



INSTITUTO FEDERAL DE EDUCAÇÃO, CIÊNCIA E TECNOLOGIA DO PIAUÍ
CAMPUS FLORIANO
CURSO TECNOLOGIA EM ANÁLISE E DESENVOLVIMENTO DE
SISTEMAS

JÉSSICA MAYARA SOARES FERREIRA

APLICAÇÕES DE TÉCNICAS DE RECONHECIMENTO ÓPTICO DE
CARACTERES PARA A EXTRAÇÃO DE TEXTO EM IMAGENS: UM ESTUDO DE
CASO NO COMÉRCIO VAREJISTA DE ARTIGOS DE VESTUÁRIO

FLORIANO

2025

JÉSSICA MAYARA SOARES FERREIRA

**APLICAÇÕES DE TÉCNICAS DE RECONHECIMENTO ÓPTICO DE
CARACTERES PARA A EXTRAÇÃO DE TEXTO EM IMAGENS: UM ESTUDO DE
CASO NO COMÉRCIO VAREJISTA DE ARTIGOS DE VESTUÁRIO**

Trabalho de Conclusão de Curso (artigo científico)
apresentado como exigência parcial para obtenção
do diploma do Curso de Tecnologia em Análise e
Desenvolvimento de Sistemas do Instituto Federal
de Educação, Ciência e Tecnologia do Piauí,
Campus Floriano.

Orientador: Prof. Dr. Willamys Rangel Nunes de
Sousa.

FLORIANO

2025

APLICAÇÕES DE TÉCNICAS DE RECONHECIMENTO ÓPTICO DE CARACTERES PARA A EXTRAÇÃO DE TEXTO EM IMAGENS: UM ESTUDO DE CASO NO COMÉRCIO VAREJISTA DE ARTIGOS DE VESTUÁRIO

Jéssica Mayara Soares Ferreira¹
Willamys Rangel Nunes de Sousa²

RESUMO

O objetivo deste estudo foi desenvolver uma ferramenta computacional para o setor varejista, visando otimizar a gestão de dados de fornecedores por meio da automação da extração de informações contidas em cartões de visita. A ferramenta utiliza técnicas de Reconhecimento Óptico de Caracteres (OCR), Processamento de Linguagem Natural (PLN) com Reconhecimento de Entidades Nomeadas (NER) e expressões regulares para identificar dados como números de telefone e perfis de Instagram. O processo foi dividido em etapas de coleta de dados, teste do modelo, validação e implementação de uma interface gráfica para organizar os dados extraídos. A base de dados foi composta por 21 imagens de cartões de visita. Após a validação, a ferramenta demonstrou bom desempenho na extração e organização das informações, com resultados satisfatórios para dados padronizados, embora a expansão da base de dados seja necessária para aprimorar a robustez do modelo. Este trabalho contribui para a digitalização e automação dos processos de gestão de fornecedores no setor varejista, abrindo possibilidades para futuras melhorias e pesquisas na área.

Palavras-chave: reconhecimento óptico de caracteres; processamento de linguagem natural; extração de dados; gestão de fornecedores.

ABSTRACT

The objective of this study was to develop a computational tool for the retail sector, aiming to optimize the management of supplier data through the automation of information extraction from business cards. The tool uses Optical Character Recognition (OCR) techniques, Natural Language Processing (NLP) with Named Entity Recognition (NER), and regular expressions to identify data such as phone numbers and Instagram profiles. The process was divided into data collection, model testing, validation, and the implementation of a graphical interface to organize the extracted data. The database consisted of 21 business card images. After validation, the tool demonstrated good performance in extracting and organizing the information, with satisfactory results for standardized data, although expanding the database is necessary to improve the model's robustness. This work contributes to the digitalization and automation of supplier management processes in the retail sector, opening opportunities for future improvements and research in the field.

Keywords: optical character recognition; natural language processing; data extraction; supplier management.

Data de aprovação: 16/07/2025 (data de apresentação do TCC)

¹ Graduando no curso de Análise e Desenvolvimento de Sistemas. IFPI- Campus Floriano. caflo.2021114tads07@aluno.ifpi.edu.br.

² Professor do Ensino Básico, Técnico e Tecnológico. IFPI- Campus Floriano. rangelnunes@ifpi.edu.br.

1 INTRODUÇÃO

No setor varejista, a gestão eficaz dos fornecedores desempenha um papel crucial para manter a competitividade e alcançar a eficiência operacional. Nesse cenário, a tecnologia assume um papel fundamental ao automatizar tarefas repetitivas e otimizar processos (Torres, 2024).

A inovação tecnológica nas empresas diz respeito à modernização dos processos produtivos, envolvendo principalmente a pesquisa e o desenvolvimento de novas soluções e tecnologias (Jiang; Qin; Khan, 2022).

Nesse contexto, a nova cadeia de suprimentos digital no setor varejista concentra-se na coleta, processamento e aplicação estratégica de dados, utilizando tecnologias como a inteligência artificial para aprimorar a tomada de decisões e otimizar a gestão operacional. O fortalecimento da inovação orientada por dados oferece um direcionamento claro para que as empresas avancem na digitalização de seus processos, especialmente na gestão da cadeia de suprimentos (Li, 2024).

Li (2024) destaca ainda que, à medida que mais empresas adotam soluções digitais, como a digitalização, o processo de transformação tende a se acelerar, e os benefícios da modernização continuarão a emergir de forma constante. Nesse cenário, a integração de ferramentas digitais avançadas e a automação de processos podem contribuir para uma gestão mais eficaz e clara (Jiang; Qin; Khan, 2022).

Contudo, mesmo diante dessas possibilidades, muitos negócios ainda dependem de métodos tradicionais, como o uso de planilhas, que podem comprometer a precisão e limitar a eficiência dessa gestão (Sartori, 2023).

Atualmente, o método de administrar cartões de visita envolve a transcrição manual das informações de contato para um formato digital, o que pode ser ineficiente e sujeito a erros. Nesse contexto, a digitalização de cartões de visitas dos fornecedores oferece vantagens significativas em comparação com o método tradicional de transcrição manual (Bhatt *et al.*, 2023).

Além da digitalização de cartões de visita, a aplicação de tecnologia de Reconhecimento Óptico de Caracteres (OCR) surge como uma alternativa promissora para otimizar a gestão de informações de fornecedores. O OCR permite a conversão automática de documentos físicos e digitalizados em dados editáveis, facilitando a extração eficiente e precisa de informações cruciais (Wang, 2023).

Portanto, este trabalho tem como objetivo propor uma ferramenta computacional para o setor varejista, aplicando técnicas de OCR, a fim de automatizar a extração e organização de informações de fornecedores a partir de imagens, como cartões de visita, contribuindo para uma gestão de dados mais eficiente e precisa.

2 REFERENCIAL TEÓRICO

Esta seção apresenta os fundamentos teóricos que sustentam o desenvolvimento da ferramenta proposta. A subseção 2.1 discute os princípios da Inteligência Artificial, enquanto a subseção 2.2 aborda o aprendizado de máquina (machine learning), incluindo suas principais abordagens: supervisionado e não supervisionado. Em seguida, a subseção 2.3 explora o Processamento de Linguagem Natural (PLN), com ênfase em técnicas utilizadas para a extração de informações em texto. A subseção 2.4 trata da Visão Computacional, seguida da subseção 2.5, que descreve o Processamento Digital de Imagens. A subseção 2.6 aborda o Reconhecimento Óptico de Caracteres (OCR), com destaque para a biblioteca EasyOCR, apresentada em 2.6.1. A subseção 2.7 apresenta as Expressões Regulares como

técnica complementar à extração textual. Por fim, a subseção 2.8 reúne trabalhos relacionados ao tema, com o objetivo de contextualizar a pesquisa no cenário atual.

2.1 INTELIGÊNCIA ARTIFICIAL

A Inteligência Artificial (IA) é uma área da computação que se concentra em desenvolver mecanismos capazes de realizar tarefas que, normalmente, exigiriam inteligência humana. Isso inclui a capacidade de aprender, raciocinar, resolver problemas e tomar decisões (Barbosa; Portes, 2023).

Algumas tecnologias essenciais na IA, como *Machine Learning* (aprendizado de máquina) e *Deep Learning* (aprendizado profundo), utilizam algoritmos para identificar padrões e características que frequentemente passam despercebidos (Jordan, Cesar, Carmo, 2022).

As tecnologias mencionadas têm encontrado aplicações práticas em diversos campos. Um exemplo significativo é a utilização de IA na área da saúde, particularmente na análise de prontuários médicos e diagnósticos. Algoritmos de *machine learning* e Processamento de Linguagem Natural (PLN) são usados para analisar grandes volumes de dados de pacientes, identificando padrões que podem indicar condições médicas como câncer. Esse processo permite diagnósticos mais rápidos e precisos, além de auxiliar na personalização de tratamentos (Brandão, 2024; Esteva *et al.*, 2017).

2.2 MACHINE LEARNING

O aprendizado de máquina, também conhecido como *machine learning*, é um componente fundamental da inteligência artificial. Por meio da criação de algoritmos, ele capacita os sistemas a aprender de maneira autônoma a partir de dados e experiências anteriores. Assim, quando novos dados são recebidos, o sistema faz previsões baseadas em dados históricos para criar modelos preditivos (Shaveta, 2023). Dentro desse campo, duas abordagens principais são reconhecidas: o aprendizado supervisionado e o aprendizado não supervisionado, cada um desempenhando funções específicas e úteis em diferentes contextos.

2.2.1 Aprendizado Supervisionado

O aprendizado supervisionado é um tipo de algoritmo preditivo no qual a máquina aprende identificando padrões em conjuntos de dados que estão rotulados. Este método é amplamente empregado em diversas aplicações, como classificação e regressão (Strauss; Villas Bôas Júnior; Ferreira, 2022).

Os algoritmos de classificação constituem uma ferramenta essencial no aprendizado supervisionado, com o propósito de prever a classe ou rótulo associado a uma variável de entrada, considerando cuidadosamente seus atributos específicos (Fontana, 2020).

Nesse contexto, uma variedade de técnicas são empregadas, tais como regressão logística, árvores de decisão e floresta randômica. Esses algoritmos oferecem diferentes abordagens para a classificação de dados, permitindo a escolha da mais adequada, conforme a natureza e complexidade do problema em questão (Fontana, 2020).

Nas tarefas de Regressão, por sua vez, os modelos são usados para realizar previsões. (Strauss; Villas Bôas Júnior; Ferreira, 2022). Métodos como regressão linear, árvores de decisão e redes neurais são comumente empregados para essas finalidades. Um exemplo disso é a aplicação de um modelo de regressão para estimar as vendas semanais de uma loja (Ogunleye, 2022).

2.2.2 Aprendizado Não Supervisionado

No aprendizado não supervisionado, a máquina é treinada com dados que não possuem rótulos predefinidos. Os algoritmos utilizados nesse tipo de aprendizado buscam descobrir estruturas e padrões nos dados, agrupando-os ou encontrando representações latentes sem informações prévias sobre categorias ou classes. Esse tipo de aprendizado é aplicado em tarefas como agrupamento e redução de dimensionalidade (Santos, 2023).

Os algoritmos de agrupamento (*clustering*) dividem os dados em grupos com características semelhantes, utilizando critérios predefinidos, o que permite identificar padrões nos dados fornecidos. Existem várias abordagens de *clustering*, que podem se basear na distância entre os pontos, em distribuições estatísticas específicas ou na densidade de pontos em áreas particulares do conjunto de dados (Fontana, 2020).

Por sua vez, a técnica de redução de dimensionalidade tem como objetivo reduzir a quantidade de características presentes em um conjunto de dados. Para essa finalidade, podem ser utilizados alguns métodos, como a Análise de Componentes Principais (PCA), t-distributed stochastic neighbor embedding (t-SNE), entre outros (Mandelli, 2023).

2.3 PROCESSAMENTO DE LINGUAGEM NATURAL

O Processamento de Linguagem Natural (PLN) é uma subárea da IA com foco na interação entre a linguagem humana e os computadores. Esse campo envolve a criação de algoritmos e técnicas que possibilitam às máquinas compreender, interpretar e gerar textos de forma semelhante aos seres humanos (Brandão, 2024).

Silva (2023) destaca diversas áreas de estudo que têm ganhado relevância crescente nas aplicações do Processamento de Linguagem Natural. Estas incluem, por exemplo, o OCR em textos manuscritos ou tipografados, o Processamento de Fala e a Classificação e Agrupamento de Documentos, entre outras.

Dentro do PLN, a tarefa de Reconhecimento de Entidades Nomeadas (REN) é crucial. Essa técnica envolve a identificação e classificação de entidades específicas em um texto, como nomes de pessoas, organizações, locais, datas, horários e outras informações relevantes. Além disso, a extração de dados através dessa técnica requer a criação de um software capaz de detectar essas entidades de maneira eficiente (Andrade *et al.*, 2023; Kirsch, Dorneles, 2023).

2.4 VISÃO COMPUTACIONAL

A visão computacional, como parte da inteligência artificial, concentra-se em permitir que os computadores compreendam e automatizem tarefas relacionadas à interpretação de imagens. Isso inclui uma variedade de desafios, desde reconhecimento de objetos até análise de cenas complexas (Pires, 2023).

Essa habilidade notável dos computadores em identificar elementos em imagens envolve uma série de cálculos matemáticos e o treinamento de redes neurais, que aprendem a reconhecer padrões, imitando o processo de aprendizagem humana (Pires, 2023).

Por ser uma área ampla, a visão computacional expande significativamente o uso de tecnologias. Dessa forma, há uma variedade de aplicações práticas. Alguns exemplos do uso desta tecnologia incluem reconhecimento facial, carros autônomos, controle de qualidade industrial, monitoramento de segurança, contagem e rastreamento de objetos, realidade aumentada, entre outros (Salles, 2022).

2.5 PROCESSAMENTO DIGITAL DE IMAGENS

O Processamento Digital de Imagens (PDI) abrange a análise e modificação de imagens digitalizadas com o objetivo de melhorar sua qualidade. As principais vantagens dessas técnicas incluem adaptabilidade, repetibilidade e precisão na preservação dos dados originais (Babu; Dasari; Yosepu, 2023).

O PDI abrange um conjunto de técnicas divididas em três grupos: compressão, aprimoramento e extração de reparo e medição. Concentrando-se no aprimoramento, essa etapa utiliza mecanismos para eliminar ruídos, controlar a nitidez e ampliar as imagens, facilitando a extração de recursos e a inspeção detalhada. Além disso, são aplicados filtros com o objetivo de destacar ou remover características específicas da imagem, tornando o processamento mais eficaz (Coady *et al.*, 2019).

Furusho *et al.* (2021) destacaram que algoritmos OCR baseados em alfabetos ocidentais frequentemente enfrentam dificuldades na detecção eficiente de caracteres japoneses, devido à grande diferença na forma de escrita desses sistemas. Para superar essa limitação, o estudo aplicou técnicas de PDI, como segmentação baseada em intervalos de cores, detecção de bordas e morfologia matemática. Essas técnicas foram essenciais para detectar placas informativas de trânsito japonesas, corrigir perspectivas distorcidas e segmentar os caracteres nelas contidos.

2.6 RECONHECIMENTO ÓPTICO DE CARACTERES

O Reconhecimento Óptico de Caracteres (OCR) é uma tecnologia fundamental no processamento de documentos digitais, permitindo a conversão de imagens ou arquivos digitalizados em texto editável (Farias, 2023).

Por meio de técnicas avançadas de OCR, é possível extrair, com eficácia, o conteúdo textual presente em imagens nítidas e com vocabulário contemporâneo (Silva, 2023). Essa funcionalidade não apenas simplifica a tarefa de transformar documentos físicos em formatos digitais, mas também viabiliza a análise e o tratamento de grandes volumes de dados de forma eficiente e precisa.

Entretanto, conforme destaca Thiyagaraj (2020), o uso de OCR enfrenta desafios significativos, especialmente em ambientes não controlados, como cenas naturais. Nesses contextos, fatores como distorções geométricas, fundos complexos e variação nas fontes tipográficas dificultam a correta interpretação dos textos.

Apesar dessas limitações, a tecnologia continua a demonstrar um potencial expressivo, sendo aplicada em diversos cenários — desde o reconhecimento de placas veiculares até a digitalização de faturas, recibos, cardápios e cartões de identificação (Thiyagaraj, 2020). Essas aplicações evidenciam a versatilidade do OCR e sua relevância crescente na modernização da gestão de documentos e informações em diferentes setores.

2.6.1 EASYOCR

O *EasyOCR* é uma biblioteca desenvolvida em Python, voltada para a extração de texto a partir de imagens. Trata-se de uma ferramenta versátil, capaz de reconhecer caracteres tanto em documentos digitalizados quanto em imagens de cenas naturais. Atualmente, oferece suporte a mais de 80 idiomas e está em constante aprimoramento (Jaided, 2024).

Segundo Rao *et al.* (2024), a integração do EasyOCR garante alta precisão na leitura de informações, mesmo em cenários de grande variabilidade. Essa abordagem abrangente não apenas proporciona uma interpretação precisa do conteúdo textual presente nas imagens, mas

também aumenta a robustez do sistema frente às condições e desafios impostos pelo mundo real.

2.7 EXPRESSÕES REGULARES

As expressões regulares, também conhecidas como *regex*, são recursos eficazes para identificar, validar e manipular padrões em textos. Muito utilizadas em diversas linguagens de programação, como o *C#*, permitem realizar operações como buscas, substituições e validações de forma rápida e precisa. No contexto da linguagem *C#*, as expressões regulares são especialmente úteis para criar soluções robustas que lidam com diferentes formatos de entrada de dados (Sabbag Filho, 2025).

A utilização conjunta de OCR e expressões regulares (*Regex*) configura uma estratégia eficiente no processamento de documentos, pois possibilita não apenas a conversão de textos contidos em imagens para formato digital, mas também a aplicação de padrões específicos para localizar informações relevantes. Dessa forma, a integração entre OCR e *Regex* se mostra uma solução eficaz para automatizar tarefas que envolvem a extração de dados a partir de documentos (Gonçalves, 2023).

2.8 TRABALHOS RELACIONADOS

Apesar dos avanços tecnológicos e das oportunidades oferecidas pela digitalização na cadeia de suprimentos, muitas empresas varejistas ainda dependem de métodos manuais para a gestão de dados de fornecedores, como o uso de cartões de visita físicos. Segundo o Sebrae (2024), com a adoção de cartões de visita digitais, a organização e o gerenciamento de contatos tornam-se significativamente mais eficientes. A integração com ferramentas de gestão permite categorizar, pesquisar e acompanhar interações de forma mais ágil e precisa, contribuindo para uma rede de contatos mais estruturada e acessível. Essa perspectiva reforça a proposta deste trabalho, que visa automatizar a extração e o armazenamento de informações contidas em cartões físicos, promovendo maior controle, eficiência e segurança na gestão dos dados de fornecedores.

Pesquisas anteriores sobre sistemas de OCR para cartões de visita têm se concentrado em várias técnicas, como a segmentação de imagens e a detecção de regiões de interesse (ROIs), para melhorar a extração de texto. Esses estudos exploraram diferentes abordagens para aumentar a precisão e a eficiência na leitura e interpretação das informações contidas nos cartões de visita.

Bhatt *et al.* (2023) desenvolveram uma aplicação que utilizou técnicas de pré-processamento de imagem, incluindo a remoção de ruídos, transformações morfológicas e a transformação de linhas de *Hough*, para limpar e remover linhas indesejadas das imagens de cartões de visita. O sistema identificou as Regiões de Interesse (ROIs) com precisão de até 80%, facilitando a extração de texto. Além disso, expressões regulares foram usadas para encontrar padrões específicos, e o método de reconhecimento de entidades nomeadas (REN) foi empregado para identificar e organizar nomes, endereços e outras informações relevantes.

Os resultados mostraram que as expressões regulares identificaram corretamente 92% dos números de telefone e 90% dos endereços de e-mail. O REN identificou 75% das entidades nomeadas no texto, destacando a eficácia do sistema na organização e extração de informações críticas. Esses resultados demonstram a capacidade do sistema em lidar com a variabilidade e complexidade dos dados presentes em cartões de visita, superando métodos tradicionais. Para validar o sistema, Bhatt *et al.* (2023) realizaram uma comparação dos resultados com rótulos extraídos manualmente das imagens dos cartões de visita.

De forma semelhante, Madan e Brindha (2019) conduziram um estudo que empregou técnicas de processamento de imagens, como o método de componentes conectados e o recorte de bordas, para segmentar e extrair texto de cartões de visita com maior precisão. Utilizaram também expressões regulares para identificar e classificar informações específicas extraídas dos cartões, alcançando uma precisão de 98% com o *Tesseract*³ OCR.

Os resultados do estudo de Madan e Brindha (2019) mostraram uma precisão notável: 97,3% na identificação dos nomes, 97% dos emails e 97,3% nos números de telefone. Esses resultados destacam a eficácia das técnicas utilizadas em comparação com métodos tradicionais de leitura linha por linha.

O método de verificação da precisão proposto por Madan e Brindha (2019) envolveu o processamento em lote de todas as imagens no diretório. Neste método, cada saída foi marcada como 1 se não fosse nula e 0 se fosse nula. A média dessas marcas foi calculada para todos os campos, permitindo a determinação da precisão geral do sistema. Essa abordagem assegurou uma avaliação rigorosa da eficácia do sistema, demonstrando sua precisão em condições práticas, com uma precisão geral de 98%.

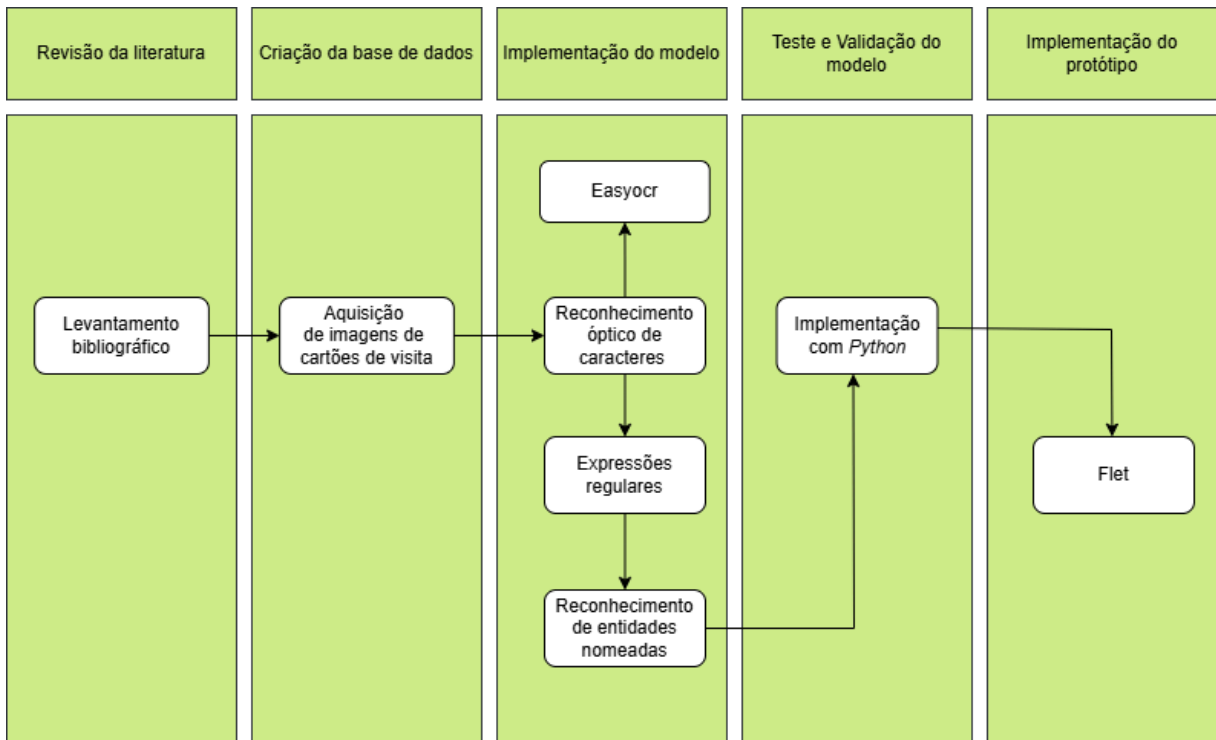
Esses trabalhos demonstram o potencial da aplicação de OCR com técnicas complementares para a extração precisa de dados em cartões de visita. No entanto, muitos ainda se concentram em etapas isoladas do processo ou não oferecem uma interface interativa para organização dos dados extraídos. A proposta deste trabalho avança nesse sentido, ao integrar um sistema de OCR com um protótipo de interface gráfica, voltada para a organização de dados de fornecedores no contexto do varejo de vestuário.

3 METODOLOGIA

Uma vez que o objetivo do projeto é implementar uma ferramenta computacional baseada em técnicas de Reconhecimento Óptico de Caracteres (OCR) para identificar e extrair informações de fornecedores no setor varejista, nesta seção serão abordados os procedimentos metodológicos utilizados no seu desenvolvimento. O processo foi dividido em 5 etapas, conforme apresentado na Figura 1.

³ <https://tesseract-ocr.github.io/>

Figura 1- Fluxograma das etapas do desenvolvimento



Fonte: Autoria própria (2025)

3.1 REVISÃO DA LITERATURA

Inicialmente, realizou-se um levantamento bibliográfico com o objetivo de identificar e reunir informações relevantes na literatura acerca da aplicação de técnicas de Reconhecimento Óptico de Caracteres (OCR), Processamento de Linguagem Natural (PLN) e Expressões Regulares na extração de informações contidas em cartões de visita. Essa etapa foi fundamental para embasar teoricamente o desenvolvimento do projeto, assegurando a adoção de boas práticas e a utilização das técnicas mais adequadas disponíveis na área.

3.2 CRIAÇÃO DA BASE DE DADOS

A segunda etapa da metodologia consiste na coleta e construção da base de dados, essencial para garantir a qualidade e a eficácia do modelo de extração de dados. Para isso, foi realizada a aquisição de um conjunto de 21 imagens representativas de cartões de visita de fornecedores no setor varejista. Essas imagens foram coletadas a partir de duas fontes distintas: algumas foram capturadas utilizando câmeras de celulares em diferentes condições de iluminação e ângulos, enquanto outras foram obtidas de fontes externas na internet. Essa abordagem assegura uma diversidade nas imagens, representando diferentes cenários e condições.

No entanto, o total de 21 imagens é um fator limitante, pois, apesar de fornecer uma boa variedade de exemplos, o número reduzido pode comprometer a capacidade do modelo de generalizar para uma maior variedade de situações e tipos de cartões. A ampliação da base de dados será crucial para melhorar a robustez e a precisão do modelo em cenários mais complexos e em imagens com diferentes características.

3.3 IMPLEMENTAÇÃO DO MODELO

A implementação do modelo foi realizada utilizando a linguagem de programação Python⁴. Para a extração de texto das imagens, foi desenvolvido um sistema de OCR com a biblioteca *EasyOCR*⁵, configurada para lidar com diferentes tipos de imagens e garantir uma extração eficiente. Após essa extração, expressões regulares são aplicadas para identificar e extrair informações específicas, como números de telefone e perfis de Instagram dos fornecedores.

Além disso, técnicas de Processamento de Linguagem Natural (PLN), como o Reconhecimento de Entidades Nomeadas (NER), foram utilizadas para identificar e extrair dados mais complexos, como os nomes das empresas, a partir dos textos extraídos. Esse conjunto de técnicas possibilita uma análise precisa e estruturada das informações dos fornecedores, otimizando a organização e o gerenciamento dos dados.

3.4 TESTE E VALIDAÇÃO

Para verificar a eficácia do sistema desenvolvido, foi conduzido um processo de teste e validação utilizando a base de dados construída previamente. Cada imagem de cartão de visita foi processada individualmente, e os textos extraídos passaram por uma análise comparativa com os valores reais esperados, previamente registrados em uma planilha de validação.

A avaliação dos resultados foi realizada com base em cinco métricas principais: confiança média, taxa de extração correta (TEC), precisão, revocação e F1-score. A confiança média foi fornecida automaticamente pela biblioteca *EasyOCR* para cada item extraído, indicando o grau de certeza do modelo. A TEC foi calculada como a proporção de campos corretamente extraídos em relação ao total de campos esperados por cartão.

Já as métricas de precisão, revocação e F1-score foram aplicadas para cada tipo de informação (telefone, nome e Instagram), considerando uma abordagem binária, onde o campo era classificado como correto apenas quando havia correspondência exata entre o valor extraído e o valor esperado. Os resultados foram organizados em planilhas para facilitar a análise quantitativa do desempenho do sistema.

3.5 IMPLEMENTAÇÃO DO PROTÓTIPO

Com o modelo se mostrando capaz de extrair e organizar as informações relevantes a partir das imagens, a etapa seguinte consistiu na construção de uma interface gráfica intuitiva e funcional, voltada para facilitar a interação do usuário com os dados extraídos. Para isso, utilizou-se a biblioteca *Flet*⁶, baseada em Python, que possibilita o desenvolvimento de interfaces leves, responsivas e com aparência moderna.

⁴ <https://www.python.org/>

⁵ <https://www.jaided.ai/easyocr/>

⁶ <https://flet.dev/>

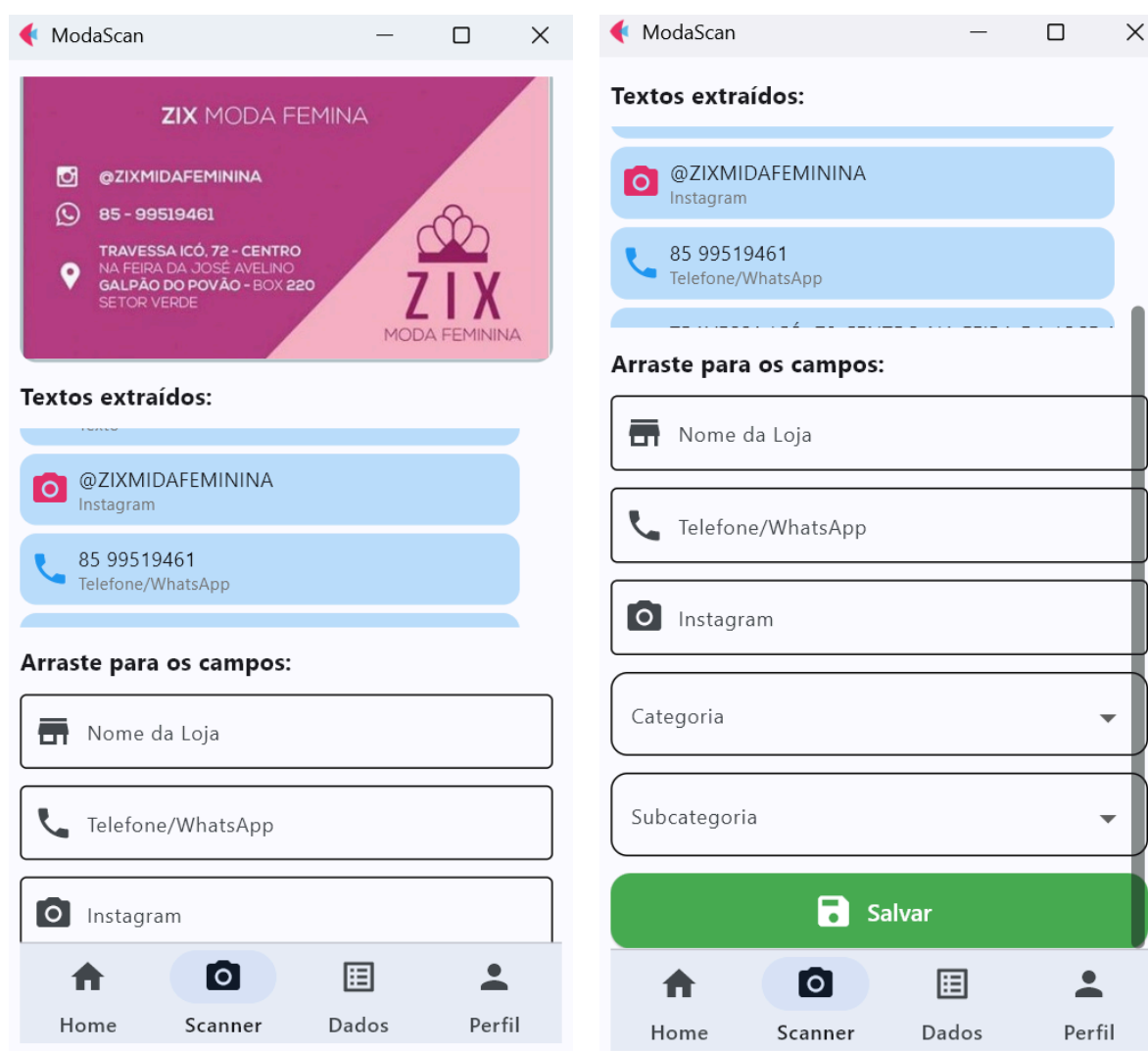
4 VALIDAÇÃO DOS RESULTADOS

A interface foi projetada com foco na simplicidade de uso, permitindo que qualquer usuário, mesmo sem conhecimentos técnicos, consiga organizar rapidamente os dados extraídos de um cartão de visita. Como ilustrado na Figura 2, logo após a digitalização, os textos identificados pelo sistema — como nome da loja, telefone e Instagram — são exibidos em uma lista, com ícones visuais que facilitam a identificação de cada categoria.

Dentre os principais recursos implementados, destaca-se a funcionalidade de “drag and drop” (arrastar e soltar), que permite ao usuário simplesmente arrastar os dados extraídos para os campos correspondentes. Por exemplo, o nome da loja pode ser arrastado para o campo “Nome da Loja”, o número de telefone para “Telefone/WhatsApp” e o perfil do Instagram para o campo correspondente. Essa abordagem torna o processo de categorização mais intuitivo, rápido e eficiente, eliminando a necessidade de digitação manual.

Além disso, a interface conta com campos adicionais, como Categoria e Subcategoria, que visam organizar melhor os dados das lojas extraídas, permitindo a futura ampliação do sistema para fins de cadastro, filtro ou segmentação de contatos.

Figura 2 - Interface gráfica do protótipo desenvolvido



Fonte: Autoria própria (2025)

Após a implementação da interface, a eficácia da ferramenta foi validada por meio de métricas quantitativas, que avaliam o desempenho do sistema na tarefa de extração de informações de cartões de visita. As métricas utilizadas incluem confiança média, taxa de extração correta (TEC), precisão, revocação e *F1-score*, sendo que cada uma dessas métricas oferece uma perspectiva complementar sobre a qualidade das extrações realizadas. A Tabela 1 apresenta a consolidação dessas métricas em uma planilha, com os valores obtidos, permitindo uma análise detalhada do desempenho do sistema.

Tabela 1 - Resultados das métricas de validação da ferramenta

MÉTRICA	%	Cartões
Confiança média	84,76%	
Taxa de extração correta (TEC)	75,79%	
Precisão Nome	42,85%	
Revocação Nome	42,85%	
F1-Score Nome	42,85%	
Precisão Instagram	71,42%	
Revocação Instagram	71,42%	
F1-Score Instagram	71,42%	
Precisão Telefone	80,95%	
Revocação Telefone	80,95%	
F1-Score Telefone	80,95%	

Fonte: Autoria própria (2025)

A confiança média é um valor atribuído pelo próprio *EasyOCR* a cada item reconhecido, variando de 0 a 1. Esse valor expressa o grau de certeza do modelo em relação ao texto extraído. No presente trabalho, esses dados foram convertidos em porcentagens para facilitar a análise e interpretação dos resultados. A confiança média obtida foi de 84,76%, o que indica um alto nível de confiabilidade nas leituras realizadas. Esse percentual demonstra que, na maioria dos casos, o modelo interpretou os textos com segurança, mesmo diante de variações visuais presentes nas imagens dos cartões.

A Taxa de Extração Correta (TEC) quantifica a proporção de campos corretamente identificados em relação ao total de campos esperados por imagem. A fórmula para calcular a TEC é apresentada na Equação 1.

$$TEC = \frac{CAMPOS\ CORRETOS}{CAMPOS\ ESPERADOS} \quad (1)$$

Essa métrica é especialmente útil para sintetizar o desempenho geral do modelo em cada cartão processado, sendo calculada pela razão entre campos corretos e campos esperados. A média geral da TEC foi de 75,79%, o que reforça a eficiência do sistema na extração automática de dados relevantes com um desempenho satisfatório.

As métricas de precisão, revocação e *F1-score* também foram aplicadas para avaliar a qualidade da extração por campo. A precisão mensura a proporção de acertos em relação ao total de elementos detectados pelo sistema, ou seja, o quão assertivo ele foi ao reconhecer uma informação como válida. Sua fórmula está representada na Equação 2.

$$\text{Precisão} = \frac{\text{Verdadeiros Positivos}}{\text{Verdadeiros Positivos} + \text{Falsos Positivos}} \quad (2)$$

Os verdadeiros positivos (VP) referem-se aos casos em que o sistema identificou corretamente a informação, como, por exemplo, um número de telefone presente na imagem. Por outro lado, os falsos positivos (FP) são os casos em que o sistema detectou algo como correto, mas estava errado. Um exemplo disso seria quando o sistema identifica um número como telefone, mas, na realidade, ele não é um número válido.

A revocação, por sua vez, verifica o quanto o sistema conseguiu identificar do total de informações que efetivamente estavam presentes na imagem. A fórmula da revocação está representada na Equação 3.

$$\text{Revocação} = \frac{\text{Verdadeiros Positivos}}{\text{Verdadeiros Positivos} + \text{Falsos Negativos}} \quad (3)$$

Os falsos negativos (FN) são os casos em que o sistema falha em identificar algo que deveria ter sido reconhecido, como quando o sistema não identifica um número de telefone que estava claramente na imagem.

O F1-score é a média harmônica entre precisão e revocação, equilibrando ambas as métricas e sendo especialmente relevante em contextos em que é necessário minimizar tanto os erros por excesso quanto os por omissão. Sua fórmula está ilustrada na Equação 4.

$$F1 - Score = 2 \times \frac{\text{Precisão} \times \text{Revocação}}{\text{Precisão} + \text{Revocação}} \quad (4)$$

Durante a análise dessas três métricas, foi possível observar que os valores se repetiram para cada campo avaliado. Esse comportamento se deve à abordagem binária adotada na validação, em que o campo é considerado correto apenas quando há uma correspondência exata entre o texto extraído e o valor esperado. Dessa forma, a inexistência de falsos positivos e falsos negativos faz com que as métricas de precisão, revocação e *F1-score* coincidam.

No campo nome, foi obtido o desempenho mais baixo, com 42,85% em todas as métricas. Esse resultado pode ser explicado pelas variações tipográficas e decorativas utilizadas nos nomes comerciais dos cartões, que frequentemente confundem o modelo de OCR. Além disso, a qualidade de algumas imagens influenciou negativamente a extração, comprometendo a acurácia do reconhecimento.

Já no campo Instagram, os resultados foram mais positivos, com 71,42% nas três métricas. Como ilustrado na Figura 3, mesmo quando o sistema apresenta alta taxa de confiança na leitura do texto — neste caso, 100% —, a ausência do caractere "@" pode impedir o reconhecimento correto de um perfil do Instagram. O exemplo mostra a detecção do nome de usuário "*modascamila*", mas o sistema não o identificou como uma rede social, tratando-o apenas como texto genérico.

Isso acontece porque o algoritmo depende de expressões regulares (regex) que buscam padrões específicos, como "*@nomeusuario*", para classificar corretamente a informação como um perfil de rede social. Sem o "@", o padrão não é reconhecido, e o dado acaba sendo tratado como pertencente a outra categoria, como "nome".

Além disso, caracteres comuns em nomes de usuário de redes sociais, como underlines, também são ignorados se o contexto não for identificado corretamente. Portanto,

mesmo com confiança elevada na leitura, o modelo pode falhar na classificação semântica do campo, afetando diretamente as métricas de revocação e *F1-score* relacionadas ao Instagram.

Figura 3 - Exemplo instagram



Fonte: Autoria própria (2025)

Por fim, o campo telefone apresentou o melhor desempenho, com 80,95% em precisão, revocação e *F1-score*. A principal dificuldade identificada ocorreu em casos onde o número telefônico foi representado com diferentes fontes entre o DDD e o restante do número, ou com variação de cor, o que afetou a consistência na segmentação do texto. Ainda assim, os resultados demonstram que o sistema é bastante robusto na identificação de dados numéricos padronizados, validando sua aplicabilidade prática no contexto varejista.

5 CONSIDERAÇÕES FINAIS

Este trabalho teve como objetivo propor uma ferramenta computacional voltada para o setor varejista, utilizando técnicas de Reconhecimento Óptico de Caracteres (OCR), Processamento de Linguagem Natural (PLN) e expressões regulares para automatizar a extração de informações contidas em cartões de visita de fornecedores.

Após o desenvolvimento e validação do sistema, constatou-se que a ferramenta é capaz de realizar a extração de dados com um bom nível de confiabilidade, especialmente em informações padronizadas, como números de telefone. No entanto, observou-se que o desempenho pode ser impactado por variações tipográficas e pela qualidade das imagens utilizadas.

Apesar dessas limitações, a ferramenta representa um avanço importante na modernização dos processos de gestão no varejo. Para aprimoramentos futuros, recomenda-se ampliar a base de dados, adotar técnicas adicionais de pré-processamento de imagens e incorporar novos recursos de inteligência artificial que permitam melhorar a precisão e ampliar as funcionalidades do sistema.

Com isso, espera-se que a solução proposta contribua para a digitalização e automação da gestão dos dados dos fornecedores, reduzindo erros manuais e otimizando o tempo dedicado à organização das informações. Além disso, este trabalho oferece um ponto de partida relevante para pesquisas futuras e desenvolvimentos que busquem fortalecer a integração entre tecnologia e práticas de gestão no comércio varejista de artigos de vestuário.

REFERÊNCIAS

ANDRADE, Claudio MV de et al. PromptNER: Uma Abordagem para Reconhecimento de Entidades Nomeadas em Dados Sensíveis a Partir de Instâncias Rotuladas Automaticamente. In: **Anais do XXXVIII Simpósio Brasileiro de Bancos de Dados**. SBC, 2023. p. 269-281.

BABU, K Ganapathi; DASARI, Sunil Kumar; YOSEPU, C. An Overview: image processing techniques and its applications. **International Journal Of Engineering Technology And Management Sciences**, [S.L.], v. 7, n. 3, p. 883-888, 2023. Mallikarjuna Infosys. <http://dx.doi.org/10.46647/ijetms.2023.v07i03.135>.

BARBOSA, Lucia Martins; PORTES, Luiza Alves Ferreira. Inteligência artificial. **Revista Tecnologia Educacional, Rio de Janeiro**, n. 236, p. 16-27, 2023.

BHATT, Chandradeep; SAYANA, Akash; CHAUHAN, Rahul; SINGH, Teekam. Text Extraction & Recognition from Visiting Cards. **2023 Third International Conference On Secure Cyber Computing And Communication (Icscce)**, [S.L.], p. 703-708, 26 maio 2023. IEEE. <http://dx.doi.org/10.1109/icscce58608.2023.10176809>.

BRANDÃO, Raul de Souza. Inteligência artificial na melhoria de textos científicos: Aplicações, benefícios e desafios. **Revista de Empreendedorismo e Gestão de Micro e Pequenas Empresas**, v. 9, n. 01, p. 141-149, 2024.

COADY, James; O'RIORDAN, Andrew; DOOLY, Gerard; NEWE, Thomas; TOAL, Daniel. An Overview of Popular Digital Image Processing Filtering Operations. **2019 13Th International Conference On Sensing Technology (Icst)**, [S.L.], p. 1-5, dez. 2019. IEEE. <http://dx.doi.org/10.1109/icst46873.2019.9047683>.

ESTEVA, Andre; KUPREL, Brett; NOVOA, Roberto A.; KO, Justin; SWETTER, Susan M.; BLAU, Helen M.; THRUN, Sebastian. Dermatologist-level classification of skin cancer with deep neural networks. **Nature**, [S.L.], v. 542, n. 7639, p. 115-118, 25 jan. 2017. Springer Science and Business Media LLC. <http://dx.doi.org/10.1038/nature21056>.

FARIAS, Walisson Nascimento de et al. Aperfeiçoando o reconhecimento óptico de caracteres em imagens de documentos pessoais. 2023.

FONTANA, Éliton. Introdução aos algoritmos de aprendizagem supervisionada. Departamento de Engenharia Química, Universidade Federal do Paraná, 2020.

FURUSHO, Rafael Yuji Hirata et al. APLICAÇÃO DE REDES NEURAIAS CONVOLUCIONAIS NO RECONHECIMENTO DE CARACTERES EM PLACAS INFORMATIVAS JAPONESAS. In: **Colloquium Exactarum**. ISSN: 2178-8332. 2021. p. 12-24.

GONÇALVES, Ana Julia P. et al. Um protótipo de software para mineração de dados de contas de prestadoras de serviços de telefonia. 2023.

JAIDED AI. EasyOCR, 2024. Disponível em: <https://www.jaided.ai/easyocr/>. Acesso em: 11 jul. 2025.

JIANG, Yuan; QIN, Jianwen; KHAN, Hayat. The Effect of Tax-Collection Mechanism and Management on Enterprise Technological Innovation: evidence from china. **Sustainability**, [S.L.], v. 14, n. 14, p. 8836, 19 jul. 2022. MDPI AG. <http://dx.doi.org/10.3390/su14148836>.
JORDAN, Dyonathan; CESAR, Victor; CARMO, Lucas do. Análise de técnicas de PLN e machine learning na detecção de Indícios de depressão. 2022.

KIRSCH, Bernardo Gularte; DORNELES, Ártton Pereira. Desenvolvimento de uma Ferramenta para Reconhecimento de Entidades Nomeadas em Certificados de Atividades Complementares de Curso utilizando spaCy. **Anais do Encontro Anual de Tecnologia da Informação**, v. 12, n. 1, p. 44-44, 2023.

LI, Peixuan. Study on Data-driven Digital New Retail Supply Chain Innovation. **Frontiers In Business, Economics And Management**, [S.L.], v. 14, n. 1, p. 246-249, 21 mar. 2024. Darcy & Roy Press Co. Ltd.. <http://dx.doi.org/10.54097/8f9sj210>.

MADAN KUMAR, C.; BRINDHA, M. Extração de texto de cartões de visita e classificação do texto extraído em classes pré-definidas. **Revista Internacional de Inteligência Computacional & IoT**, v. 3, 2019.

MANDELLI, PEDRO. Aprendizado Não Supervisionado: Entenda essa técnica de ML, 2023. Disponível <https://domineia.com/aprendizado-nao-supervisionado/>. Acesso em: 02 Jun 2024.

OGUNLEYE, Julius Olufemi. Predictive data analysis using linear regression and random forest. In: **Data integrity and data governance**. IntechOpen, 2022
PIRES, Ricardo Lucas. Visão computacional aplicada na gestão da tecnologia da informação. 2023.

RAO, Tabish; RAWAT, Shivam; GUPTA, Atika; MEHRA, Nidhi; DHONDIYAL, Shiv Ashish; KAPIL, Divya. Comprehensive study on automatic number plate recognition with python using OpenCV and EasyOCR. **Challenges In Information, Communication And Computing Technology**, [S.L.], p. 501-505, 19 nov. 2024. CRC Press.
<http://dx.doi.org/10.1201/9781003559092-86>.

SABBAG FILHO, Nagib. Aplicações Práticas de Expressões Regulares para Validação em Csharp. **Leaders Tec**, v. 2, n. 1, 2025.

SALLES, A. O que é Visão Computacional e para que serve? 2022. Disponível em: <https://santodigital.com.br/o-que-e-visao-computacional-e-para-que-serve/>. Acesso em: 26 de maio de 2024.

SANTOS, Andrei Camilo dos. Aprendizado de máquina aplicado na detecção de fraudes em cartão de crédito. 2023.

SARTORI, A. Gestão de fornecedores: 9 práticas para fazer com eficácia, 2023 . Disponível em: <https://qualyteam.com/pb/blog/descubra-como-gerenciar-fornecedores-com-eficacia/>. Acesso em: 25 jun 2024.

SEBRAE. Por que devo aderir aos novos cartões de visita digitais?, 2024. Disponível em: <https://respostas.sebrae.com.br/por-que-devo-aderir-aos-novos-cartoes-de-visita-digitais/>. Acesso em: 13 jul 2025.

SHAVETA. A review on machine learning. **International Journal Of Science And Research Archive**, [S.L.], v. 9, n. 1, p. 281-285, 30 maio 2023. GSC Online Press. <http://dx.doi.org/10.30574/ijrsra.2023.9.1.0410>.

SILVA, Diana Marcela da. Provendo acesso ao conteúdo de documentos centenários: um processo de correção e melhoria do texto extraído de imagens utilizando técnicas de processamento de linguagem natural. 2023. Dissertação de Mestrado. Universidade Federal de Pernambuco.

SILVA, Thomas S. e; FERNANDES, Arlindo; FERNANDES, Sílvio. GrayScaleAccel: acelerador de escala de cinza em fpga. **Anais Estendidos do XXI Simpósio em Sistemas Computacionais de Alto Desempenho (Sscad Estendido 2020)**, [S.L.], p. 14-21, 21 out. 2020. Sociedade Brasileira de Computação - SBC. http://dx.doi.org/10.5753/wscad_estendido.2020.14084.

STRAUSS, Edilberto; VILLAS BÔAS JÚNIOR, Manoel; FERREIRA, Wagner Luiz Lobo. A importância de utilizar métricas adequadas de avaliação de performance em modelos preditivos de machine learning. *Projectus*, v. 7, n. 2, p. 52-62, 202

THIYAGARAJ, S. A comprehensive guide to OCR with Tesseract, OpenCV and Python, 2020. Disponível em: <https://medium.com/nanonets/a-comprehensive-guide-to-ocr-with-tesseract-opencv-and-python-fd42f69e8ca8>. Acesso em: 09 maio 2024

TORRES, Marcelo. Eficiência operacional e o caminho para a competitividade, 2024. Disponível em: <https://conectecomunicacao.com.br/eficiencia-operacional-e-o-caminho-para-a-competitividade/>. Acesso em: 25 jun 2024.

WANG, Junmiao. A Study of The OCR Development History and Directions of Development. **Highlights In Science, Engineering And Technology**, [S.L.], v. 72, p. 409-415, 15 dez. 2023. Darcy & Roy Press Co. Ltd.. <http://dx.doi.org/10.54097/bm665j77>.