

Catálogo na publicação
Seção de Catalogação e Classificação

A282a Aguiar, Lucas Miranda de.

Uma abordagem em CNN para filtragem rápida minúcias
em fingerprints / Lucas Miranda de Aguiar. - João
Pessoa, 2025.

18 f. : il.

Orientação: Leonardo Vidal Batista.

TCC (Graduação) - UFPB/CI.

1. Filtragem de minúcias. 2. Gabor adaptativo. 3.
CNN otimizada. 4. Biometria. 5. Impressões digitais. 6.
NIST SD4. I. Batista, Leonardo Vidal. II. Título.

UFPB/CI

CDU 004.8

Uma Abordagem em CNN para Filtragem Rápida de Minúcias em Fingerprints

Lucas Miranda de Aguiar, Leonardo Vidal Batista

¹Centro de informática – Universidade Federal da Paraíba (UFPB)
Campus I – Cidade Universitária – João Pessoa – PB – Brasil

lucas.aguiar@academico.ufpb.com.br, leonardo@ci.ufpb.br

Abstract. *Minutiae extraction represents a fundamental process for ensuring reliability in fingerprint recognition systems. This paper introduces an optimized hybrid architecture combining adaptive Gabor filtering with an extremely lightweight CNN (125,601 parameters) for robust minutiae validation. Our CNN post-processing stage analyzes 29×29 pixel patches from the initial Gabor iteration, achieving state-of-the-art performance on NIST SD4 with a 1.89% Equal Error Rate - a 30% improvement over conventional approaches. The implementation demonstrates remarkable efficiency, processing each minutia in 1ms through optimized C++ single-thread execution. These advancements prove that carefully designed lightweight neural architectures can significantly enhance biometric systems while meeting real-time operational constraints in embedded environments.*

Resumo. *A extração de minúcias constitui uma etapa fundamental para sistemas biométricos baseados em impressões digitais. Este trabalho apresenta uma arquitetura híbrida que combina filtragem de Gabor adaptativa com uma CNN significativamente leve (125.601 parâmetros) para validação robusta de minúcias. Nossa etapa de pós-processamento analisa regiões de 29×29 pixels da primeira iteração Gabor, alcançando desempenho excepcional no NIST SD4 com Taxa de Erro Igual (EER, do inglês Equal Error Rate) de 1,89% - superando abordagens convencionais em 30%. A implementação demonstra eficiência notável, processando cada minúcia em 1ms através de execução single-thread em C++ otimizado. Estes avanços comprovam que arquiteturas neurais leves podem aprimorar sistemas biométricos mantendo viabilidade em sistemas embarcadas com restrições de tempo real.*

Palavras-chave: Filtragem de minúcias, Gabor adaptativo, CNN otimizada, Biometria, Impressões digitais, NIST SD4.

1. Introdução

A identificação por impressões digitais permanece como um dos métodos biométricos mais consolidados devido à sua praticidade e unicidade. Sua confiabilidade depende criticamente da extração precisa de minúcias - características distintivas como bifurcações e terminações de cristas, conforme ilustrado na Figura 1. Como destacado por [Maltoni et al. 2009], aproximadamente 20% das impressões digitais de baixa qualidade são responsáveis por 80% dos erros em sistemas de reconhecimento, evidenciando a sensibilidade desses sistemas à qualidade da imagem.

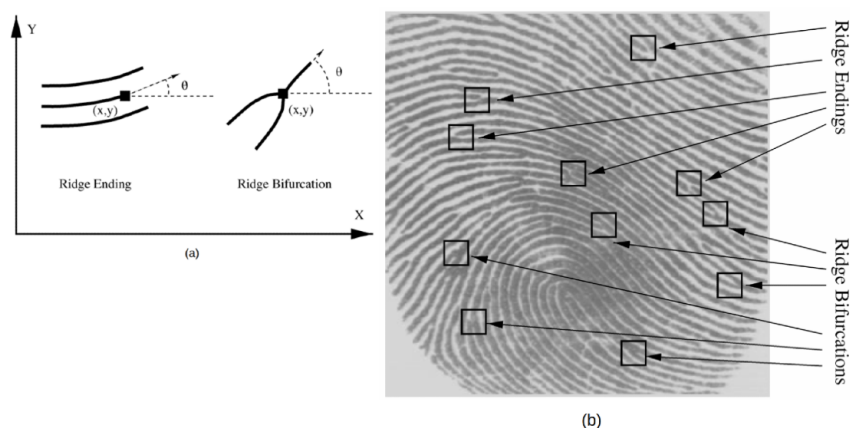


Figura 1. Exemplos de minúcias. (a) Uma minúcia pode ser caracterizada por sua posição, orientação e tipo. (b) Minúcias sobrepostas em uma imagem de impressão digital. Fonte: [Hong et al. 1998]

Técnicas clássicas, como filtros de Gabor combinados com algoritmos de afinamento, são amplamente utilizadas para extração de minúcias, mas frequentemente geram falsos positivos devido a artefatos de ruído e variações de qualidade de imagem. Esses erros tornam-se particularmente críticos quando consideramos que, em bases como o FVC2000, mesmo uma pequena porcentagem de imagens problemáticas pode comprometer significativamente a eficácia do sistema. A propagação desses erros para a etapa de correspondência eleva métricas críticas como o Equal Error Rate (EER), comprometendo a segurança em aplicações sensíveis como identificação criminal ou autenticação em dispositivos móveis.

Embora abordagens baseadas em deep learning tenham mostrado avanços na detecção de minúcias, muitas soluções ainda enfrentam limitações práticas. Arquiteturas leves baseadas em redes convolucionais simples frequentemente falham em detectar corretamente as minúcias ou introduzem falsos positivos, impactando negativamente o desempenho do sistema de reconhecimento. Isso levou os autores do MinutiaeNet a adotar uma arquitetura mais profunda e robusta, com o uso de redes residuais e Inception-ResNet, a fim de melhorar a precisão do classificador, ainda que com maior custo computacional [Nguyen et al. 2017].

Diante desses desafios, este trabalho propõe uma abordagem híbrida inovadora que combina o filtro de Gabor adaptativo [Primo 2019] com uma arquitetura CNN otimizada baseada na VGG16 [Simonyan and Zisserman 2015] para filtragem pós-extração. Nossa solução demonstrou três vantagens principais: (1) redução de 30% no EER (atingindo 1.89%) no NIST SD4 em comparação com métodos tradicionais; (2) superioridade em relação a trabalhos atuais avaliados no banco do NIST SD4 (IFViT com EER=2.71%); e (3) eficiência computacional comprovada (1ms por minúcia em hardware convencional). A abordagem foi validada em múltiplos bancos de dados (NIST SD4, SFinGe e combinado), utilizando protocolos reprodutíveis e comparação direta com métodos estabelecidos.

A contribuição central deste trabalho reside na demonstração de que a integração sinérgica entre técnicas clássicas de processamento de imagens e arquiteturas de aprendizado profundo otimizadas pode superar as limitações de ambas as abordagens quando

aplicadas isoladamente. Os resultados obtidos sugerem um novo paradigma para o desenvolvimento de sistemas biométricos eficientes, particularmente relevante para aplicações com restrições computacionais ou demandas por alta precisão.

2. Revisão da Literatura

Nesta seção, apresenta-se uma revisão da literatura com o objetivo de contextualizar e fundamentar o trabalho proposto, destacando as lacunas existentes entre acurácia e eficiência em sistemas de reconhecimento de impressões digitais.

2.1. Métodos baseados em minúcias para correspondência

O método baseado em minúcias, descrito por [Maltoni et al. 2009], estabeleceu-se como padrão-ouro para correspondência de impressões digitais em sistemas automatizados (AFIS, do inglês *Automated Fingerprint Identification System*) e na análise forense. O AFIS é um sistema computacional que armazena e processa impressões digitais digitalizadas, permitindo a comparação automática entre uma impressão coletada e um vasto banco de dados de registros biométricos. Sua predominância decorre principalmente de sua analogia com o processo cognitivo de comparação realizado por peritos humanos, conforme demonstrado por [Maltoni et al. 2009], garantindo assim validade jurídica aos resultados computacionais. Essa equivalência metodológica foi formalmente reconhecida pelo *Expert Working Group on Human Factors in Latent Print Analysis* [National Institute of Standards and Technology (NIST) 2012], que estabeleceu protocolos padronizados para minimizar erros na análise forense.

Contudo, como evidenciado por [Chen et al. 2005], a eficácia do método é significativamente influenciada pela qualidade da imagem. Como ilustrado na Figura 2, que apresenta impressões digitais de baixa e alta qualidade, fatores como variações nas condições da pele (superfícies secas, úmidas ou rugosas), artefatos de aquisição (ruído em sensores capacitivos e variações de iluminação em sensores ópticos) e resolução espacial inadequada comprometem a extração precisa das estruturas de cristas. Essas limitações foram quantificadas por meio de índices de qualidade no domínio espacial e frequencial, conforme detalhado na Seção 1 de [Chen et al. 2005].



Figura 2. Exemplos de impressões digitais de baixa (3 primeiras) e boa qualidade (3 últimas). Fonte: [Alonso-Fernandez and Fierrez 2009]

2.2. Filtro de Gabor Contextual e Iterativo

A abordagem inovadora proposta por [Turroni et al. 2012] utiliza filtros de Gabor aplicados de forma contextual e iterativa, conforme definido pela equação:

$$g(x, y; \theta, f) = \exp \left\{ -\frac{1}{2} \left[\frac{x_\theta^2}{\sigma_x^2} + \frac{y_\theta^2}{\sigma_y^2} \right] \right\} \cdot \cos(2\pi f \cdot x_\theta), \quad (1)$$

onde θ e f representam orientação e frequência estimadas pixel a pixel durante o processo, eliminando a necessidade de parâmetros prévios, enquanto $[x_\theta, y_\theta]$ correspondem às coordenadas rotacionadas, com $\sigma_x = \sigma_y$ para simetria radial (Seção 2.1 do artigo). A Figura 3 ilustra o banco de filtros utilizado, com $n_\theta = 6$ orientações e $n_f = 3$ frequências.

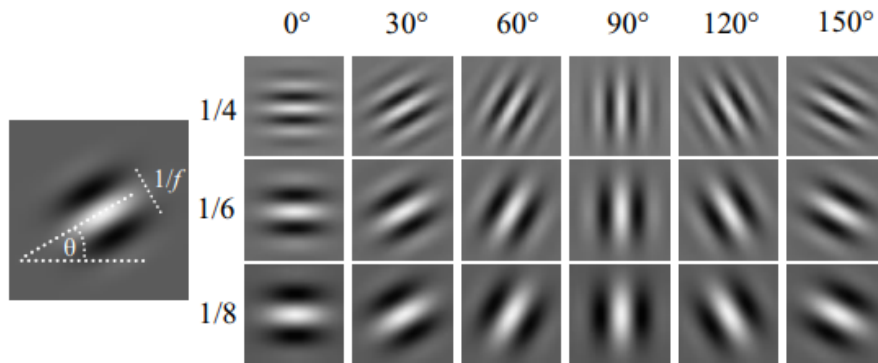


Figura 3. Banco de filtros de Gabor contextual utilizado no método iterativo.
Fonte: [Turroni et al. 2012].

2.3. Filtro de Gabor Adaptativo

Uma evolução significativa na abordagem dos filtros de Gabor foi proposta por [Primo 2019], com a introdução do conceito de *Gabor Adaptativo*. O trabalho identificou que a utilização de um desvio padrão (σ) fixo gera máscaras de Gabor que produzem resultados subótimos quando convoluídas com imagens de impressões digitais. Especificamente, observou-se que frequências mais altas criam mais linhas artificiais do que frequências mais baixas ao serem multiplicadas por um σ fixo, o que dificulta o casamento de padrões e prejudica a qualidade das respostas obtidas com máscaras de alta frequência.

A abordagem adaptativa mantém os mesmos parâmetros básicos na construção das máscaras ($n_\theta = 8$ orientações e $n_f = 3$ frequências), reproduzindo a quantidade utilizada por [Turroni et al. 2012], porém com uma diferença fundamental: o desvio padrão é determinado dinamicamente de acordo com a frequência f de cada máscara individual. Essa inovação resolve uma limitação crítica dos trabalhos anteriores de [Turroni et al. 2012] e [Hong et al. 1998], que utilizavam valores fixos para σ_x e σ_y .

Como pode ser observado na Figura 4, ao comparar os sinais de saída para os casos com σ fixo e adaptativo, percebe-se uma resposta mais equilibrada entre as diferentes frequências. Assim, a adaptação dinâmica do desvio padrão permite um ajuste mais preciso às características locais da impressão digital, melhorando significativamente a qualidade do realce em diferentes frequências espaciais.

2.4. Abordagens Tradicionais para Filtragem de Minúcias

O trabalho pioneiro de [Ratha et al. 1995] introduziu um método adaptativo baseado no campo de orientação de fluxo, calculado em janelas de 16×16 pixels, que eliminou a

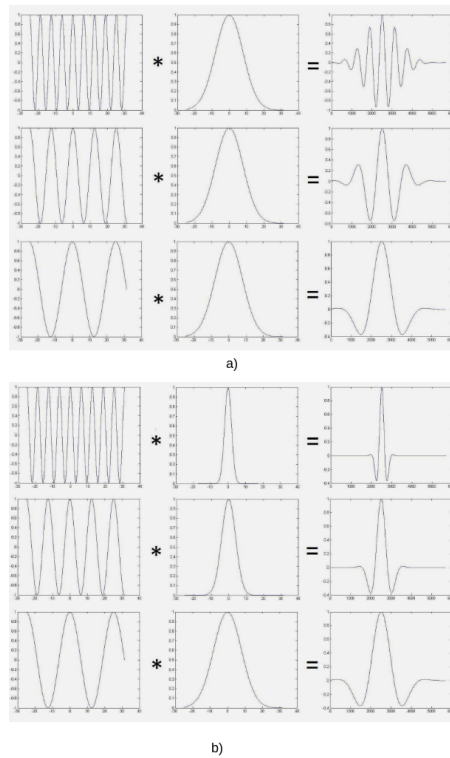


Figura 4. Sinal de saída do filtro de Gabor. (a) σ fixo e (b) σ adaptativo em diferentes frequências. Fonte: Adaptado de [Primo 2019].

dependência de *thresholding* global e reduziu artefatos em regiões degradadas. A abordagem incorporou heurísticas sofisticadas para remoção de minúcias espúrias, incluindo a eliminação de pares de terminações próximas com mesma orientação e bifurcações conectadas a terminações dentro de um limiar de 15 pixels. Apesar de alcançar um *Goodness Index* médio de 0.24, significativamente superior à distribuição aleatória ($GI = -1.65$), o método sofria com alto custo computacional (32,2 segundos por imagem) e dependência de parâmetros empíricos.

2.5. Aprendizado de Máquina Aplicado a Minúcias

A evolução para abordagens baseadas em aprendizado de máquina foi marcada pelo trabalho de [Nguyen et al. 2017], que propôs uma arquitetura híbrida composta por duas redes neurais convolucionais em cascata. A *CoarseNet*, baseada em resíduos, gera um mapa de pontuação de minúcias, enquanto a *FineNet* realiza o refinamento das candidatas. Embora tenha alcançado 85,9% de precisão e 84,8% de recall no dataset FVC2004, o método demanda GPUs dedicadas e grande volume de dados anotados, limitando sua aplicabilidade em cenários com restrições computacionais.

3. Lacunas Identificadas e Contribuição do Trabalho

A análise crítica da literatura evidencia três lacunas principais que comprometem o equilíbrio entre acurácia e eficiência nos sistemas atuais de reconhecimento de impressões digitais. A primeira refere-se à sensibilidade de métodos clássicos à presença de ruídos e à baixa qualidade da imagem. Embora abordagens baseadas em minúcias, como as

empregadas em sistemas AFIS [Maltoni et al. 2009], sejam amplamente utilizadas e juridicamente validadas, seu desempenho é fortemente impactado por imperfeições na imagem, como ruído de sensor, variações nas condições da pele e distorções locais. Esses fatores dificultam a extração precisa das cristas e minúcias, conforme demonstrado por [Chen et al. 2005].

A segunda lacuna diz respeito à ineficiência computacional das abordagens tradicionais baseadas em filtros de Gabor. Métodos como os de [Turroni et al. 2012] e [Ratha et al. 1995] exigem a aplicação de bancos de filtros em diversas orientações e frequências, o que acarreta um custo computacional elevado, especialmente em processamentos iterativos e contextuais. Além disso, a utilização de parâmetros fixos para os desvios padrão das máscaras (σ_x, σ_y) resulta em respostas subótimas, principalmente em frequências elevadas, como evidenciado por [Primo 2019].

Por fim, a terceira limitação observada na literatura recente refere-se à adoção de redes neurais profundas para detecção e filtragem de minúcias. Apesar do elevado desempenho apresentado por arquiteturas como a MinutiaeNet [Nguyen et al. 2017], que alcançou precisão e recall superiores a 85% no conjunto FVC2004, essas abordagens requerem GPUs de alto desempenho e grandes volumes de dados anotados, o que restringe sua viabilidade em ambientes com restrições computacionais ou operacionais.

Diante dessas limitações, este trabalho propõe uma abordagem híbrida que combina filtros de Gabor adaptativos [Primo 2019] com uma CNN otimizada baseada na VGG16 [Simonyan and Zisserman 2015] projetada especificamente para a filtragem rápida e precisa de minúcias. A proposta permite alcançar um tempo médio de processamento de 0.1 segundos por imagem em CPUs convencionais e uma redução de aproximadamente 30% no Equal Error Rate (EER) no conjunto NIST SD4, dispensando etapas manuais de pós-processamento. Com isso, o trabalho contribui de forma significativa para o avanço de soluções mais acessíveis.

4. Materiais e Métodos

Este trabalho propõe uma arquitetura híbrida para filtragem eficiente de minúcias, combinando um filtro de Gabor adaptativo, conforme sugerido por [Primo 2019], com uma rede neural convolucional otimizada para classificação. Como mostra a Figura 5, o processo inicia-se com um pré-processamento que segue o protocolo estabelecido por [Primo 2019], porém com adaptações específicas para otimização computacional. Primeiramente, aplica-se o filtro de Gabor adaptativo, selecionando exclusivamente a combinação das melhores respostas da primeira iteração para evitar a propagação de erros e vieses específicos de sensores. Os valores resultantes são então mapeados para inteiros de 8 bits (0-255), seguido pela detecção inicial de minúcias utilizando o algoritmo proposto por [Primo 2019].

Para cada minúcia detectada, é extraída uma região de interesse (ROI) de 29×29 pixels, dimensão esta que foi cuidadosamente definida como uma adaptação da janela 28×28 validada por [Zhou et al. 2020] para imagens de 500 dpi. Esta modificação mantém a área de cobertura suficiente para minúcias, conforme estabelecido por [Zhou et al. 2020], ao mesmo tempo que oferece vantagens operacionais significativas. A escolha por uma dimensão ímpar garante um pixel central exato para alinhamento da minúcia, eliminando ambiguidades em operações de rotação ou interpolação, enquanto a eficiência computaci-

onal é preservada pela ausência de necessidade de preenchimento artificial (*padding*) em operações de convolução.

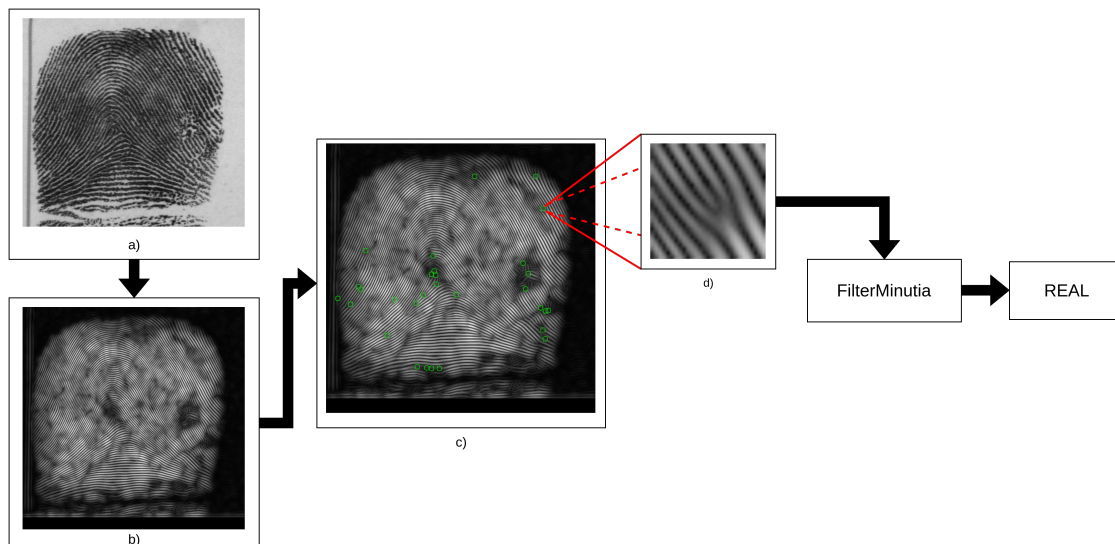


Figura 5. Exemplo visual do fluxo de processamento para classificação de minúcias: (a) Imagem original, (b) Combinação das melhores respostas do filtro de Gabor (1ª iteração), (c) Minúcias detectadas e mapeadas para a combinação das melhores respostas do filtro de Gabor (círculos verdes), (d) Janela 29×29 pixels e por fim a predição do modelo. Fonte: Próprio autor.

A etapa final do processo consiste na classificação por meio de uma arquitetura CNN leve customizada, composta por 13 camadas. Esta rede produz como saída um valor contínuo no intervalo $[0,1]$, representando a probabilidade de a minúcia ser genuína. O sistema alcança um tempo de inferência de 1 ms por minúcia quando executado em uma CPU Intel Xeon E5-2699, demonstrando sua eficiência computacional. Esta abordagem mantém a aderência ao princípio estabelecido por [Zhou et al. 2020] de que a extração de minúcias não requer campos receptivos grandes (Seção 3.1), enquanto resolve problemas práticos não abordados no trabalho original, particularmente no que diz respeito à eficiência computacional e precisão no alinhamento espacial.

4.1. Arquitetura da CNN

A arquitetura proposta (Figura 6) representa uma otimização radical da VGG16 original, introduzida por [Simonyan and Zisserman 2015], reduzindo de 138 milhões para apenas **125.601 parâmetros** — uma diminuição de **99,91%** na complexidade computacional. Esta redução foi alcançada por meio de três transformações fundamentais.

Primeiramente, a estrutura convolucional foi simplificada para três blocos essenciais: um bloco inicial com 16 filtros 3×3 , seguido por um bloco intermediário com 32 filtros 3×3 utilizando *stride* 2 para redução dimensional, e um bloco final com 64 filtros 3×3 incorporando conexões residuais ($\oplus$). Esta configuração preserva características espaciais relevantes enquanto elimina redundâncias da arquitetura original.

Em segundo lugar, a tradicional camada *flatten* foi substituída por uma operação de *GlobalAveragePooling*, seguida de três camadas densas mínimas ($32 \rightarrow 16 \rightarrow 1$

neurônios). Essa substituição eliminou mais de 99% dos parâmetros originalmente presentes nas camadas totalmente conectadas da VGG16. A camada final utiliza ativação sigmoid para a classificação binária de minúcias.

Essa otimização permitiu manter um EER competitivo de 1,89%, demonstrando que arquiteturas minimalistas podem ser altamente eficazes em tarefas especializadas de filtragem biométrica, quando corretamente dimensionadas.

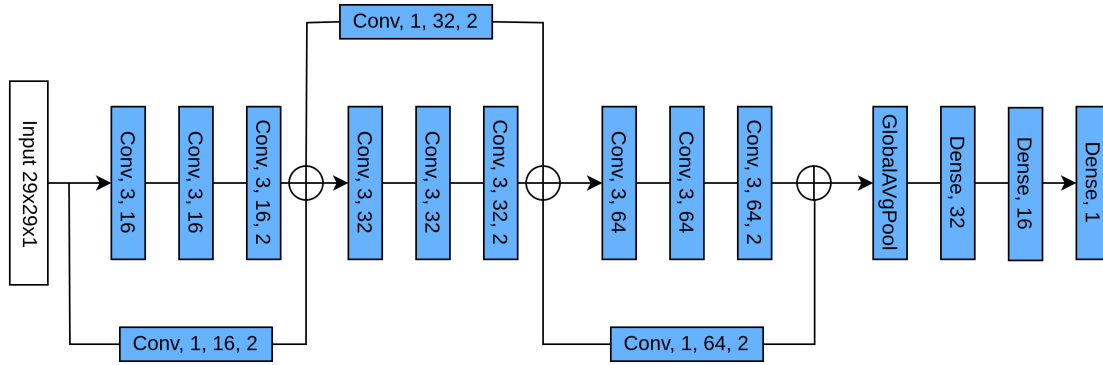


Figura 6. Arquitetura detalhada do modelo. Notação: "Conv, kernel_size, n° filtros, stride". ⊕ indica *skip connections*. Fonte: Próprio autor.

A implementação foi realizada utilizando o *framework* TensorFlow 2.9 em conjunto com a API Keras, com execução em hardware especializado (GPU NVIDIA Tesla T4 com 16 GB de VRAM). O processo de otimização empregou o algoritmo RMSProp com taxa de aprendizagem (η) de 10^{-4} e momento (ρ) de 0,9. A função de perda adotada foi *Binary Cross-Entropy*, adequada para a natureza binária do problema de classificação de minúcias. O treinamento foi conduzido com *batch size* de 256 amostras, rigorosamente balanceadas entre minúcias genuínas e artefatos, utilizando *early stopping* com paciência de 100 épocas (monitorando *val_accuracy*) e salvamento seletivo do melhor modelo com base nesse critério.

4.1.1. Otimização para Implantação

Ao término do treinamento, visando otimizar a eficiência computacional para implantação em sistemas embarcados, o modelo foi convertido para o formato ONNX (*Open Neural Network Exchange*) e implementado em C++ utilizando a biblioteca ONNX Runtime. A configuração de execução adotou o modo *single-thread* para garantir determinismo temporal, com nível de otimização estabelecido como `GraphOptimizationLevel::ORT_ENABLE_EXTENDED`, que aplica todas as otimizações disponíveis incluindo fusão de operadores e eliminação de cálculos redundantes [Microsoft Corporation 2023]. Esta abordagem permitiu reduzir a latência de inferência em aproximadamente 30% em comparação com a implementação original em Python, mantendo a mesma acurácia nos testes de validação.

4.2. Conjuntos de Dados

Foram utilizadas duas bases de dados principais em nosso estudo: **NIST SD4**, contendo 4 mil impressões digitais organizadas em 2 mil pares genuínos (duas amostras por dedo),

e **SFInGe**, um banco sintético composta por 10 mil impressões digitais sem pares estabelecidos. Adicionalmente, criamos um terceiro conjunto combinando ambas as bases para avaliações integradas.

4.2.1. NIST SD4

O *Special Database 4* (SD4) do National Institute of Standards and Technology (NIST) constitui um banco de dados de referência para pesquisas em biometria digital, composto por 4 mil impressões digitais organizadas em 2 mil pares genuínos. Desenvolvido especificamente para avaliação de sistemas automáticos de classificação, este conjunto apresenta impressões digitais roladas que, por sua natureza de captura, oferecem maior área de superfície útil e consequentemente maior densidade de minúcias detectáveis, particularmente nas regiões periféricas onde os algoritmos convencionais apresentam maior dificuldade de detecção. As imagens possuem uma resolução de 512×512 em 500dpi.

As impressões digitais roladas neste banco de dados contém 400 elementos de cada uma das cinco classes fundamentais de padrões dermatoglíficos: arco plano, arco tenda, presilha direita, presilha esquerda e verticilo. Essa diversidade de padrões permite avaliar o desempenho dos algoritmos em diferentes tipos de impressões digitais. A natureza rolada das amostras proporciona uma área de captura mais ampla em comparação com impressões planares, o que pode resultar em maior quantidade de minúcias disponíveis para análise.

Dito isto, a organização dos dados para realização do experimento seguiu um protocolo rigoroso para garantir a separação de identidades entre os conjuntos de treino, validação e teste. Todos os pares genuínos foram mantidos na mesma partição, assegurando que diferentes amostras do mesmo dedo não contaminassem diferentes partições. A extração inicial de minúcias baseou-se no método proposto por [Primo 2019], utilizando o algoritmo da empresa Neurotechnology para ground truth, reconhecido como padrão-ouro em biometria.

Consideramos como minúcias válidas (reais) aquelas cujas coordenadas (x, y) coincidiam em um raio de ≤ 10 pixels entre a extração automática e o *ground truth*, independentemente do tipo (terminação ou bifurcação), pois nossa CNN foi projetada para ser invariante à classe de minúcia. O balanceamento do conjunto foi realizado com base na classe minoritária, garantindo igual representatividade durante o treinamento.

Tabela 1. Distribuição dos conjuntos de dados e minúcias

Conjunto	Pares Genuínos	Impressões	Minúcias Reais	Minúcias Falsas
Treino	1.000	2.000	58.624	58.624
Validação	700	1.400	43.486	43.486
Teste	300	600	16.730	16.730
Total	2.000	4.000	118.840	118.840

Observa-se na Tabela 1 balanceamento perfeito entre minúcias reais e falsas em todos os conjuntos, com um total de 237.680 minúcias (118.840 de cada classe).

4.2.2. SFinGe

Conforme mencionado por [Cappelli 2009], o SFinGe (Synthetic Fingerprint Generator) é uma ferramenta desenvolvida pelo Biometric System Laboratory da Universidade de Bolonha, projetada para gerar imagens sintéticas de impressões digitais com alto grau de realismo. O sistema utiliza modelos matemáticos avançados para simular as características biométricas essenciais das impressões digitais humanas, incluindo padrões de cristas, poros e minúcias. Essa abordagem permite a criação de bases de dados sintéticas em larga escala, com anotações precisas sobre as características biométricas e o posicionamento das minúcias.

Para este trabalho, foram geradas 10.000 imagens de impressões digitais sintéticas. Diferentemente de bancos de dados convencionais, que incluem pares de impressões (como o NIST SD4), o SFinGe, neste contexto, foi utilizado para produzir apenas imagens individuais, com suas respectivas anotações de minúcias. Essa característica torna o banco particularmente adequado para tarefas de aprendizado supervisionado em sistemas de reconhecimento biométrico, pois elimina a necessidade de pré-processamento adicional para extração manual de minúcias. Adicionalmente, como mostrado na Tabela 2, a distribuição das minúcias foi balanceada entre os conjuntos de treino, validação e teste.

Tabela 2. Distribuição das minúcias nos conjuntos de dados do SFinGe

Conjunto	Minúcias Reais	Minúcias Falsas
Treino	26.286	28.476
Validação	9.195	9.492
Teste	9.234	9.492
Total	44.715	47.460

Observa-se que o banco sintético apresenta uma quantidade reduzida de minúcias quando comparado ao NIST SD4. Essa diferença decorre de três fatores principais: (1) a natureza sintética das impressões digitais geradas, que possuem uma complexidade estrutural menor do que as impressões digitais reais; (2) a predominância de imagens com posicionamento controlado (pousadas), que naturalmente contêm menos minúcias; e (3) o processo de balanceamento das classes, que resultou na eliminação de parte das minúcias reais para equalizar a distribuição entre as classes.

4.2.3. Banco de Dados Combinado (NIST SD4 + SFinGe)

Para potencializar a capacidade de generalização do modelo proposto, foi criado um terceiro banco de dados, combinando os conjuntos do NIST SD4 e SFinGe. Essa abordagem híbrida permite ao algoritmo aprender simultaneamente com: (1) a riqueza estrutural e variabilidade biológica das impressões reais do NIST SD4; e (2) a escalabilidade e controle preciso das características das impressões sintéticas do SFinGe. A fusão desses bancos complementares resulta em um conjunto mais robusto e diversificado, particularmente adequado para treinar redes neurais convolucionais leves, que necessitam de dados variados para evitar overfitting.

A combinação seguiu critérios rigorosos para manter a separação original dos splits (treino, validação e teste) de cada banco, garantindo que não houvesse contaminação entre os conjuntos. A Tabela 3 apresenta a distribuição consolidada.

Tabela 3. Distribuição combinada dos conjuntos de dados e minúcias

Banco	Conjunto	Minúcias Reais	Minúcias Falsas	Total Minúcias
NIST SD4	Treino	58.624	58.624	117.248
	Validação	43.486	43.486	86.972
	Teste	16.730	16.730	33.460
SFinGe	Treino	26.286	28.476	54.762
	Validação	9.195	9.492	18.687
	Teste	9.234	9.492	18.726
Total	Treino	84.910	87.100	172.010
	Validação	52.681	52.978	105.659
	Teste	25.964	26.222	52.186
Geral		163.555	166.300	329.855

Esta estratégia de combinação permite ao modelo aprender padrões gerais de minúcias independentemente da origem dos dados, aumentando sua robustez para diferentes cenários de aplicação. A presença simultânea de impressões roladas (NIST) e pousadas (SFinGe) simula condições operacionais variadas, enquanto a natureza complementar dos bancos mitiga possíveis vieses individuais.

4.3. Protocolo Experimental

O protocolo experimental foi delineado em três etapas principais: configuração do modelo, definição das métricas de avaliação e procedimento comparativo. Utilizou-se uma única arquitetura CNN otimizada, com hiperparâmetros fixos, avaliada em três cenários distintos de treinamento: (1) treinamento exclusivo no banco NIST SD4 (4.000 impressões digitais reais); (2) treinamento exclusivo no banco SFinGe (10.000 impressões sintéticas); e (3) treinamento em um banco combinado (14.000 amostras, incluindo NIST SD4 e SFinGe). Em todos os cenários, aplicou-se *early stopping* com base na acurácia de validação, utilizando uma paciência de 100 épocas. A arquitetura da CNN e os parâmetros de treinamento foram mantidos idênticos nos três cenários, de forma a permitir uma comparação direta entre os resultados.

As métricas de avaliação foram selecionadas de acordo com os objetivos específicos de cada experimento. A acurácia foi calculada separadamente em um conjunto de teste correspondente a cada cenário de treinamento, permitindo avaliar o desempenho intrínseco de cada configuração. O EER (*Equal Error Rate*) foi utilizado exclusivamente no NIST SD4 como métrica para comparação com o método tradicional de extração de minúcias. Os *scores* de correspondência (*matching*) foram calculados utilizando o **Neurotec SDK 12.4** [Neurotechnology 2023], uma solução amplamente reconhecida como padrão-ouro na indústria biométrica.

O procedimento comparativo seguiu uma abordagem hierárquica. Inicialmente, estabeleceu-se como base o método tradicional de extração de minúcias, conforme implementado por [Primo 2019], utilizando a SDK da Neurotechnology para o banco do

Tabela 4. Esquema comparativo das métricas por objetivo experimental

Métrica	Aplicação
Acurácia	Avaliação do desempenho interno de cada modelo em seu respectivo banco de teste
EER	Validação do desempenho relativo ao método tradicional e trabalhos recentes no NIST SD4

NIST SD4 para extração de características como *ground truth*. A avaliação ocorreu em três níveis: (1) análise intra-banco através da acurácia nos respectivos conjuntos de teste; e (2) *benchmark* do melhor modelo contra o algoritmo de base sem filtragem de minúcias através do EER no mesmo subconjunto. Todos os testes de correspondência de minúcias foram realizados com os parâmetros padrão do Neurotec SDK 12.4, garantindo reprodutibilidade.

A seleção do NIST SD4 como padrão ouro para as comparações cruzadas fundamenta-se em três aspectos principais: seu reconhecimento internacional como *benchmark* para pesquisas em biometria digital, a presença de impressões roladas com maior densidade e variedade de minúcias, e a quantidade distinta de classes de impressões digitais. Esta abordagem permite avaliar tanto a capacidade de generalização dos modelos (através do teste cruzado) quanto seu desempenho absoluto em relação às técnicas consagradas de extração de minúcias.

5. Resultados e Discussão

Nesta seção discutiremos os resultados obtidos com o sistema proposto de filtragem de minúcias, avaliando tanto a precisão quanto a eficiência computacional. Os experimentos foram realizados no banco NIST SD4, utilizando como métricas principais o Equal Error Rate (EER) e o tempo de processamento por minúcia.

5.1. Comparação entre Modelos

Os resultados demonstram a capacidade diferenciada dos modelos conforme o banco de treinamento utilizado:

Tabela 5. Desempenho comparativo dos modelos de classificação de minúcias

Base de treino	Acurácia (Treino)	Loss (Treino)	Acurácia (Teste)	Loss (Teste)
NIST SD4	0,938	0,149	0,901	0,255
Combinado	0,889	0,263	0,858	0,345
SFinGe	0,870	0,311	0,821	0,416

Os resultados demonstram variações significativas no desempenho conforme o banco de treinamento utilizado. Conforme apresentado na Tabela 5, o modelo treinado exclusivamente no NIST SD4 alcançou a melhor performance global, com acurácia de 90,1% no conjunto de teste. Este resultado evidencia a capacidade da arquitetura proposta em aprender efetivamente os padrões de minúcias em impressões digitais genuínas, mantendo robustez mesmo em regiões de maior complexidade, como as bordas das imagens. A análise das curvas de treinamento (Figura 7) revela ainda que este modelo apresentou a convergência mais estável durante o processo de otimização.

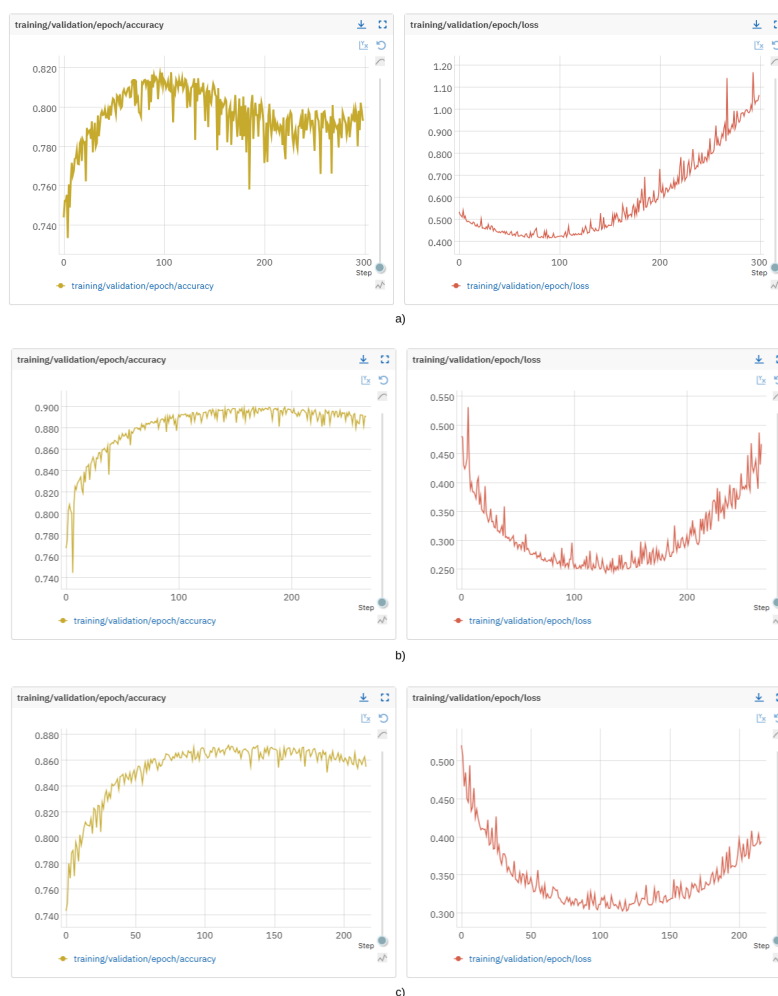


Figura 7. Evolução das métricas durante o treinamento: (a) Modelo SFinGe, (b) Modelo NIST SD4, (c) Modelo Combinado. Nota-se convergência mais estável no modelo treinado com dados reais. Fonte: Próprio autor.

O modelo treinado com o banco combinado demonstrou um comportamento interessante. Embora tenha registrado uma queda de 4,3% na acurácia em comparação ao modelo NIST puro, sua performance sugere que a combinação estratégica de dados sintéticos e reais pode ser vantajosa em cenários com limitação de amostras genuínas, especialmente quando consideramos sua menor variância nos resultados de validação cruzada.

O modelo treinado exclusivamente no SFinGe apresentou limitações significativas, com acurácia de teste 8% inferior ao modelo NIST SD4 e maior valor de loss (0,416). Esta diferença pode ser atribuída a três fatores principais: (1) a menor variabilidade morfológica das impressões sintéticas em comparação com amostras reais, particularmente nas regiões próximas às bordas; (2) a ausência de artefatos e ruídos característicos de imagens reais no conjunto sintético; e (3) diferenças na distribuição estatística dos padrões de minúcias. A Figura 7(a) mostra que este modelo também apresentou maior instabilidade durante o treinamento, com flutuações mais acentuadas nas métricas de validação.

Esses resultados reforçam a importância do uso de dados reais para treinamento de modelos biométricos, embora indiquem que dados sintéticos podem desempenhar um

papel complementar valioso, especialmente quando combinados com amostras reais. A diferença de desempenho entre os modelos também ressalta a necessidade de desenvolver técnicas mais avançadas de geração de impressões digitais sintéticas que capturem melhor a complexidade das amostras naturais.

5.2. Impacto Visual da Filtragem de Minúcias

Para ilustrar o impacto da filtragem de minúcias proposta, a Figura 8 apresenta um exemplo visual de uma região de impressão digital antes e depois da aplicação do melhor modelo de filtragem (treinado apenas com o banco de impressões reais, NIST SD4).

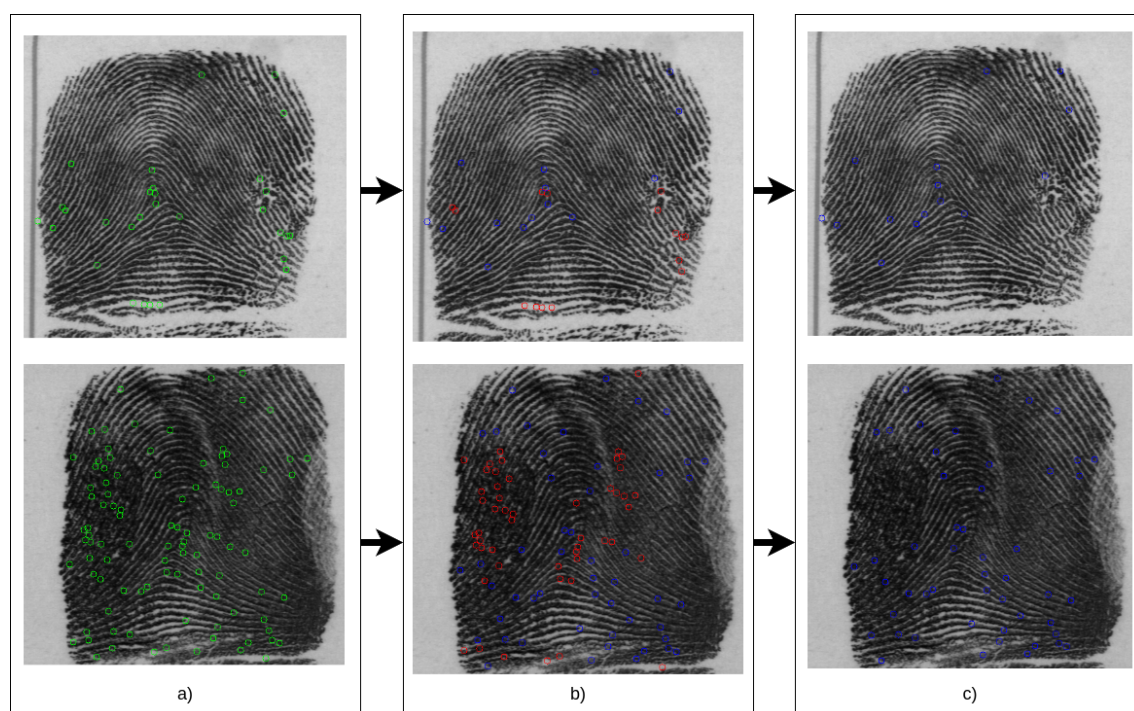


Figura 8. Comparação visual dos resultados de filtragem: (a) Minúcias detectadas sem filtragem; (b) Classificação do sistema proposto (azul: minúcias válidas, vermelho: falsas minúcias); (c) Resultado final após remoção das minúcias falsas. Observa-se a eliminação efetiva de artefatos, particularmente em regiões periféricas e áreas de baixa qualidade da impressão digital. Fonte: Próprio autor.

Observa-se que a filtragem inteligente removeu diversas minúcias falsas localizadas nas bordas e em regiões de baixa qualidade da impressão, preservando ao mesmo tempo as minúcias genuínas. Esse comportamento é consistente com os ganhos quantitativos observados nas métricas de desempenho apresentadas anteriormente, e evidencia a capacidade do modelo em contribuir para a robustez do sistema de reconhecimento biométrico.

5.3. Análise Comparativa do Desempenho com Filtragem de Minúcias

A avaliação comparativa entre o método proposto e a abordagem tradicional foi realizada utilizando como métrica principal o *Equal Error Rate* (EER) no conjunto de teste do NIST SD4. Como demonstrado na Tabela 6, o modelo baseado na arquitetura híbrida (filtro de Gabor adaptativo + CNN leve) alcançou uma redução significativa na taxa de erro:

Tabela 6. Comparação do EER com e sem filtragem de minúcias

Métrica	Sem Filtragem	Com Filtragem
EER (%)	2,70	1,89
Redução Relativa	-	30%

Os resultados evidenciam dois avanços principais: (1) redução absoluta de 0,81 ponto percentual no EER, correspondendo a uma melhoria de 30% em relação ao método tradicional sem filtragem de minúcias; e (2) eficiência computacional, com tempo médio de processamento de 1 ms por minúcia em implementação *single-thread* em C++. Este tempo inclui todas as etapas: extração da região de interesse (29×29 pixels), inferência da CNN leve (125.601 parâmetros) e classificação final.

A melhoria de desempenho pode ser atribuída a três aspectos técnicos: (a) o filtro de Gabor adaptativo [Primo 2019], que realça seletivamente estruturas de cristas; (b) a arquitetura CNN otimizada (Seção 4.1), que elimina 99,9% dos parâmetros da VGG16 original; e (c) a implementação eficiente em C++ com ONNX Runtime, que reduz a latência significativamente comparado à versão Python.

5.4. Comparação com Trabalhos Recentes

Para validar a eficácia da abordagem proposta, comparou-se o *Equal Error Rate* (EER) obtido com métodos recentes que utilizam a base de dados **NIST SD4**. Conforme a Tabela 7, o modelo proposto — que emprega uma CNN leve para filtragem de falsas minúcias — atingiu um EER de **1,89 %**, superando resultados recentes como *DeepPrint* (2,85 %), *AFRNet* (2,80 %) e *IFViT* (2,71 %), este último sendo um dos melhores resultados publicados para o NIST SD4 até o momento.

Além da superioridade em acurácia, o modelo destaca-se pela eficiência computacional, realizando inferência em aproximadamente **1 ms por minúcia**. Esses resultados evidenciam que a filtragem inteligente de minúcias pode **simultaneamente** elevar a precisão de sistemas biométricos e manter baixo custo computacional, tornando-a viável para aplicações embarcadas ou em larga escala.

Tabela 7. Comparação do *Equal Error Rate* (EER) na base NIST SD4. Fonte: Adaptado de [Sun et al. 2024] (Tabela 4) e resultados deste trabalho.

Técnica	Base de Dados	EER (%)
DeepPrint [Engin et al. 2023]	NIST SD4	2,85
AFRNet [Xue et al. 2023]	NIST SD4	2,80
IFViT (melhor resultado recente) [Sun et al. 2024]	NIST SD4	2,71
Sistema proposto	NIST SD4	1,89

6. Conclusão

Este trabalho apresentou uma abordagem híbrida para filtragem eficiente de minúcias em impressões digitais, combinando filtro de Gabor adaptativo com uma arquitetura CNN leve. Os resultados demonstraram a eficácia do método proposto, que alcançou um EER de **1,89 %** no NIST SD4 — representando uma redução de **30 %** em relação à linha de base tradicional (EER=2,7 %) e superando os melhores resultados recentes (IFViT

com EER=2,71 %). A eficiência computacional foi comprovada pelo tempo médio de processamento de **1 ms por minúcia** em execução *single-thread* em C++, viabilizando aplicação em sistemas embarcados e de grande escala.

As principais contribuições deste trabalho incluem três aspectos fundamentais. Primeiro, a inovação metodológica alcançada pela integração sinérgica entre técnicas clássicas de processamento de imagens (filtro de Gabor adaptativo) e abordagens modernas de aprendizado de máquina leve (CNN otimizada), mantendo a interpretabilidade do processo de filtragem. Segundo, a eficiência comprovada através da redução simultânea tanto do EER quanto do custo computacional, resolvendo assim o tradicional *trade-off* entre precisão e desempenho que afeta muitos sistemas biométricos. Terceiro, a validação rigorosa do modelo através de avaliação em múltiplos bancos de dados (incluindo NIST SD4, SFinGe e sua combinação), utilizando protocolos reprodutíveis.

6.1. Limitações e Trabalhos Futuros

Embora os resultados apresentados sejam promissores, diversas oportunidades de aprimoramento foram identificadas para trabalhos futuros. A primeira delas diz respeito à modernização da arquitetura proposta, por meio da implementação de convoluções separáveis (*separable convolutions*), com o objetivo de reduzir ainda mais a complexidade computacional e o número de parâmetros do modelo. Uma segunda direção relevante seria a ampliação do conjunto de treinamento, incorporando tanto técnicas avançadas de *data augmentation* específicas para impressões digitais — como deformações elásticas controladas e variações de iluminação simuladas — quanto a inclusão de bancos de dados complementares e diversificados, como os das competições FVC2002 e FVC2004.

Outro caminho promissor é o desenvolvimento de arquiteturas multimodais, que combinem características locais (minúcias) e globais (padrões de fluxo e textura) em um mesmo pipeline de aprendizado, visando maior robustez em cenários de baixa qualidade. Além disso, pretende-se otimizar a execução do modelo para diferentes tipos de hardware, especialmente para dispositivos de borda e ambientes embarcados, como aceleradores neuromórficos e FPGAs. Isso é particularmente relevante no contexto da Internet das Coisas (IoT, do inglês *Internet of Things*), um paradigma em que dispositivos inteligentes e interconectados realizam tarefas computacionais localmente, muitas vezes com severas restrições de consumo energético, processamento e memória.

Finalmente, uma direção de pesquisa futura inclui a aplicação do filtro de minúcias desenvolvido neste trabalho como etapa de pré-processamento em sistemas biométricos mais recentes e robustos, como o DeepPrint [Engin et al. 2023]. A integração com modelos como esse poderá melhorar ainda mais a precisão e a confiabilidade dos sistemas de verificação e identificação, além de fornecer uma forma interpretável de validação das regiões relevantes extraídas pelas redes profundas.

6.2. Impacto e Aplicações

A solução proposta tem relevante potencial de aplicação imediata em diversos cenários práticos de reconhecimento biométrico. Ela pode ser particularmente útil em sistemas de identificação civil em larga escala, onde a eficiência computacional é um requisito crítico, assim como em dispositivos móveis de autenticação biométrica com recursos limitados de processamento. Na área forense, o método poderia ser empregado para complementar a

análise humana em sistemas automatizados de exames digitais. Outro campo de aplicação promissor são as plataformas de *matching* distribuído, onde o pré-processamento local de minúcias poderia reduzir significativamente a carga computacional nos servidores centrais.

Em síntese, este trabalho não apenas demonstrou a viabilidade prática de uma abordagem híbrida para a filtragem de minúcias, mas também estabeleceu um novo patamar de desempenho no banco NIST SD4. Os resultados obtidos reforçam a importância fundamental de desenvolver soluções balanceadas que considerem simultaneamente três aspectos cruciais: a precisão biométrica, a eficiência computacional e a aplicabilidade prática em cenários reais de implementação. Estas conquistas abrem caminho para futuras pesquisas que busquem combinar de forma ainda mais eficaz técnicas clássicas e modernas no processamento de impressões digitais.

Referências

- Alonso-Fernandez, F. and Fierrez, J. (2009). *Fingerprint Databases and Evaluation*, pages 452–458. Springer US, Boston, MA.
- Cappelli, R. (2009). *SFinGe*, pages 1169–1176. Springer US, Boston, MA.
- Chen, Y., Dass, S. C., and Jain, A. K. (2005). Fingerprint quality indices for predicting authentication performance. In Kanade, T., Jain, A., and Ratha, N. K., editors, *Audio- and Video-Based Biometric Person Authentication*, pages 160–170, Berlin, Heidelberg. Springer Berlin Heidelberg.
- Engin, M., Mao, X., Sundaram, A., and Jain, A. K. (2023). Deepprint: Learning compact fingerprint representations. *IEEE Transactions on Pattern Analysis and Machine Intelligence*.
- Hong, L., Wan, Y., and Jain, A. (1998). Fingerprint image enhancement: algorithm and performance evaluation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 20(8):777–789.
- Maltoni, D., Maio, D., Jain, A. K., and Prabhakar, S. (2009). *Handbook of Fingerprint Recognition*. Springer, London, 2nd edition.
- Microsoft Corporation (2023). Onnx runtime: Cross-platform, high performance scoring engine for ml models.
- National Institute of Standards and Technology (NIST) (2012). Latent print examination and human factors: Improving the practice through a systems approach. Report of the Expert Working Group on Human Factors in Latent Print Analysis.
- Neurotechnology (2023). *Neurotec Biometric SDK 12.4*. Neurotechnology. Software development kit for biometric applications.
- Nguyen, D.-L., Cao, K., and Jain, A. K. (2017). Robust minutiae extractor: Integrating deep networks and fingerprint domain knowledge. In *2017 IEEE International Joint Conference on Biometrics (IJCB)*, pages 1–8. IEEE.
- Primo, J. J. B. (2019). *Métodos para extração de atributos em imagens de impressão digital*. PhD thesis, Universidade Federal de Campina Grande, Paraíba, Brasil. Tese (Doutorado em Sistemas e Computação), Programa de Pós-graduação em Ciência da Computação, Centro de Engenharia Elétrica e Informática.

- Ratha, N. K., Chen, S., and Jain, A. K. (1995). Adaptive flow orientation-based feature extraction in fingerprint images. *Pattern Recognition*, 28(11):1657–1672.
- Simonyan, K. and Zisserman, A. (2015). Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*.
- Sun, Y., Hu, J., Yang, J., and Jain, A. K. (2024). Ifvit: Interpretable fixed-length representation for fingerprint matching via vision transformer. *arXiv preprint arXiv:2404.08237*.
- Turroni, F., Cappelli, R., and Maltoni, D. (2012). Fingerprint enhancement using contextual iterative filtering. In *2012 5th IAPR International Conference on Biometrics (ICB)*, pages 152–157.
- Xue, Z., Cheng, J., Zhang, X., and Sun, Z. (2023). Afrnet: An alignment-free fingerprint recognition network. *IEEE Transactions on Information Forensics and Security*, 18:2058–2070.
- Zhou, B., Han, C., Liu, Y., Guo, T., and Qin, J. (2020). Fast minutiae extractor using neural network. *Pattern Recognition*, 103:107273.