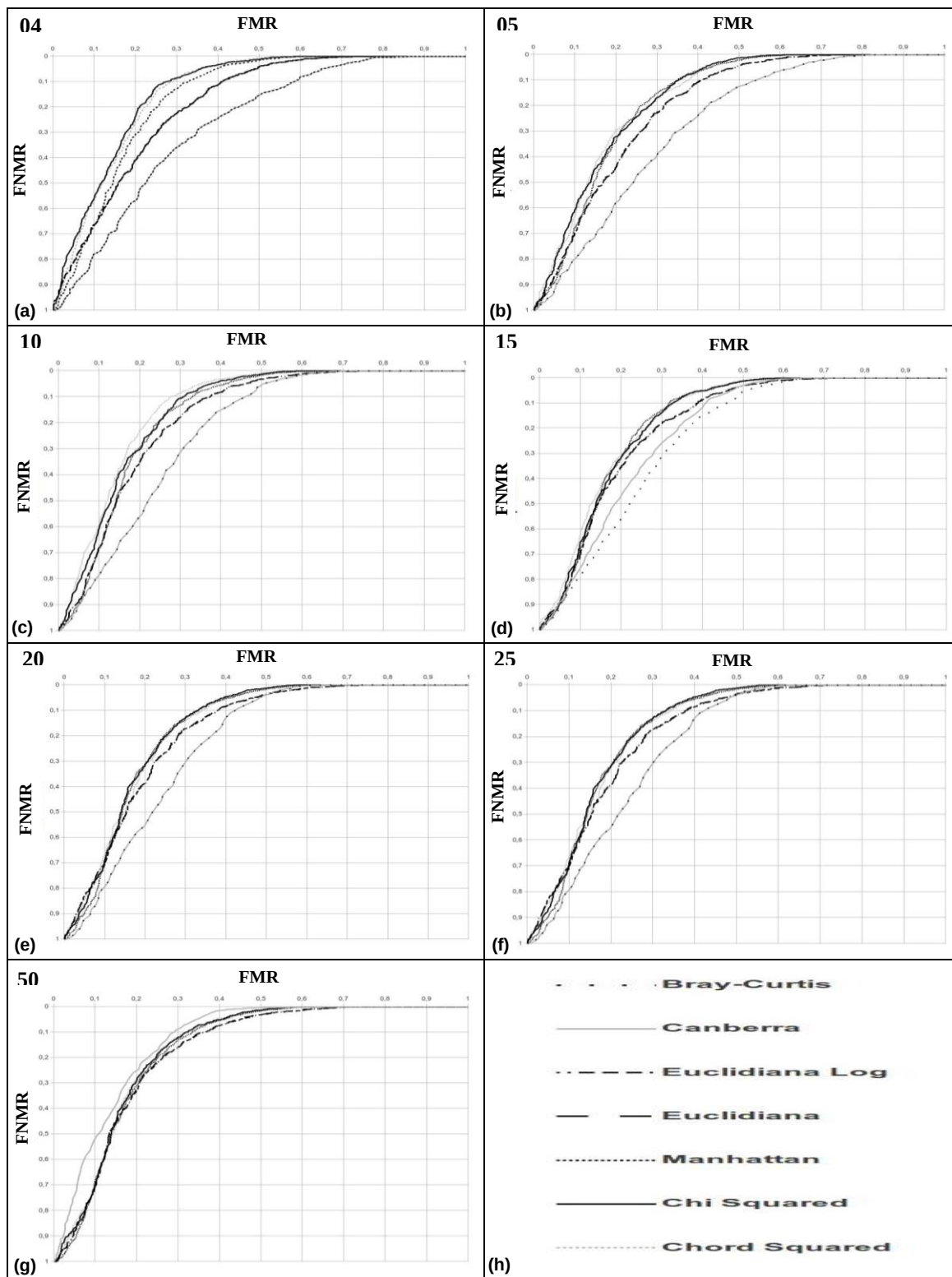


Figura 9.2. OPF com abordagem por votação com descritores *EBGM*.