



UNIVERSIDADE EDUARDO MONDLANE
FACULDADE DE ENGENHARIA
DEPARTAMENTO DE ENGENHARIA ELECTROTÉCNICA
CURSO DE ENGENHARIA INFORMÁTICA

**DESENVOLVIMENTO DE UMA EXTENSÃO DE NAVEGADOR PARA ANÁLISE DE
CREDIBILIDADE DE ANÚNCIOS DE EMPREGO NO LINKEDIN UTILIZANDO
APRENDIZAGEM DE MÁQUINA E PROCESSAMENTO DE LINGUAGEM NATURAL (NLP)**

Autor:

Uqueio Jr., Delfim Luís

Supervisor

Eng.º Rúben Manhiça

Maputo, Maio de 2025



UNIVERSIDADE EDUARDO MONDLANE
FACULDADE DE ENGENHARIA
DEPARTAMENTO DE ENGENHARIA ELECTROTÉCNICA
CURSO DE ENGENHARIA INFORMÁTICA

**DESENVOLVIMENTO DE UMA EXTENSÃO DE NAVEGADOR PARA ANÁLISE DE
CREDIBILIDADE DE ANÚNCIOS DE EMPREGO NO LINKEDIN UTILIZANDO
APRENDIZAGEM DE MÁQUINA E PROCESSAMENTO DE LINGUAGEM NATURAL (NLP)**

Autor:

Uqueio Jr., Delfim Luís

Supervisor

Eng.º Rúben Manhiça

Maputo, Maio de 2025



UNIVERSIDADE EDUARDO MONDLANE
FACULDADE DE ENGENHARIA
DEPARTAMENTO DE ENGENHARIA ELECTROTÉCNICA
CURSO DE ENGENHARIA INFORMÁTICA

TERMO DE ENTREGA DE RELATÓRIO DE TRABALHO DE LICENCIATURA

Declaro que o estudante Delfim Luís Uqueio Júnior entregou no dia ____, as 03 cópias do relatório do seu trabalho de licenciatura com referência “REFERÊNCIA” intitulado: **Desenvolvimento de uma extensão de navegador para Análise de Credibilidade de anúncios de emprego no LinkedIn utilizando aprendizagem de máquina e processamento de linguagem natural (NLP)**

Maputo, Maio de 2025

Chefe da secretaria



UNIVERSIDADE EDUARDO MONDLANE
FACULDADE DE ENGENHARIA
DEPARTAMENTO DE ENGENHARIA ELECTROTÉCNICA
CURSO DE ENGENHARIA INFORMÁTICA

DECLARAÇÃO DE HONRA

Declaro sob compromisso de honra que o presente trabalho é resultado da minha investigação e que foi concebido para ser submetido apenas para obtenção do grau de Licenciatura em Engenharia Informática na Faculdade de Engenharia da Universidade Eduardo Mondlane

Maputo, Maio de 2025

O Autor

(Delfim Luís Uqueio Júnior)

DEDICATÓRIA

Aos meus Pais, Delfim Luís Uqueio e Olga Francisco Massuco

Às minhas irmãs, Yúnila, Ladinice, Zígnia e Ana

E à minha avó Olinda

Agradecimentos

Em primeiro lugar, bendito seja o Deus e Pai do nosso Senhor Jesus Cristo, que nos abençoou com todas as bênçãos espirituais nas regiões celestiais (Efésios 1.3).

Agradeço aos meus Pais, Delfim Luís Uqueio e Olga Francisco Massuco pelo suporte total em todo este processo desde o primeiro ano da Faculdade até ao fim. Obrigado Papá e Mamã pelo apoio em todos os sentidos. Lembro-me como se fosse hoje quando tomei a decisão de sair de uma outra universidade pública onde me encontrava e voltar a concorrer para Engenharia Informática na UEM; foi uma decisão que lhes abalou pois “perderia” um ano, mas eles acreditaram em mim e no que eu estava a fazer.

Agradeço também às minhas irmãs Ana e Zígnia pelo suporte e em especial as mais novas, Ladinice e Yúnila, por me apoiarem cuidando de mim e por se preocuparem comigo; espero que o culminar deste curso lhes sirva como inspiração.

Aos meus colegas de turma, alguns dos quais se tornaram grandes amigos. Em especial, à Mónica Macamo, que foi uma colega extraordinária e cuja contribuição foi fundamental para a conclusão deste trabalho. Mónica tornou-se não apenas na melhor amiga que eu poderia desejar, mas também na mulher com quem decidi partilhar a vida.

Ao corpo docente, especialmente ao meu supervisor Ruben Manhiça, pela paciência e alta disponibilidade sempre. Mesmo diante das suas ocupações, sempre arranjou um tempo para me ajudar neste processo.

Epígrafe

“Ninguém precisa continuar onde está”

- Clóvis de Barros Filho

Resumo

A crescente incidência de fraudes em anúncios de emprego disponibilizados em plataformas online tem vindo a configurar um desafio significativo, tanto para candidatos em busca de novas oportunidades profissionais como para empregadores que procuram recrutar de forma segura e eficiente. Este fenómeno compromete a confiança nas práticas de recrutamento digital, expondo os utilizadores a riscos de natureza financeira, emocional e reputacional.

Face a este cenário, o presente trabalho propõe o desenvolvimento de uma extensão de navegador destinada a auxiliar os utilizadores do *LinkedIn* na avaliação da credibilidade dos anúncios de emprego. A metodologia adoptada integra pesquisas documentais e a aplicação de formulários para levantamento de requisitos e necessidades dos utilizadores. A solução desenvolvida envolve ainda a construção de um modelo de aprendizagem automática, treinado com dados rotulados, capaz de identificar padrões e características indicativas de fraude em anúncios de emprego.

Os resultados esperados incluem o fortalecimento da segurança no ambiente de recrutamento online, oferecendo aos candidatos um instrumento de apoio que permita a tomada de decisões mais informadas, contribuindo assim para a mitigação dos riscos associados à exposição a anúncios fraudulentos.

Palavras-chave: fraude em recrutamento, anúncios de emprego, plataformas digitais, aprendizagem automática, extensão de navegador, *LinkedIn*.

Abstract

The increasing incidence of fraudulent job advertisements on online platforms has become a significant challenge for both candidates seeking new professional opportunities and employers aiming to recruit safely and efficiently. This phenomenon undermines trust in digital recruitment practices, exposing users to financial, emotional, and reputational risks.

In response to this scenario, the present study proposes the development of a browser extension designed to assist LinkedIn users in efficiently evaluating the credibility of job advertisements. The adopted methodology comprises documentary research and the application of user surveys to gather requirements and identify user needs. The proposed solution further involves the development of a machine learning model, trained on a labelled dataset, capable of identifying patterns and characteristics indicative of fraudulent job postings.

The expected outcomes include strengthening security within the online recruitment environment by providing candidates with a practical tool to support more informed decision-making, thereby helping to mitigate the risks associated with exposure to fraudulent advertisements.

Keywords: recruitment fraud, job advertisements, digital platforms, machine learning, browser extension, LinkedIn.

Índice

1	Capítulo I – Introdução	1
1.1	Contextualização	1
1.2	Definição do problema	3
1.3	Motivação	3
1.4	Objectivos	4
1.4.1	Objectivo Geral	4
1.4.2	Objectivos Específicos	4
1.5	Metodologia	5
1.5.1	Metodologia de pesquisa	5
1.5.2	Classificação da Metodologia	5
1.5.3	Técnicas de colecta de dados	8
1.5.4	Metodologia de desenvolvimento do trabalho prático	9
1.6	Estrutura do trabalho	10
2	Capítulo II - Revisão de Literatura	13
2.1	Plataformas online de anúncio de emprego	13
2.1.1	Conceito	13
2.1.2	Principais plataformas online de anúncio de emprego	13
2.1.3	Conceito de fraude em anúncios de emprego	14
2.1.4	História das fraudes em anúncios de emprego	15
2.1.5	Mecanismos de detecção de fraudes usados pelo LinkedIn	16
2.1.6	Constrangimento nos mecanismos actuais de detecção de fraudes em anúncios de emprego do LinkedIn	17
2.2	Principais navegadores utilizados para aceder ao LinkedIn	18
2.3	Extensão de navegadores e seu desenvolvimento para Google Chrome	18
2.4	Aprendizagem de máquina (Machine Learning - ML)	19
2.4.1	Inteligência artificial (IA)	19
2.4.2	Conceito de aprendizagem de máquina	20
2.4.3	Tipos de aprendizagem de máquina	21
2.4.4	Algoritmos de aprendizagem de máquina	24

2.4.5	Processamento de linguagem natural (NLP)	25
2.4.6	Relevância da inteligência artificial e aprendizagem de máquina na análise de credibilidade de anúncios de emprego	26
2.4.7	Fontes de dados para análise de credibilidade em anúncios de emprego	27
2.4.8	Ferramentas para implementação de aprendizagem de máquina	28
3	Capítulo III – Proposta de solução	30
3.1	Relevância da busca por uma solução	30
3.2	Escolha do navegador para implementação da extensão	31
3.2.1	Análise de utilização de navegadores	31
3.3	Funcionamento da extensão de navegador	32
3.3.1	Preferência do público-alvo	32
3.3.2	Limitações das abordagens nativas do LinkedIn	33
3.3.3	Fluxo de funcionamento da extensão	34
3.4	Repositório de dados e Dataset	34
3.5	Algoritmo de aprendizagem de máquina a utilizar	35
3.5.1	Fundamentacao com base na literatura técnica	35
3.5.2	Fundamentação com base nas características da solução	36
3.5.3	Fundamentação com base na pesquisa com o público-alvo	36
4	Capítulo IV - Treinamento do modelo computacional	38
4.1	Descrição dos dados	38
4.2	Pré-processamento dos dados	38
4.3	Treinamento e avaliação	39
4.4	Armazenamento do modelo	39
5	Capítulo V - Desenvolvimento da solução proposta	40
5.1	Elicitação dos requisitos	40
5.2	Análise da solução	40
5.2.1	Stackholders da solução	40
5.2.2	Requisitos do sistema	41
5.2.3	Modelagem da solução proposta	43
5.2.4	Modelos de casos de uso	44

5.2.5 Casos de uso	44
5.3 Projecto	44
5.4 Codificação do protótipo	47
5.4.1 Estrutura do protótipo	48
5.4.2 Testes do protótipo	49
6 Capítulo VI – Teste do algoritmo	52
6.1 Contextualização do Dataset	52
6.2 Metodologia	52
1. Pré-processamento textual	52
2. Divisão dos dados	52
6.3 Algoritmos avaliados e resultados	53
6.3.1 Regressão logística (modo padrão)	53
6.3.2 Regressão logística (modo balanceado)	54
7 Capítulo VII – Discussão de resultados	56
7.1 Revisão de literatura	56
7.2 Proposta de solução	57
8 Capítulo VIII – Considerações finais	59
8.1 Conclusões	59
8.2 Recomendações	60
Referências bibliográficas	62
Bibliografia	62
Outras bibliografias consultadas	67
Anexos	68
Anexo 1: Formulário	68

Lista de figuras

Figura 1: Participação de mercado dos navegadores para desktop (2024). Fonte: Adaptado de StatCounter Global Stats, disponível em: gs.statcounter.com	18
Figura 2: Aprendizagem de máquina como sub-área da inteligência artificial. Fonte: CSLEE Tech blog, 2022	20
Figura 3: Principais tipos de aprendizagem de máquina (Fonte: https://beatrizmaiads.medium.com)	21
Figura 4: Aprendizagem supervisionada (Fonte: Estatsite)	22
Figura 5: Teorema de Bayes.....	25
Figura 6: Repositório de dados (Fonte: Super Annotate).....	28
Figura 7: Triangulação metodológica (fonte: chatgpt)	30
Figura 8: Navegador utilizado para aceder ao LinkedIn.....	31
Figura 9: Forma escolhida para verificação de credibilidade em anúncios.....	32
Figura 10: Diagrama de sequência da solução proposta.....	34
Figura 11: Diagrama de casos de uso	44
Figura 12: Arquitectura de uma extensão do google chrome (Fonte: Publicação de Yoshi no Medium, 2022).....	45
Figura 13: Selecção do texto a ser analisado.....	50
Figura 14: a extensão do navegador, antes da realização da análise.	50
Figura 15: Resultado da análise de credibilidade de anúncio confiável.....	51
Figura 16: Resultado de análise de credibilidade de um anúncio não confiável.....	51
Figura 17: Divisão dos dados para teste (fonte: Guido Mascia)	53
Figura 18: Resultados da classificação (regressão logística padrão)	54
Figura 19: Resultados da classificação (regressão logística com balanceamento).....	54

Lista de Tabelas

Tabela 1: Comparação entre abordagens possíveis.....	33
Tabela 2: Comparação entre algoritmos de aprendizagem de máquina.....	36
Tabela 3: Stakeholders e os benefícios da solução proposta.....	41
Tabela 4: Descrição dos requisitos funcionais	42
Tabela 5: Descrição dos requisitos não funcionais	42
Tabela 6: Estrutura do protótipo.....	48
Tabela 7: Resultados da classificação (regressão logística padrão).....	54
Tabela 8: Resultados da classificação (regressão logística com balanceamento).....	55

Lista de abreviaturas e acrónimos

- **API:** *Application programming Interface*
- **CSS:** *Cascading*
- **DOM:** *Document Object Model*
- **EUA:** *Estados Unidos da América*
- **FTC:** *Federal Trade Comission*
- **HTML:** *HyperText Markup Language*
- **IA:** *Inteligência artificial*
- **JS:** *JavaScript*
- **JSON:** *JavaScript Object Notation*
- **MacOS:** *Macintosh Operating System*
- **ML:** *Machine Learning*
- **MVP:** *Minimum Viable Product*
- **NLP:** *Natural Language Processing*
- **OIT:** *Organização Internacional do Trabalho*
- **RF:** *Requisitos Funcionais*
- **RNF:** *Requisitos Não Funcionais*
- **SVM:** *Máquinas de vectores de suporte*
- **TD-IDF:** *Term Frequency – Inverse Document Frequency*
- **UEM:** *Universidade Eduardo Mondlane*
- **UI:** *User Interface*
- **UML:** *Unified Modelling Language*

Glossário de termos

- **Algoritmo:** conjunto de regras ou instruções sequenciais definidas para resolver um problema ou executar uma tarefa.
- **Anotação:** processo de rotular dados com informações relevantes.
- **API:** conjunto de rotinas e padrões de programação que permitem a comunicação e troca de dados entre diferentes componentes de *software*.
- **Backend:** A parte da aplicação responsável pelo processamento dos dados do negócio e comunicação maior parte das vezes com a base de dados.
- **Dataset:** Conjunto organizado de dados utilizado para treinar ou avaliar modelos de aprendizagem de máquina
- **Extensão:** módulo de *software* que adiciona funcionalidades a um navegador *web*, personalizando a experiência de navegação.
- **Feature:** atributo ou variável de entrada usada para treinar um modelo de *machine learning*.
- **Frontend:** a parte da aplicação com a qual utilizador interage directamente.
- **Inferência:** acto de aplicar um modelo treinado para prever resultados com novos dados.
- **Modelo:** representação matemática criada por algoritmos de aprendizagem de máquina a partir de dados.
- **Overfitting:** problema em que o modelo aprende demais os dados de treino, perdendo capacidade de generalização
- **Rótulo:** valor ou categoria associado a um dado de entrada, que indica a resposta correcta que o modelo deve aprender a prever.
- **Ruído:** informações irrelevantes ou inconsistentes que podem atrapalhar a análise de dados.
- **Stopwords:** palavras comuns que geralmente são removidas em análise textual.
- **Tokenização:** processo de dividir um texto em partes menores (*tokens*), como palavras ou símbolos.
- **Vectorização:** conversão de texto em representações numéricas para uso em modelos de aprendizagem de máquina.

1 Capítulo I – Introdução

1.1 Contextualização

A busca por emprego é um processo essencial para garantir a segurança financeira das pessoas. Em todo o mundo, as pessoas buscam diariamente por oportunidades de emprego; uma pesquisa feita no Reino Unido mostra que empregadores recebem, em média, 140 candidaturas por vaga aberta (Global Recruiter, 2024). Só nos Estados Unidos de América (EUA), o número de empregos publicados anualmente variou entre 28 e 35 milhões nos últimos anos e este número tende a aumentar (OBERLO, 2021).

Com o aumento da busca por emprego, surgem novos desafios no que diz respeito a gestão de anúncios de vagas de emprego, bem como o processo de candidatura em si (Feng, S. et al, 2021). Antigamente, as formas mais comuns de se anunciar vagas eram por meio de jornais, canais televisivos e rádios, e as pessoas precisavam submeter as suas candidaturas fisicamente. Actualmente, com o advento da *internet*, plataformas *online* como *LinkedIn*, *Indeed*, *Glassdoor*, emprego Moz, entre outras, publicam diversos anúncios de vagas, facilitando o acesso a oportunidades de emprego de formas mais rápida e fácil (FAROOQ, 2021; *International Labour Organization*, 2020).

Em Moçambique, o uso dessas plataformas digitais tem vindo a crescer, principalmente entre jovens recém-formados e profissionais à procura do primeiro emprego. De acordo com dados do Instituto Nacional de Estatística (INE), a taxa de desemprego jovem no País atingia cerca de 30% em 2023, o que intensifica a procura por oportunidades no sector formal e informal (INE, 2023). Nesse cenário, plataformas como emprego.co.mz, emprego.infromoz.com, LinkedIn e grupos de emprego em redes sociais como Facebook e o WhatsApp tornam-se alternativas populares. Contudo, essa maior acessibilidade também abre espaço para surgimento de práticas fraudulentas, uma vez que muitos utilizadores não possuem meios eficazes de verificação da autenticidade de anúncios (TechZim, 2022)

Segundo Yusoff et al. (2022), apesar da utilização de tecnologias digitais no recrutamento trazer benefícios significativos, como maior eficiência e alcance global, criou desafios relacionados à triagem e verificação da autenticidade das vagas publicadas. Com uma escala massiva de anúncios, segundo Zhao e Wang (2019), torna-se cada vez mais difícil filtrar as publicações fraudulentas. Além disso, hoje qualquer pessoa pode publicar uma vaga ou se passar por empresas e espalhar vagas ilegítimas com objectivo de obter vantagens financeiras, enganando os candidatos. Com a crescente falta de emprego a nível global, muitos candidatos acabam acreditando que encontraram a oportunidade ideal, mas que no final resulta em dados financeiros por conta destes golpes. De acordo com um relatório da *Federal Trade Commission* (FTC), as fraudes envolvendo oportunidades de emprego aumentaram entre 2019 e 2023, com perdas financeiras chegando a US\$ 490.7 milhões, refletindo um problema crescente de fraudes no sector de recrutamento (FTC, 2023).

Há relatos frequentes de burlas relacionadas a ofertas de emprego falsas, principalmente em plataformas informais e redes sociais, onde os candidatos são instruídos a pagar valores para garantir uma vaga ou participar em supostos processos selectivos (Jornal Notícias, 2022). A fraca regulação digital e a limitação de mecanismos automatizados de verificação contribuem para o agravamento do problema.

Dessa forma, surge a necessidade de buscar mitigar essas fraudes; para Muller, J. et al (2020), o uso de tecnologias de detecção automática de fraudes é essencial. O *LinkedIn*, uma das maiores plataformas *online* de emprego, tem investido em mecanismos tecnológicos e humanos para combater esse problema. Este trabalho propõe o desenvolvimento de uma solução eficaz para fazer análise de credibilidade em anúncios de empregos no *LinkedIn*, com recurso a inteligência artificial, concretamente a aprendizagem de máquina e uma extensão de navegador, tendo em conta que 42.8% dos utilizadores tem acesso ao *LinkedIn* por meio de um navegador (CHARLE AGENCY, 2024) e muitos dos navegadores, por sua vez, são abertos ao desenvolvimento de extensões de terceiros (MOZILLA, 2024).

1.2 Definição do problema

As plataformas de anúncio emprego ganharam uma grande notoriedade que actualmente ocupam maior percentagem das vagas que são publicadas no mundo inteiro, substituindo métodos tradicionais como jornais e rádios (BERNHARD, 2022). Existe um inúmero volume de vagas publicadas diariamente em plataformas como *LinkedIn*, o que torna cada vez mais difícil realizar o filtro de publicações falsas, pois os algoritmos responsáveis por esta tarefa têm o desafio de lidar com um volume gigantesco de dados. De acordo com McKinsey & Company (2017), o uso de algoritmos avançados para tarefas críticas, como moderação de conteúdo ainda exige intervenções humanas em situações que envolvem alta complexidade ou incertezas, por isso o *LinkedIn* recorre até a revisão manual de algumas publicações denunciadas pelos utilizadores ou sinalizadas pelos algoritmos como possíveis fraudes (MOLLY, 2021; MEDIA TODAY, 2023).

Não obstante, o facto de *LinkedIn* ser uma rede social, dificulta mais ainda a capacidade de detecção dos anúncios de emprego fraudulentos, pois sendo uma rede social, qualquer utilizador tem permissões suficientes para publicar o que bem entender gerando uma quantidade de dados ainda maior. Estes factos, aliados também as taxas de desemprego que tendem a subir, principalmente em países em via de desenvolvimento (OIT, 2023), criam um ambiente propício para que os anúncios de vagas falsas triunfem. Portanto, é perceptível que os utilizadores acabam sendo prejudicados caindo em fraudes que chegam a causar danos financeiros e quiçá, outros tipos de danos.

No entanto, faltam ferramentas acessíveis que possam auxiliar na detecção desses anúncios falsos nas plataformas de emprego (Federal Trade Commission, 2023). Daí que, o presente trabalho procura mitigar o problema, criando uma ferramenta que seja acessível para os utilizadores do *LinkedIn* para que consigam mensurar a credibilidade dos anúncios que a plataforma apresenta.

1.3 Motivação

A crescente digitalização de processos de recrutamento, impulsionada pelo uso de plataformas como *LinkedIn*, *Indeed* e *Glassdoor*, transformou o acesso a

oportunidades de emprego, facilitando a conexão entre candidatos e empregadores em escala global. No entanto, esse avanço também trouxe um aumento significativo de fraudes em anúncios de emprego, com entidades mal-intencionadas a publicarem vagas falsas para obter vantagens financeiras e informações pessoais dos candidatos. Esse cenário se agrava com a situação de alta competitividade e a necessidade de emprego, onde muitos profissionais, ansiosos por encontrar uma oportunidade, podem se tornar vítimas de golpes.

Dado o impacto económico e emocional que essas fraudes causam, há uma necessidade urgente de implementar soluções eficazes para mitigar os riscos e criar um ambiente seguro para os candidatos. As plataformas de emprego, embora contem com algumas ferramentas de segurança, ainda enfrentam dificuldades para lidar com o volume de anúncios e a verificação constante de sua autenticidade. Nesse sentido, o desenvolvimento de tecnologias de inteligência artificial e aprendizagem de máquina para identificar possíveis fraudes é fundamental, especialmente em um contexto em que a confiança e a segurança dos utilizadores dessas plataformas estão em jogo.

Este trabalho busca contribuir para esse campo emergente, oferecendo uma análise aprofundada das tecnologias disponíveis e propondo uma solução eficaz para detectar fraudes em anúncios de emprego, com o objetivo de proteger os candidatos e aprimorar a confiabilidade das plataformas online de emprego.

1.4 Objectivos

1.4.1 Objectivo Geral

Desenvolver uma extensão de navegador para análise de credibilidade de anúncios de emprego no *LinkedIn* utilizando aprendizagem de máquina e processamento de linguagem natural (NLP)

1.4.2 Objectivos Específicos

- Descrever o estado actual do *LinkedIn* em relação à ocorrência fraudes em anúncios de emprego e os mecanismos para o seu combate;
- Descrever os conceitos da inteligência artificial e aprendizagem de máquina, destacando sua relevância no contexto de análise de credibilidade em anúncios de emprego;

- Identificar os navegadores mais utilizados no mundo para aceder ao *LinkedIn*, e um repositório de dados aberto sobre ocorrência de fraudes na plataforma;
- Treinar um modelo de aprendizagem de máquina para análise de credibilidade em anúncios de emprego no *LinkedIn* com base nos dados disponibilizados no repositório identificado;
- Desenvolver a extensão para o navegador de maior uso no mundo para aceder ao *LinkedIn* que permita à análise da credibilidade em anúncios de emprego na plataforma.

1.5 Metodologia

Para Richardson (1999, p.22), o método é o caminho ou a maneira para se chegar a um determinado objectivo, e metodologia, por sua vez, pode ser entendida como um conjunto de procedimentos e regras utilizadas por determinado método.

1.5.1 Metodologia de pesquisa

Para Lakatos e Marconi (2017), a pesquisa aplicada é fundamental para encontrar soluções tecnológicas que respondam a demanda da sociedade, utilizando os avanços científicos de maneira prática e inovadora. Tendo isto exposto, pela natureza do trabalho, segue uma abordagem exploratória e aplicada, pois visa investigar um problema emergente, que são fraudes em anúncios de emprego no *LinkedIn* e propor uma solução prática para este desafio através do desenvolvimento de uma extensão de navegador que com recurso a aprendizagem de máquina e processamento de linguagem natural, vai analisar a credibilidade dos anúncios.

1.5.2 Classificação da Metodologia

O trabalho de monografia deve ser sustentado por uma metodologia de pesquisa e para tal existem diversas classificações, sendo necessárias para este trabalho, as abordagens sugeridas por Marconi e Lakatos (2008) e Gil (2003 e 2008), que as classificam: a) quanto a abordagem; b) quanto ao método; c) quanto a natureza; d) quanto aos seus objectivos gerais e por fim, e) quanto aos procedimentos técnicos.

- **Quanto a abordagem**

Quanto a abordagem, as pesquisas científicas podem ser qualitativas ou quantitativas, ou ainda mista, quando agrega os dois tipos. A escolha depende da

área, objecto e dos objectivos da pesquisa. Lakatos e Marconi (2017) afirmam que a pesquisa qualitativa é ideal para explorar questões novas ou pouco compreendidas, fornecendo uma visão mais subjectiva e interpretativa. Já, a pesquisa quantitativa busca medir o fenómeno estudado, utilizando métodos estatísticos para analisar dados numéricos. Ela é útil quase se quer analisar relações causais, comparar variáveis e validar teorias existentes de maneira mais objectiva e replicável (Lakatos e Marconi, 2017).

Portanto, este trabalho segue uma abordagem mista, fazendo uso tanto da pesquisa qualitativa para melhor entender os aspectos subjectivos e as nuances dos anúncios fraudulentos, e da quantitativa para medir a prevalência dos padrões desses anúncios no *LinkedIn*.

- **Quanto à natureza**

A natureza da pesquisa diz respeito a finalidade, à contribuição que ela trará à ciência e pode ser classificada como básica ou aplicada (Even3 blog).

Gil (2008), afirma que a pesquisa básica se destina a produção de conhecimento puro, visando aprofundar a compreensão de um fenómeno específico, se distinguindo pela ausência de compromisso directo com a aplicação prática. Enquanto a pesquisa aplicada é fundamental para encontrar soluções tecnológicas que respondam a demandas da sociedade, utilizando os avanços científicos de maneira prática e inovadora (Lakatos e Marconi, 2017).

Exposto isto, o presente trabalho é classificado como uma pesquisa aplicada pelo facto de propor uma solução prática para os desafios relacionados a credibilidade das vagas do *LinkedIn*, através do desenvolvimento de uma extensão de navegador que com recurso a aprendizagem de máquina e processamento de linguagem natural, vai analisar a credibilidade dos anúncios.

- **Quanto aos objectivos**

Uma pesquisa, quando classificada quanto aos objectivos pode ser exploratória, descritiva ou explicativa (Gil, 2010). Para Vergana (2005), uma pesquisa exploratória visa proporcionar maior familiaridade com o problema, por meio de uma análise detalhada e sistemática de literatura e fontes secundárias. A pesquisa descritiva tem como objectivo descrever um fenómeno ou situação, sem se preocupar com suas causas (Lakatos e Marconi, 2017). Para Gil (2010), a pesquisa explicativa está preocupada em entender as causas de determinado fenómeno e analisar as relações de causas e feitos permitindo explicar como certos eventos ou características

acontecem. Tendo isto exposto, pela natureza do trabalho, segue uma abordagem exploratória, pois visa investigar um problema emergente, que são fraudes em anúncios de emprego no *LinkedIn*, proporcionando maior familiaridade com o problema.

- **Quanto aos procedimentos**

Segundo Gil (2008), uma pesquisa quanto aos procedimentos técnicos pode ser classificada como: pesquisa bibliográfica, documental, experimental, levantamento, estudo de campo e estudo de caso.

1. **Pesquisa bibliográfica**

A pesquisa bibliográfica é aquela que se realiza, segundo Severino (2007), a partir do registo disponível, decorrente de pesquisas anteriores, em documentos impressos, como, livros, artigos, teses, etc. Utilizam-se dados de categorias teóricas já trabalhadas por outros pesquisadores e devidamente registados. Os textos tornam-se fontes dos temas a serem pesquisados. O pesquisador trabalha a partir de contribuições dos autores dos estudos analíticos constantes dos textos.

2. **Pesquisa experimental**

A pesquisa experimental é aquela que envolve algum tipo de experimento, geralmente em laboratórios, onde o pesquisador trabalha em variáveis manipuladas.

Para Gil (1999), esta categoria de pesquisa apoia-se na determinação de um objecto de estudo, na selecção das variáveis capazes de constituir influência ao mesmo e na definição das normas de controlo e de observação dos efeitos que a variável produz no objecto.

3. **Pesquisa documental**

A pesquisa documental, segundo Gil (1999) é muito semelhante à pesquisa bibliográfica, diferenciando-se na natureza das fontes: enquanto a bibliográfica se utiliza fundamentalmente das contribuições de diversos autores, a documental vale-se de materiais que não receberam, ainda, um tratamento analítico, podendo ser reelaboradas conforme os objectivos de pesquisa.

4. **Pesquisa de levantamento**

Para Medeiros (2019), a pesquisa de levantamento é uma categoria de pesquisa que se realiza para a obtenção de dados ou informações sobre características, ou opiniões de um grupo de pessoas selecionado como representante de uma

população. Procede-se à solicitação de informações a um grupo significativo de pessoas acerca do problema estudado para, em seguida, mediante análise quantitativa, obterem-se as conclusões correspondentes aos dados colectados. (Gil, 2008).

5. Pesquisa de estudo de campo

O estudo de campo procura muito mais o aprofundamento das questões propostas do que a distribuição das características da população. Desse modo, seu planeamento torna-se mais complexo e, em simultâneo, sua aplicação é mais flexível do que o levantamento (RUIZ, 2006). Estuda-se um único grupo ou comunidade relativamente a sua estrutura social, ou seja, ressaltando a interação de seus componentes. Utiliza-se com maior frequência técnicas de observação.

6. Pesquisa de estudo de caso

O estudo de caso é caracterizado pelo estudo exaustivo de um ou de poucos objectos, de maneira a permitir o seu conhecimento amplo e detalhado. Para a realização de um estudo de caso podem ser utilizadas diferentes fontes de investigação como: entrevistas, questionários e observação (GIL, 1999). Em linhas gerais, estudo de caso procura o aprofundamento de uma realidade específica.

Para o alcance dos objectivos, o trabalho segue uma pesquisa bibliográfica, documental e de estudo de caso. Bibliográfica porque recorre a revisão das literaturas existentes acerca das fraudes em anúncios de emprego e como o *LinkedIn* lida com estes problemas por exemplos. Documental pelo facto de analisar documentos técnicos sobre a implementação de tecnologias de detecção de fraudes em plataformas *online* de anúncio de emprego. E, por fim, estudo de caso pois tem o seu foco num caso específico que é plataforma do *LinkedIn*.

1.5.3 Técnicas de colecta de dados

Para a realização trabalho foram utilizadas duas técnicas, nomeadamente, a colecta documental, observação, entrevista e questionário.

- **Colecta documental:** Segundo Gil (2008), a colecta documental permite que o pesquisador compreenda melhor o contexto do estudo a partir de materiais previamente colectados e analisados.
- **Observação:** Para pesquisas de campo, a observação directa dos sujeitos ou fenómenos pode ser uma técnica valiosa. De acordo com Gil (2008), a

observação pode ser participativa ou não participativa, dependendo do envolvimento do pesquisador com o grupo estudado.

- **Entrevistas:** Para Cervo e Bervian (2002), a entrevista é uma das principais técnicas de colecta de dados e pode ser definida como conversa face a face pelo pesquisador junto ao entrevistado, seguindo um método para se obter informações sobre determinado assunto.
- **Questionário:** Para Gil (2009), esta técnica é importante especialmente quando se pretende obter uma grande quantidade de dados quantitativos, permitindo uma colecta sistemática a partir de um grande número de pessoas, utilizando questões fechadas ou abertas, conforme o tipo de análise desejada.

Tendo sido expostos os conceitos inerentes as técnicas de colecta de dados, o trabalho seguiu uma abordagem de colecta documental, buscando informações em um *dataset* do *Kaggle* para criação do modelo de aprendizagem.

1.5.4 Metodologia de desenvolvimento do trabalho prático

A metodologia de desenvolvimento utilizada neste trabalho baseia-se no modelo classificado *Waterfall*, amplamente descrito na literatura (Pressman, 2010; Sommerville, 2011). No entanto, Pressman (2010) destaca que a metodologia pode ser adaptada para incluir ciclos iterativos em situações que exige, maior flexibilidade ou revisões intermediárias. Portanto, considerando as características específicas deste projecto académico, o modelo foi adaptado, dando origem ao que será referido a seguir como metodologia ***waterfall adaptada***. Essas adaptações incluem a realização de revisões intermediárias para alinhamento com o supervisor e maior flexibilidade em relação ao cronograma das fases.

Enquanto a metodologia waterfall tradicional segue um fluxo linear, onde cada etapa é concluída antes de iniciar a próxima, a adaptação desta não precisa necessariamente seguir este processo. A seguir são descritas as fases que compreendem esta metodologia:

- **Eliciação de requisitos:** consistiu na colecta e documentação das necessidades do sistema, garantindo que os requisitos sejam claros e compreensíveis.

- **Análise da solução:** caracterizada pela análise dos requisitos do sistema; definição dos requisitos funcionais e não funcionais e modelagem da solução proposta.
- **Projecto ou desenho da solução:** nesta fase foi feito o desenvolvimento do design da solução, incluindo a arquitectura do sistema.
- **Codificação:** implementação do sistema com base no *design* definido através da linguagem de programação e ferramentas seleccionadas.
- **Teste e validação:** foi realizado o teste do modelo de aprendizagem de máquina de acordo com o algoritmo seleccionado.

Paradigma de programação

Existem diversos paradigmas de programação, dependendo do objectivo a ser alcançado. Para este trabalho, se pretende realizar o processamento do texto dos anúncios enviados pelo utilizador do *LinkedIn*, analisá-lo com recurso ao modelo de aprendizagem de máquina e retornar o nível de credibilidade. Desta forma, foi utilizado o paradigma orientado a objectos pelo facto de este ter a capacidade de modularização e reutilização do código, especialmente em sistemas escaláveis e integrados a *APIs* modernas (SOMMERVILLE, 2011).

1.6 Estrutura do trabalho

O trabalho é composto por cinco (5) capítulos, devidamente enumerados e mais uma secção destinada à bibliografia e outra última destinada aos anexos.

• **Capítulo I — Introdução**

Neste capítulo é apresentado o assunto que será tratado ao longo do trabalho, trazendo clareza em relação aos tópicos relevantes, bem como a relevância do mesmo para a sociedade e para o ecossistema de desenvolvimento de sistemas de informação. São partes deste capítulo a contextualização, definição do problema, motivação, objectivos e metodologia.

• **Capítulo II — Revisão de Literatura**

A revisão de literatura, também conhecida como revisão conceitual ou bibliográfica é o capítulo que contém o levantamento bibliográfico preliminar que dará suporte e

fundamentação teórica ao estudo. Nela são citados os principais conceitos relacionados ao trabalho de modo dissertativo, mostrando relação entre os mesmos. E para o presente trabalho é apresentada uma sequência lógica dos conceitos relacionados as fraudes em plataformas *online* com mais enfoque ao *LinkedIn*, bem como embasamento teórico para a solução proposta.

- **Capítulo III – Proposta de solução**

Neste capítulo, ocorre um processo de decisão para definir como e que solução de facto poderá ser implementada. Este processo de decisão é essencialmente baseado nas metodologias definidas que para este trabalho consistem em pesquisa documental e levantamento de formulário. Portanto, é neste capítulo onde ocorre a fusão dos conceitos expostos na revisão de literatura e do formulário com vista a obter uma solução que é viável sob ponto de visto técnico, mas que também vá de encontro com as necessidades dos utilizadores.

- **Capítulo IV – Treinamento do modelo computacional**

A solução a ser desenvolvida é dividida em duas partes: a extensão que é a parte visual e o modelo. Sendo assim, é neste capítulo onde se descrevem os dados que serão utilizados para o treinamento do modelo, bem como a sua proveniência. Se faz, também, uma descrição do processo de preparação dos dados para garantir que o modelo seja eficaz e responda as métricas propostas no mesmo capítulo.

- **Capítulo V — Desenvolvimento da solução proposta**

Neste capítulo, detalha-se a implementação da solução concebida a partir da problemática levantada e fundamentada pelo embasamento teórico discutido nos capítulos anteriores. O objetivo é demonstrar como as diretrizes metodológicas e conceituais foram aplicadas para atender às exigências da questão central do estudo.

- **Capítulo VI – Teste do algoritmo**

Neste capítulo, ocorre uma demonstração passo a passo sobre como é que o algoritmo escolhido para o modelo se comporta, bem como se ele responde às métricas pré-estabelecidas no capítulo IV. Se descreve também quais os desafios identificados e que mecanismos foram aplicados para ultrapassá-los.

- **Capítulo VII — Discussão de Resultados**

Neste capítulo, são discutidos os principais achados do estudo, correlacionando-os com os objetivos previamente estabelecidos. Além disso, analisa-se o impacto potencial da implementação da proposta, considerando suas contribuições e limitações.

- **Capítulo VIII — Considerações finais**

Neste capítulo são apresentadas a síntese de toda a reflexão, as limitações do trabalho e as sugestões para futuras pesquisas. É a parte final do trabalho, onde se apresentam conclusões correspondentes aos objetivos, e se verifica se os objectivos inicialmente levantados foram cumpridos, deixando recomendações para futuras pesquisas.

- **Secção de Bibliografias**

Nesta secção são apresentadas todas obras utilizadas na pesquisa, podendo elas terem sido citadas ou não.

- **Secção de Anexos**

Nesta secção se encontram elementos esclarecedores sobre o sistema e seu processo de formatação, incluindo elementos necessários para melhor compreensão de outras partes do trabalho.

2 Capítulo II - Revisão de Literatura

2.1 Plataformas online de anúncio de emprego

2.1.1 Conceito

Segundo Laudon e Laudon (2020), a tecnologia desempenha um papel fundamental na optimização de processos e na melhoria da produtividade. Na área de recursos humanos, é visível o impacto da tecnologia, pois actualmente existem diversas plataformas *online* que facilitam a publicação de oportunidades de trabalho, bem como a gestão de todo o processo de candidatura e recrutamento (TURBAN et al., 2018). Então, essencialmente, plataformas *online* de anúncio de emprego são ferramentas que permitem que as entidades, quer seja individual ou empresarial, publiquem as oportunidades disponíveis nas suas empresas e permitem também que os candidatos se candidatem directamente pela plataforma ou por via de um outro canal informado.

Actualmente, ao contrário do que acontecia há alguns anos em que as entidades partilhavam as oportunidades de emprego nos jornais, vitrines e outros canais mais tradicionais, as entidades, quer sejam particulares ou institucionais têm feito divulgação maioritariamente através das redes sociais ou em plataformas dedicadas a este fim, e, por sua vez, o candidato às vagas poderá candidatar-se através de um simples email ou submissão em plataformas online, reflectindo a transformação digital que tem estado a revolucionar o mercado de trabalho (MARRA, 2020).

Portanto, apesar de os canais tradicionais ainda serem utilizados para um nicho específico, como setores de alta rotatividade, é notável a relevância das plataformas digitais na sociedade, visto que a maior parte das oportunidades são partilhadas por meio destas.

2.1.2 Principais plataformas *online* de anúncio de emprego

Existem diversas plataformas *online* de anúncio de emprego, mas a nível global há três (3) que têm bastante relevância e é importante destacar. De acordo com Kaplan e Haenlein (2010), plataformas como estas transformaram o mercado de trabalho ao promover conexões rápidas e eficientes. Relatórios de Nikolaou (2014) e Statista

(2024), apontam que o uso massivo e a confiabilidade destas ferramentas consolidam suas posições como líderes no sector de recrutamento *online*.

- **Indeed:** É uma plataforma gigante pois possui uma base de dados extensa de vagas a nível global. Na plataforma é possível encontrar oportunidades para diversas indústrias e oferece funcionalidades extremamente avançadas para filtro e triagem de forma a facilitar a busca tanto para recrutadores quando para candidatos (KUMAR, 2021).
- **Glassdoor:** É uma plataforma sobejamente conhecida por um dos seus principais recursos que consiste em os utilizadores poderem fornecer uma avaliação anónima de empresas, fornecendo aos candidatos informações relevantes sobre a cultura organizacional, salários e processos selectivos. Além disso, a plataforma divulga vagas de emprego, permitindo candidaturas directas (POTTER, 2019, p. 52).
- **LinkedIn:** É a maior plataforma profissional do mundo; é na verdade uma rede social profissional com mais de 850 milhões de utilizadores (LINKEDIN). Ela permite que os candidatos se conectem com empregadores; possui funcionalidades de candidatura directa para vagas, ou seja, a empresa partilha uma vaga de emprego e o utilizador tem a capacidade de se candidatar directamente do *LinkedIn* sem a necessidade de aceder ao *site* oficial. E por fim, possui ferramentas como cursos e recomendações para melhorar perfis profissionais (LINKEDIN). Além disso, o *LinkedIn* possui uma taxa de 40.29% de candidaturas para cada anúncio publicado, destacando a sua eficácia na atracção de candidatos.

2.1.3 Conceito de fraude em anúncios de emprego

A busca pelo emprego pode ser uma jornada desafiadora e cheia de incertezas, e onde há alta demanda, há sempre também oportunistas querendo se aproveitar da situação. Os falsos anúncios de emprego têm como objectivos a exploração financeira, roubo de identidade, entre outros. A mensagem subjacente em maior parte destes anúncios vende uma promessa de salários elevados, modelo de trabalho com

maior flexibilidade e ambientes de trabalho estáveis de forma atrair maior número de candidatos. Segundo Frank Abagnale (2019), os autores destes actos procuram explorar as vulnerabilidades humanas através da elaboração de ofertas de emprego convincentes que induzem o candidato a partilhar informações pessoais ou realizar pagamentos.

Após os candidatos manifestarem algum interesse pela oportunidade falsa, inicia o processo de exploração que pode ser através da cobrança de taxas para participação no processo de candidatura, solicitação de informações pessoais cobrança para realização de falsas formações, etc.

2.1.4 História das fraudes em anúncios de emprego

A prática de enganar candidatos através de falsos anúncios de emprego, apesar de ser mais frequente actualmente, não é tão contemporânea pois teve o seu início ainda na época antes da internet e foi evoluindo consoante o tempo.

- **Época pré-internet:** apesar da inexistência de meios tão velozes como actualmente, nesta época algumas entidades faziam anúncios de emprego falsos em meios de comunicação impressos, como jornais, onde maior parte das vezes cobravam pelos formulários de candidatura e até para realização de entrevistas.
- **Início da internet (entre 1990 e 2000):** com o início da *internet*, a propagação de falsos anúncios de emprego tornou-se cada vez mais predominante, pois começou o surgimento de painéis online criando um espaço em que os protagonistas deste fenómeno tivessem maior alcance. Mitnick & Simon (2002), debruçam-se sobre este assunto afirmando que os agentes de falsos anúncios ganharam novas oportunidades com a ascensão da internet pois facilitou a partilha de informações confidenciais através de manipulação dos candidatos.
- **Redes sociais e plataformas de emprego (década de 2010):** À medida em que plataformas conhecidas como fidedignas como *LinkedIn*, *Indeed* e outras foram crescendo, os autores começaram a explorá-las. Maior parte se faz passar por recrutador ou empresa de forma a criar publicações falsas, mas que

pareçam genuínas. Hadnagy (2018) esclarece como os enganadores utilizam plataformas como o *LinkedIn* para construir a confiança dos candidatos a emprego, tornando mais difícil a identificação de ofertas de emprego falsas.

- **Pandemia da COVID-19:** A mudança global para o trabalho remoto e a incerteza económica durante a pandemia de COVID-19 levaram a um aumento de anúncios de emprego falsos. Muitos dos autores destes falsos empregos atacaram candidatos vulneráveis, oferecendo empregos remotos fraudulentos.

2.1.5 Mecanismos de detecção de fraudes usados pelo LinkedIn

O *LinkedIn* procura implementar mecanismos de segurança para detectar falsos anúncios de emprego. Estes mecanismos muitas vezes consistem na identificação automática de elementos comuns em falsos anúncios, como pedido antecipado de pagamentos, promessa de salários extremamente altos, descrições incompletas, entre outros (GAO; LIU; ZHAO, 2020). De acordo com a *Cybersecurity Journal* LinkedIn (2023), em geral, estes são os mecanismos utilizados:

- **Inteligência artificial (IA) e aprendizagem de máquina:** O *LinkedIn* faz uso de IA baseada em algoritmos para filtrar publicações e anúncios fraudulentos. Basicamente, estes algoritmos fazem análise de padrões comuns em anúncios de emprego falsos, como solicitação de pagamentos antecipados, promessas de salários irrealistas, descrições vagas ou incompletas e inconsistência na informação de contactos
- **Denúncia de anúncios e publicações:** Uma estratégia de combate às fraudes é a denúncia de anúncios suspeitos de fraude. O *LinkedIn* incentiva seus utilizadores a denunciarem os anúncios caso tenham visto algo que lhes remeta a fraude e combinação disto com as automatizações ajudam a reduzir o número de anúncios fraudulentos na plataforma.
- **Moderação humana e revisão manual:** Apesar do investimento em automatizações, o *LinkedIn* combina a IA com uma equipa de moderadores humanos para fazerem a revisão dos anúncios sinalizados como suspeitos através dos algoritmos ou com base nas denúncias feitas. Este mecanismo é comum porque em grandes plataformas há maior complexidade em identificar

fraudes mais sofisticadas, necessitando as vezes de um profissional dedicado a isso.

2.1.6 Constrangimento nos mecanismos actuais de detecção de fraudes em anúncios de emprego do *LinkedIn*

As plataformas de anúncio de emprego têm evoluído e buscado constantemente combater as fraudes em anúncios de emprego, porém existem diversos desafios neste processo.

Um problema comum nos sistemas automatizados é a ocorrência de falsos positivos, onde anúncios legítimos são marcados como fraudulentos, e falsos negativos, onde anúncios falsos passam despercebidos. Segundo CAMPOS (2020), sistemas baseados em regras simples são particularmente susceptíveis a esses problemas, pois não conseguem capturar nuances complexas.

Um dos maiores constrangimentos das estratégias actuais de combate às fraudes, é a rápida adaptação dos autores de fraudes para escapar a detecção pelos sistemas, o que exige que haja uma actualização constante dos algoritmos e das regras de detecção. Segundo Pereira e Silva (2021), os autores de fraudes ajustam o vocabulário e o estilo dos anúncios para enganar os sistemas, o que torna a eficácia dos mecanismos automatizados temporária.

Quando se fala de plataformas como *LinkedIn* estamos a falar de conjunto gigantesco de dados e apesar de os algoritmos ajudarem na remoção de milhões de anúncios que não atendam aos padrões estabelecidos, outros milhões de anúncios continuam na plataforma, pois, como foi dito anteriormente, também existe uma dependência de revisão humana das publicações. A revisão manual pode ser eficaz, mas além de envolver custos elevados, possui limitações no que diz respeito a escalabilidade. Plataformas que contam com equipas de revisão manual enfrentam dificuldades para lidar com grande volume de anúncios. Isso pode atrasar a remoção de conteúdos fraudulentos e prejudicar a experiência do utilizador (GOMES, 2019).

2.2 Principais navegadores utilizados para aceder ao LinkedIn

Primeiro, o que seriam navegadores *web*? De acordo Tanenbaum (2003), o navegador é um programa responsável por buscar e exibir os recursos da *web*, através da interpretação de documentos *HTML* e uma interface que permite a interação com o conteúdo do *site*, como textos, imagens e multimédia. Webb (2017) completa afirmando que os navegadores modernos incorporam funcionalidades como controlo de segurança, funcionalidade de privacidade, extensões e ferramentas para desenvolvedores.

Apesar de existirem algumas fontes que fazem menção do facto do *LinkedIn* ser mais acedido pelo navegador *Google Chrome*, o *LinkedIn* não tem nenhuma informação oficial a respeito disso, mas segundo a WORLDMETRICS (2024), a nível global, o *Google Chrome* lidera o nível de utilização com uma participação dominante de 68.78% entre os utilizadores *desktop*, tornando-o no navegador ideal para o presente trabalho.

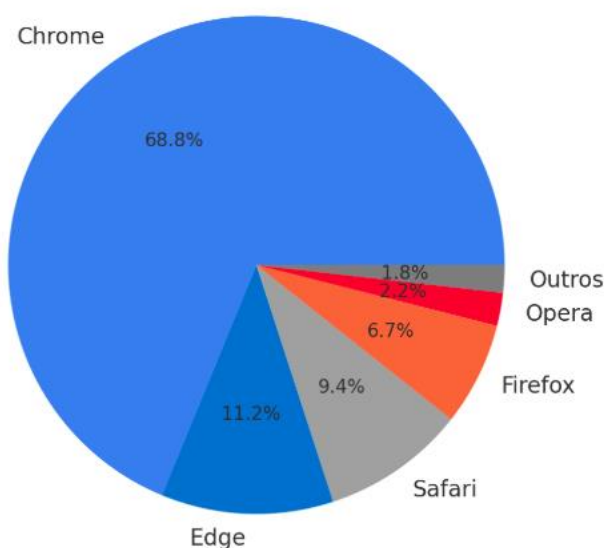


Figura 1: Participação de mercado dos navegadores para desktop (2024). Fonte: Adaptado de StatCounter Global Stats, disponível em: gs.statcounter.com

2.3 Extensão de navegadores e seu desenvolvimento para Google Chrome

Segundo Chakrabarti e Singh (2020), extensões de navegadores são módulos de software que adicionam funcionalidades específicas a um navegador web permitindo

aos utilizadores realizarem personalizações na sua experiência de navegação e executarem tarefas específicas no navegador. A *Google Developers* (2023) acrescenta que uma extensão de navegador é um software que melhora as funcionalidades existentes, permitindo interações e melhorias directamente na interface gráfica do utilizador, geralmente através de ícones ou opções integradas na interface do navegador.

No que diz respeito ao desenvolvimento das extensões no *Google Chrome*, o desenvolvimento de extensões é realizado por meio de *APIs* específicas e tecnologias comuns da *web*, como o *HTML*, *CSS* e *Javascript*, o que torna mais acessível a desenvolvedores de *software* de qualquer nível pois são ferramentas abertas ao público e de fácil aprendizagem.

2.4 Aprendizagem de máquina (*Machine Learning* - *ML*)

O conceito de aprendizagem de máquina é um dos temas centrais deste trabalho, porém, antes de debruçar-se acerca deste, é importante referir que a aprendizagem de máquina é uma sub-área da **inteligência artificial (IA)**.

2.4.1 Inteligência artificial (IA)

De acordo com Russel e Norvig (2010, p. 34) pode ser definida como o estudo de agentes que recebem percepções do ambiente e realizam acções. Uma outra definição relevante seria a de Minsky (1968) que diz que a inteligência artificial é a ciência que torna as máquinas capazes de realizar tarefas que necessitariam de inteligência se fossem executadas por seres humanos. Portanto, a inteligência artificial faz com que a partir de entradas do ambiente, os sistemas tornem-se capazes de tomar decisões e executar tarefas de forma inteligente.

A inteligência artificial é fundamental, promovendo inovações e revolucionando vários sectores industriais. Seu valor pode ser ressaltado em diversos aspectos distintos como automação de processos, tomada de decisão baseada em dados, personalização na experiência do utilizador, saúde, entre outros.

2.4.2 Conceito de aprendizagem de máquina

- **Definição**

De acordo com Russel e Norving (2003), a aprendizagem de máquina é o campo de estudo que dá aos computadores a capacidade de aprender sem serem explicitamente programados. Diferente dos sistemas convencionais que são programados para executar actividades específicas, os algoritmos de ML conseguem aprimorar sua eficácia ao serem expostos a novas informações, realizando previsões e decisões com base em padrões detectados nos dados. Essecialmente, a aprendizagem de máquina possibilita que máquinas obtenham conhecimento sem precisarem ser programadas explicitamente para executar determinadas acções.

- **Relação entre aprendizagem de máquina e inteligência artificial**

Como mencionado anteriormente, a aprendizagem de máquina é uma subárea da inteligência artificial, sendo amplamente reconhecido como um dos seus principais pilares.

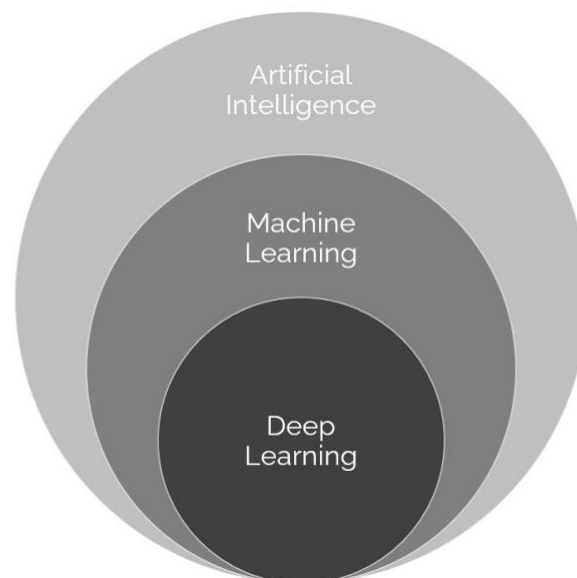


Figura 2: Aprendizagem de máquina como sub-área da inteligência artificial. Fonte: CSLEE Tech blog, 2022

Enquanto a inteligência artificial é uma área abrangente que busca desenvolver sistemas com a capacidade de replicar funções cognitivas humanas, como raciocínio, decisão e identificação de padrões, a aprendizagem de máquina como técnica fundamental, permite que os sistemas aprendam independentemente, sem precisar

de instruções codificadas para cada tarefa. Dessa forma, embora a IA seja um conceito mais amplo, o ML foca especificamente em como os sistemas podem aprender e aprimorar suas funções de forma autónoma.

2.4.3 Tipos de aprendizagem de máquina

A aprendizagem de máquina pode ser categorizada em diversos tipos, dependendo de como os algoritmos adquirem conhecimento a partir dos dados. Cada tipo de ML apresenta diferentes maneiras de resolver problemas, de acordo com a disponibilidade e características dos dados, e também conforme os objectivos da aplicação. A escolha adequada do tipo de aprendizagem é fundamental para o êxito dos modelos ML, uma vez que afeta directamente o desempenho, a precisão e a generalização das soluções criadas. Geralmente, os tipos de aprendizagem de máquina podem ser classificados em três categorias principais: aprendizagem, não aprendizagem e por esforço, os quais serão abordados a seguir.



Figura 3: Principais tipos de aprendizagem de máquina (Fonte: <https://beatrizmaiads.medium.com>)

1. Aprendizagem supervisionada

Segundo Witten et al. (2016), a aprendizagem supervisionada é um tipo de ML que consiste no treinamento de um modelo através de um conjunto de dados previamente rotulados contendo pares de entradas e saídas conhecidas. Através deste treinamento, o modelo deverá adquirir capacidade de fazer previsões sobre novos

dados. É importante mencionar que este processo envolve ajuste dos parâmetros internos do modelo para reduzir a margem de erros nas previsões.

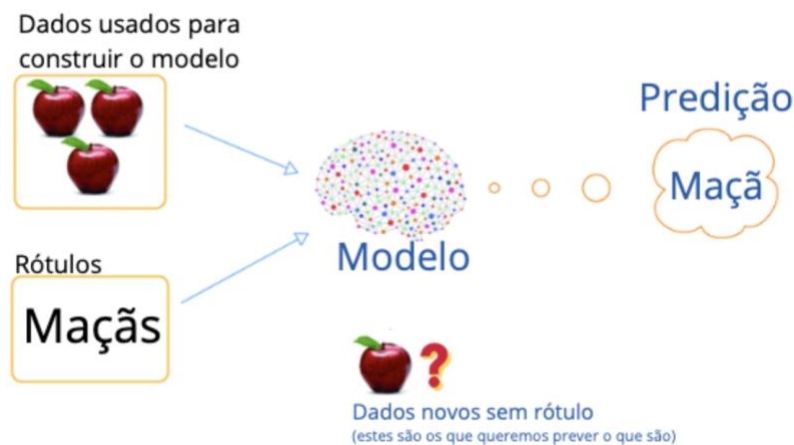


Figura 4: Aprendizagem supervisionada
(Fonte: Estatsite)

Na figura acima, são fornecidos dados rotulados como maçã na expectativa de treinar o modelo para ter capacidade de fazer previsões quando novos dados que não tenham sido rotulados forem apresentados.

Principais tarefas da aprendizagem supervisionada

A aprendizagem supervisionada destina-se a resolução de diversos tipos de tarefas com base em dados rotulados. Para cada tarefa, existe um objectivo específico e diferentes abordagens para o treinamento e avaliação do modelo. Apesar de existirem outras tarefas como detecção de anomalias, reconhecimento de padrões, as principais tarefas da aprendizagem supervisionada são a **classificação** e a **regressão**.

A. Classificação

Para Mitchell (1997), classificação é o processo de prever um rótulo categórico para uma instância com base em atributos previamente observado, destacando que a aprendizagem ocorre quando o modelo melhora sua capacidade de associar instâncias a classes correctas com base em dados rotulados de treinamento. Essencialmente, a classificação consiste na categorização de uma nova observação com base em um conjunto de dados rotulados. A classificação pode ser denominada **binária** ou **multiclasse**. A classificação é conhecida por **binária** quando o objectivo

é atribuir uma entrada de dados a apenas duas classes, por exemplo: ao determinar se uma lâmpada está ou não, as classes seriam “acesa” ou “apagada”. Por outro lado, a classificação **multiclasse** envolve mais de duas (2) classes, por exemplo: classificação de gêneros musicais que pode ser “hip-hop”, “jazz”, “marrabenta”, etc.

B. Regressão

Para Russel e Norving (2016) a regressão é um problema de aprendizagem de máquina, aprendizagem em que a tarefa é prever um valor contínuo com base em um conjunto de variáveis de entrada. O problema de regressão consiste na previsão de valores numéricos a partir de dados de entrada. Ao contrário da classificação, que atribui classes específicas, a regressão trabalha com valores contínuos. Quando se diz que a saída da regressão é um valor contínuo, significa que pode assumir qualquer valor em um intervalo. Por exemplo, prever o preço de uma casa, a temperatura do dia ou o tempo que uma pessoa levará para percorrer uma determinada distância.

2. Aprendizagem não supervisionada

Para Mitchell (1997), aprendizagem não supervisionada é uma abordagem onde o modelo aprende padrões e estruturas nos dados sem a necessidade de rótulos pré-definidos. Diferente do que acontece na aprendizagem supervisionada, o objectivo desta é explorar a estrutura subjacente dos dados, identificando grupos ou padrões relevantes sem supervisão explícita. Por exemplo, em *marketing*, são usadas técnicas de agrupamentos de clientes com comportamentos semelhantes, permitindo que as empresas criem estratégias direccionadas a grupos específicos sem a necessidade de rótulos definidos previamente.

3. Aprendizagem por reforço

Para Sutton e Barto (2018), aprendizagem por reforço é uma área da aprendizagem de máquina onde um agente toma decisões sequenciais em um ambiente, buscando maximizar uma recompensa acumulada ao longo do tempo. O agente aprende por tentativa e erro, ajustando suas acções com base no *feedback* recebido (recompensas ou punições). Por exemplo, um robô de limpeza que aprende a mapear

um ambiente e realizar tarefas de limpeza mais eficientes com base nas suas recompensas.

2.4.4 Algoritmos de aprendizagem de máquina

Para resolver problemas de classificação, existem diversos algoritmos que podem ser aplicados, cada um com suas particularidades e casos de uso. A seguir, poderemos explorar alguns dos algoritmos mais comuns:

1. Árvores de decisão

Segundo Mitchel (1997), árvores de decisão são uma estrutura em que cada nó interno corresponde a um teste de um atributo, cada ramo a um resultado do teste e cada nó folha mantém um valor de classe. É uma representação hierárquica que orienta a classificação de instâncias, onde decisões sucessivas levam à classificação final. Uma árvore de decisão possui a seguinte estrutura típica: (1) nós internos não rotulados com atributos; (2) folhas são rotuladas com classes e (3) ramos são rotulados com valores (atributos categóricos) ou com intervalos (atributos numéricos).

2. Máquinas de vectores de suporte (SVM)

As SVM são classificadoras lineares que buscam encontrar um hiperplano de separação ótimo entre as classes. Hiperplano, em um espaço bidimensional, é uma linha que separa os dados em duas classes e em espaços de dimensões superiores torna-se um termo generalizado. Esse hiperplano é escolhido para maximizar a margem, ou seja, a distância entre os pontos mais próximos de ambas classes, chamados de vectores de suporte. Em termos gerais, as SVM procuram o hiperplano que melhor separa as classes de modo a minimizar o erro de classificação (CORTES; VAPNIK, 1995; BURGESS, 1998).

Essa abordagem é largamente utilizada em problemas de classificação pois tem uma boa capacidade de lidar com dados de alta dimensionalidade e é robusta contra overfitting, especialmente em casos onde a relação entre as classes é linearmente separável (CRISTIANINI, 2000).

3. Regressão logística

Segundo Hosmer e Lemeshow (2000), a regressão logística é uma técnica estatística usada para modelar a probabilidade de ocorrência de um evento em função de uma

ou mais variáveis independentes. Hosmer e Lemesho acrescentam que a regressão logística utiliza uma transformação sigmoide para modelar a probabilidade de uma variável dependente binária. De acordo com o Kleinbaum e Klein (2010), a regressão logística é ideal para estudar a relação entre uma variável dependente categórica e uma ou mais variáveis independentes, sejam estas categóricas ou contínuas.

4. Redes neurais

Hinton (2012), define redes neurais como sistemas computacionais que tentam aprender a partir de dados por meio de uma rede de unidades conectadas, de forma semelhante aos neurónios do cérebro humano. O autor destaca o uso de *backpropagation* para ajustar os pesos das conexões entre neurónios, a fim de minimizar os erros de previsão e melhorar a capacidade de generalização do modelo.

5. Naive Bayes

Russel e Norving (2009), definem Naive Bayes com um algoritmo de aprendizagem baseado no teorema de Bayes, que assume que as variáveis preditoras são condicionalmente independentes entre si, dado a classe. Apesar da simplicidade e da suposição ingénua de independência, ele é amplamente utilizado e pode ser altamente eficiente. Essencialmente, este algoritmo usa o teorema de Bayes para calcular a probabilidade de uma classe, dada as características (*features*) de entrada.

$$P(A | B) = \frac{P(B | A) \cdot P(A)}{P(B)}$$

A, B	= eventos
$P(A B)$	= probabilidade de A se B for verdade
$P(B A)$	= probabilidade de B se A for verdade
$P(A), P(B)$	= as probabilidades independentes de A e B

Figura 5: Teorema de Bayes

2.4.5 Processamento de linguagem natural (NLP)

Processamento de linguagem natural é uma área linguística computacional focada na interação entre computadores e linguagem humana, especialmente na capacidade de

entender, interpretar e gerar linguagem humana de maneira que seja útil (Jurafsky e Martin, 2021). Os autores enfatizam que NLP utiliza técnicas de aprendizagem supervisionada e não supervisionada para analisar grandes volumes de dados textuais e extrair significados complexos, como contextos e relações entre palavras. Essencialmente, o processamento de linguagem natural é uma área fundamental para o desenvolvimento de tecnologias que permitem aos computadores compreender e gerar texto de maneira humana.

2.4.6 Relevância da inteligência artificial e aprendizagem de máquina na análise de credibilidade de anúncios de emprego

A crescente proliferação de anúncios de emprego fraudulentos em plataformas digitais tem gerado desafios significativos para candidatos e empregadores. Com o aumento da quantidade de informações disponíveis online, identificar anúncios falsos tornou-se uma tarefa complexa, demandando soluções mais avançadas e automatizadas. Neste contexto, a inteligência artificial e aprendizagem de máquina surgem como ferramentas poderosas para análise de credibilidade desses anúncios.

A IA e o ML permitem a análise de grandes volumes de dados em tempo real, possibilitando a identificação de padrões e comportamentos fraudulentos que seriam difíceis de detectar manualmente. Segundo Xia et al. (2018), técnicas de processamento de linguagem natural (NLP) são especialmente eficazes para entender e classificar conteúdo textual dos anúncios, identificando inconsistências, palavras-chave suspeitas e até mesmo a linguagem utilizada. Essas técnicas podem ser aplicadas para analisar a estrutura do anúncio, verificar a coerência das informações apresentadas, além de cruzar dados com históricos de anúncios semelhantes para prever a veracidade de um novo anúncio.

A utilização de algoritmos de aprendizagem supervisionada, como a regressão logística, máquinas de vetores de suporte (SVM) e redes neurais artificiais, têm mostrado particularmente eficiente na tarefa de classificação de anúncios de emprego, distinguindo entre os legítimos e fraudulentos com base em atributos específicos do anúncio, como o título, a descrição e a localização da vaga (Liu, et al., 2019). Além disso, métodos de aprendizagem não supervisionada também têm sido

aplicados que permitem agrupar anúncios com características semelhantes e detectar padrões anómalos associados a fraudes.

Além disso, Raghu et al (2020), destacam que a utilização de IA e ML não se limita à análise textual, mas também se estende à análise de dados históricos e comportamentais, como o tempo de postagem do anúncio, a frequência de atualizações e o comportamento dos utilizadores interagindo com o anúncio. Essas variáveis oferecem uma visão mais holística da credibilidade do anúncio, ajudando a identificar tentativas de manipulação ou práticas enganosas que possam passar despercebidas por uma análise puramente textual.

Portanto, a aplicação dessas tecnologias não só melhora a precisão na detecção de anúncios fraudulentos, mas também aumenta a confiança dos utilizadores nas plataformas de emprego, garantindo uma experiência mais segura e eficiente. A capacidade de automatizar e escalar essa análise torna a AI e o ML indispensáveis para enfrentar os desafios actuais da fraude em anúncios de emprego.

2.4.7 Fontes de dados para análise de credibilidade em anúncios de emprego

A eficiência de um modelo de Machine Learning (ML) na detecção de fraudes em anúncios de emprego depende directamente da qualidade e relevância dos dados utilizados para o seu treinamento e avaliação. Dessa forma, a selecção de uma fonte confiável de dados é um aspecto fundamental para garantir que a solução desenvolvida possa identificar padrões consistentes de fraude.

Repositórios de dados para detecção de fraudes

Diversos repositórios de dados abertos oferecem conjuntos de dados relacionados a fraudes online, incluindo fraudes em anúncios de emprego. Esses *datasets* são essenciais para o desenvolvimento de modelos preditivos, pois contêm registos históricos de anúncios fraudulentos e legítimos, permitindo a identificação de padrões específicos associados a cada categoria.

Entre os principais repositórios de dados utilizados para estudos relacionados a fraudes em plataformas digitais, destacam-se:

- **Kaggle:** um dos maiores repositórios de *datasets* públicos, contendo diversos conjuntos de dados sobre fraudes, incluindo aqueles voltados para anúncios de emprego (Chowdhury et al., 2022)

- **UCI Machine Learning Repository:** uma base de dados amplamente utilizada na pesquisa académica, oferecendo conjuntos de dados relacionados a fraudes online e análise de credibilidade de informações (Dua & Graff, 2019).
- **OpenML:** Uma plataforma colaborativa que disponibiliza *datasets* para aprendizado de máquina, incluindo alguns relacionados à detecção de fraudes.



Figura 6: Repositório de dados (Fonte: Super Annotate).

2.4.8 Ferramentas para implementação de aprendizagem de máquina

Aqui são enumeradas as principais ferramentas que serão utilizadas ao longo do desenvolvimento do projecto, principalmete no que diz respeito ao modelo de aprendizagem de máquina:

- **Python:** Para Allen Downey (2002), Python é uma linguagem ideal para iniciantes em programação devido à sua sintaxe intuitiva, mas suficientemente poderosa para aplicações complexas como ciência de dados e inteligência artificial.
- **Pandas:** É a biblioteca para manipulação e análise de dados em Python, oferecendo uma interface fácil de usar e funcionalidades poderosas para operações de transformação, agregação e análise de dados (Daniel Chen, 2013).
- **Jupyter Notebook:** Jupyter Notebook, de acordo com a documentação oficial da ferramenta, é uma aplicação *web* que permite a criação e partilha de documentos que contém código executável, visualizações e texto explicativo. Ele é amplamente usado em ciência de dados, aprendizagem de máquina, pesquisas científicas e educação.

- ***Scikit-learn***: Segundo os criadores, é uma biblioteca de aprendizagem de máquina que oferece uma interface simples e consistente para construção de modelos preditivos, incluindo classificação, regressão, *clustering* e redução de dimensionalidade (Pedregosa et al, 2011).
- ***NLTK (Natural Language Toolkit)***: Raschka (2015), afirma que é uma das bibliotecas mais populares para o processamento de linguagem natural, fornecendo uma ampla variedade de recursos, incluindo modelos linguísticos e ferramentas para análise semântica e sintática.

3 Capítulo III – Proposta de solução

Conforme discutido nos capítulos anteriores, a presença de falsos anúncios de emprego pode gerar prejuízos para candidatos e empresas, tornando essencial a implementação de mecanismos que avaliem a veracidade dessas publicações.

Neste contexto, este capítulo apresenta a proposta de uma solução baseada em inteligência artificial e aprendizado de máquina para análise de credibilidade dos anúncios de emprego no *LinkedIn*.

Com vista minimizar vieses e aumentar a confiabilidade, a abordagem adotada para a proposta da solução é a triangulação metodológica que consiste na utilização de múltiplas fontes de informação ou métodos de pesquisa (Denzin, 1978). Sendo assim, a proposta de solução é fundamentada tanto na revisão da literatura quanto em dados colectados por meio de um formulário aplicado ao público-alvo, permitindo identificar as melhores decisões para implementação da solução.

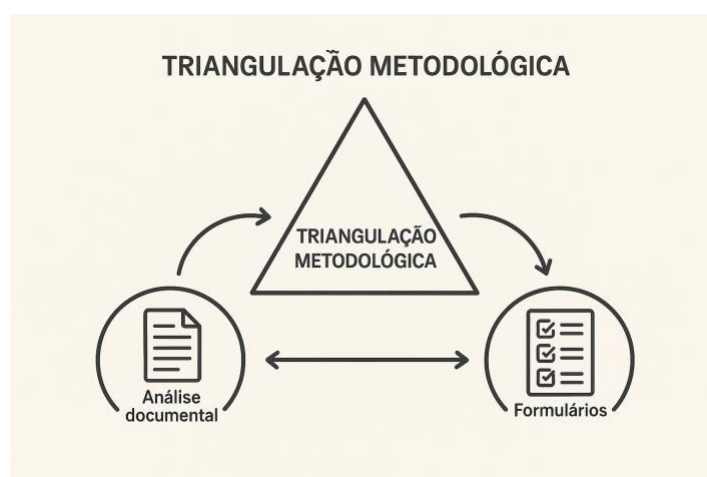


Figura 7: Triangulação metodológica (fonte: chatgpt)

3.1 Relevância da busca por uma solução

O crescimento das plataformas de recrutamento *online* trouxe consigo desafios significativos relacionados à credibilidade dos anúncios de emprego, especialmente no *LinkedIn*. Conforme abordado na revisão da literatura, a plataforma conta com aproximadamente 850 milhões de utilizadores, consolidando-se como uma das redes sociais profissionais mais utilizadas globalmente.

Além das evidências levantadas na revisão de literatura, a pesquisa realizada reforça essa preocupação. Conforme ilustram os anexos, dos participantes do formulário, 93.9% eram públicos de Moçambique e 87.9% afirmaram já ter utilizado o *LinkedIn* para procurar emprego. Esses dados evidenciam a relevância do problema e a

necessidade de uma solução eficaz para garantir maior segurança e confiabilidade nas oportunidades divulgadas na plataforma.

3.2 Escolha do navegador para implementação da extensão

A escolha do navegador para a implementação da extensão é um factor crucial para garantir a adopção e eficácia da solução proposta. Por isso, há uma necessidade fazer uma análise para ir de encontro com as expectativas dos utilizadores.

3.2.1 Análise de utilização de navegadores

Conforme apontado na revisão da literatura, diversos estudos apontam que a maioria dos acessos a plataformas *online* ocorre por meio de navegadores populares como *Google Chrome*, *Microsoft*, *Microsoft Edge*, *Mozilla Firefox* e *Safari*; E segundo StatCounter (2024), o **Google Chrome** lidera com cerca de 68.78% do mercado global de navegadores.

Além disso, um estudo realizado por McDonald et al (2021) sobre acessibilidade a redes sociais reforça que a escolha do navegador pode impactar a usabilidade e a adoção de extensões para análise de credibilidade, visto que alguns navegadores possuem suporte mais amplo para *WebExtensions*.

Para validar a literatura com dados empíricos, foi aplicado um formulário ao público-alvo de pesquisa. Conforme ilustrado nos anexos, os resultados indicaram que a maioria dos respondentes utiliza o Google Chrome para aceder o *LinkedIn*, seguindo o *Mozilla Firefox*.

Qual é o navegador que você usa com mais frequência para procurar vagas de emprego no LinkedIn?

33 responses

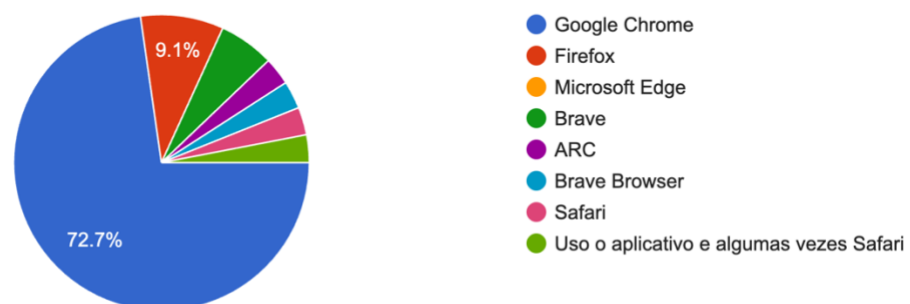


Figura 8: Navegador utilizado para aceder ao *LinkedIn*

Portanto, conforme se pode observar, estes dados reforçam a relevância de desenvolver a extensão para **Google Chrome** inicialmente, pois atinge maior parte dos utilizadores e facilita a adopção da ferramenta sem exigir mudanças de hábitos.

3.3 Funcionamento da extensão de navegador

Ao propor uma solução para o problema da presença de anúncios de emprego fraudulentos no *LinkedIn*, é necessário considerar não apenas a viabilidade técnica, mas também a experiência desejada pelo utilizador e as limitações das abordagens existentes. Assim, este subtítulo visa analisar possíveis formas de funcionamento de um sistema de verificação de credibilidade, confrontando preferências do público-alvo com evidências científicas, de modo a justificar o funcionamento mais adequado para a extensão a ser desenvolvida.

3.3.1 Preferência do público-alvo

De acordo com o formulário aplicado, a maioria dos participantes (54,5%) preferia um sistema que evitasse exibição de anúncios falsos, utilizando filtros automáticos disponibilizados pelo *LinkedIn*. Outros 30,3% sugeriram a existência de um mecanismo interno do próprio *LinkedIn* que pudesse verificar a autenticidade em tempo real. Apenas 15,2% dos respondentes demonstraram preferência por um sistema independente, como uma extensão de navegador, que oferecesse maior autonomia ao utilizador.

De que forma acha que a verificação de credibilidade de anúncios de emprego no LinkedIn deve funcionar?

33 responses

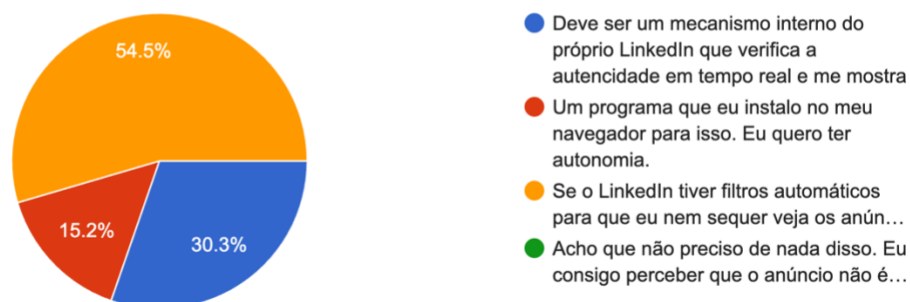


Figura 9: Forma escolhida para verificação de credibilidade em anúncios.

3.3.2 Limitações das abordagens nativas do LinkedIn

A literatura especializada em segurança digital e detecção de fraudes (Phua et al., 2012; Kumar et al., 2016; Sillito et al., 2022) aponta para a limitação dos filtros internos implementados por grandes plataformas. Estes mecanismos, embora úteis, raramente são suficientes para impedir totalmente a circulação de conteúdos enganosos. Isso se dá, sobre tudo, pela rápida evolução das técnicas utilizadas por entidades fraudulentas e pela natureza fechada (não auditável) dos sistemas internos das plataformas.

Além disso, os próprios termos de uso e políticas de API do LinkedIn não oferecem suporte directo para desenvolvedores auditarem o conteúdo dos anúncios via meios oficiais, o que reforça a necessidade de solução externas.

Tabela 1: Comparação entre abordagens possíveis

Critério	Filtros internos do LinkedIn	Solução automática integrada	Extensão de navegador proposta
Transparência da avaliação	Baixa	Baixa	Alta
Controlo e autonomia do utilizador	Nenhum	Baixo	Alto
Personalização e adaptação contínua	Baixa	Média	Alta
Facilidade de actualização da solução	Baixa	Baixa	Alta
Abertura para auditoria e melhoria	Nenhuma	Baixa	Alta
Preferência dos utilizadores	Alta	Alta	Média
Viabilidade técnica e legal	Restrita	Restrita	Viável
Baseada em IA/ML	Sim	Possível	Sim

3.3.3 Fluxo de funcionamento da extensão

Considerando os factores expostos anteriormente, propõe-se que a extensão opere da seguinte forma:

1. Utilizador selecciona o texto a ser analisado.
2. A extensão é informada para extrair o texto seleccionado.
3. O texto é enviado para um *backend* que contém o modelo de aprendizado de máquina treinado para classificar a credibilidade do anúncio.
4. O *backend* retorna uma pontuação ou categoria de credibilidade.
5. A extensão exibe a pontuação de credibilidade ao utilizador de forma visual

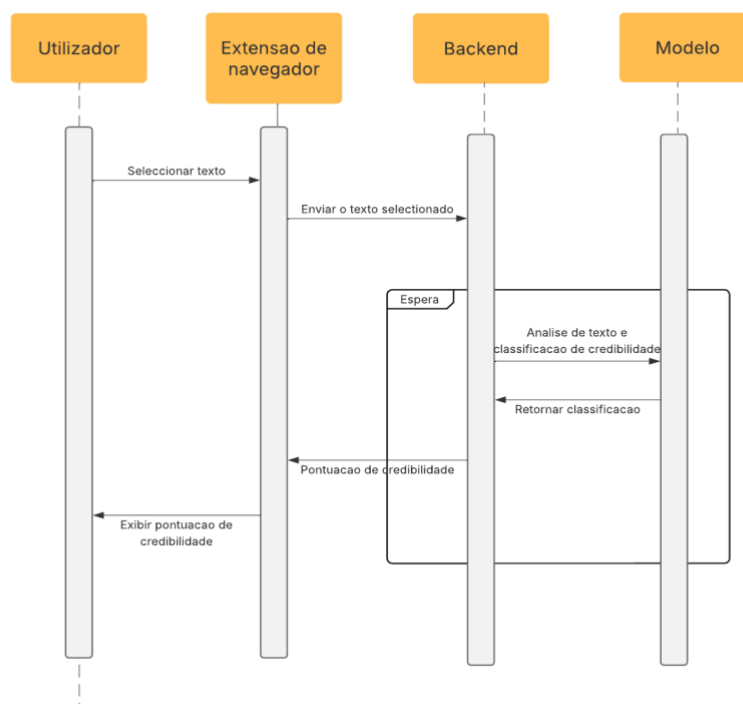


Figura 10: Diagrama de sequência da solução proposta.

3.4 Repositório de dados e *Dataset*

Para treino do modelo de aprendizagem, optou-se pela utilização de dados rotulados disponíveis no **Kaggle**, uma das plataformas mais reconhecidas no domínio da ciência de dados e inteligência artificial. O Kaggle, oferece *datasets* com curadoria comunitária, acessíveis, bem documentados e amplamente utilizados na investigação académica.

O *dataset* selecionado foi o *Fake Job Posting Prediction* criado por Shivam Bansal, contendo cerca de 18000 anúncios de emprego rotulados como reais ou falsos. Este *dataset* servirá de base para treinar e testar o modelo que será integrado no *backend* da extensão proposta.

3.5 Algoritmo de aprendizagem de máquina a utilizar

Para a construção de um modelo de classificação capaz de prever a credibilidade de anúncios de emprego do *LinkedIn*, foram considerados quatro algoritmos amplamente adotados em problemas de aprendizagem supervisionada com dados textuais rotulados: Naives Bayes, regressão logística, random forest e support vector machine (SVM). A decisão sobre o algoritmo mais adequado não se baseia apenas na literatura, mas também considera a realidade prática da solução proposta, o ambiente de uso (extensão de navegador), e as preferências identificadas junto ao público-alvo, conforme indicado na triangulação metodológica adotada neste trabalho.

3.5.1 Fundamentacao com base na literatura técnica

Estudos clássicos e recentes (Aggarwal & Zhai, 2012; Bansal, 2020) apontam que algoritmos lineares como regressão logística e Naive Bayes apresentam excelente desempenho em tarefas de classificação de texto com alto volume de dados esparsos, especialmente quando vectorizados com técnicas com TF-IDF.

- O **Naive Bayes**, embora extremamente eficiente e leve, assume independencia entre as palavras, o que pode limitar sua acurácia, especialmente em contextos de linguagem ambígua ou com dependências sintáticas complexas.
- A **SVM**, embora poderosa e precisa, é mais adequada para contextos onde não se tem limitações de performance e recursos computacionais, o que não é o caso de uma extensão de navegador leve e responsiva.
- Já, o **Random Forest**, embora robusto e eficaz com atributos mistos, apresenta dificuldades com dados textuais esparsos, além de demandar mais memória e dificultar interpretação dos resultados.

Tabela 2: Comparação entre algoritmos de aprendizagem de máquina

Algoritmo	Vantagens	Desvantagens
Naive Bayes	Simples, rápido, eficaz para texto com TF-IDF, boa interpretabilidade	Assume independência entre as características, o que pode não ser realista
Regressão logística	Interpretação clara das probabilidades, leve, bom desempenho com TF-IDF	Pode ter dificuldades com relações não lineares
Random Forest	Alta precisão, robusto a overfitting, bom com dados com ruídos	Mais pesado, mais lento, menos interpretável
SVM	Alta precisão, bom espaço de alta dimensão (TF-IDF)	Requer mais tempo e recursos; difícil de calibrar probabilidades

3.5.2 Fundamentação com base nas características da solução

Dado que a solução proposta será incorporada em uma extensão de navegador, que visa fornecer uma análise rápida e clara ao utilizador, o modelo precisa ser:

- Leve e com **tempo de inferência rápido**;
- **Fácil de integrar** com *backend* simples
- **Interpretável**, para facilitar futuros ajustes;
- **Treinável com dados textuais vectorizados**, como aqueles extraídos das descrições dos anúncios via TF-IDF.

A regressão logística atende de forma ideal a todos esses critérios. Ela apresenta excelente desempenho em tarefas de classificação binária com dados textuais e permite ajustar probabilidades de forma intuitiva, o que é essencial para apresentar percentuais de credibilidade ao utilizador, como desejado na interface da extensão.

3.5.3 Fundamentação com base na pesquisa com o público-alvo

Os dados colectados por meio do formulário aplicado ao público-alvo também reforçam a escolha por uma abordagem leve e eficiente. Quando questionados sobre suas expectativas em relação à experiência de uso da extensão, maior parte

dos respondentes demonstraram interesse por uma solução que fosse rápida, simples de usar e integrada à sua rotina de navegação, sem causar atrasos ou impacto significativo na experiência com plataforma.

Tais preferências descartam, em um primeiro momento, o uso de algoritmos mais pesados como o Random Forest ou SVM, embora ofereçam alta precisão, demandam maior capacidade computacional e tempo de processamento, o que poderia prejudicar a experiência do utilizador em um ambiente web.

Além disso, a própria natureza da aplicação, uma extensão de navegador para uso cotidiano, demanda um modelo com baixo custo computacional, que possa ser executado de forma eficiente por meio de requisições ao backend, sem sobrecarregar os recursos da máquina do utilizador ou servidor.

Portanto, aliando os benefícios técnicos documentados na literatura à preferência expressa pelos futuros utilizadores, conclui-se que o algoritmo de regressão logística representa melhor escolha para o desenvolvimento de uma versão funcional, com boa precisão, interpretação e baixo consumo de recursos.

4 Capítulo IV - Treinamento do modelo computacional

O processo de construção do modelo computacional proposto neste trabalho baseia-se na utilização de técnicas de aprendizagem de máquina supervisionada, sendo o objectivo final estimar a credibilidade de anúncios de emprego publicados em plataformas *online*. Este capítulo descreve as etapas essenciais para o treinamento do modelo, abordando **a preparação dos dados**, **a escolha de algoritmos**, **a avaliação do desempenho** e o **armazenamento do modelo final**.

4.1 Descrição dos dados

O *dataset* escolhido conforme apresentado em 3.4, apresenta uma variedade de campos textuais, tais como: título do anúncio, descrição, requisitos, localização e outros atributos relevantes. O mesmo encontra-se organizado e com as etiquetas verificadas, o que facilita o desenvolvimento de modelos de classificação supervisionada.

As principais vantagens do *dataset* incluem:

- boa qualidade dos dados;
- equilíbrio relativo entre as classes “*fake*” e “*real*”;
- compatibilidade com algoritmos de classificação textual;
- utilização reconhecida em pesquisas académicas e estudos comparativos.

4.2 Pré-processamento dos dados

Antes do treinamento do modelo, é necessário realizar o pre-processamento dos dados, que é uma etapa fundamental em tarefas de NLP, como o objectivo de transformar o texto bruto em uma representação estruturada que possa ser compreendida por algoritmos de aprendizagem de máquina (CAMACHO-COLLADOS; PILEHVAR. 2018). O pré-processamento inclui os seguintes passos:

- **Limpeza dos textos:** remoção de caracteres especiais, pontuações, números e espaçamentos desnecessários, considerados ruídos em tarefas de classificação textual (INDURKHYA; DAMERAU, 2010);
- **Normalização:** conversão de textos para palavras minúsculas, para evitar que palavras idênticas, porém com grafias diferentes sejam tratadas como distintas (JURAFSKY; MARTIN, 2021);

- **Remoção de stopwords:** palavras muito frequentes que não agregam valor semântico significativo, como artigos e preposições, são removidas para reduzir a dimensionalidade dos dados e eliminar ruídos (MANNING; RAGHAVAN; SCHÜTZE, 2008);
- **Tokenização e vectorização:** o texto é transformado em vetores numéricos por meio do método TF-IDF (*Term Frequency-Inverse Document Frequency*), que avalia a relevância de cada termo no documento em relação ao conjunto de dados como um todo (RAJARAMAN; ULLMAN, 2021).

4.3 Treinamento e avaliação

Para determinar o modelo mais adequado à proposta, foram comparados quatro algoritmos clássicos de aprendizagem de máquina, comumente utilizados em tarefas de classificação de texto, conforme apresentado no capítulo 3, mais especificamente em 3.5.

O treinamento do modelo foi realizado localmente em um computador portátil de marca Apple, modelo MacBook Pro, com o processador Apple M1 Pro 3.2 GHz e memória de 16GB, utilizando algumas bibliotecas *Python*: *pandas*, *scikit-learn*, *seaborn*, *joblib*.

As métricas adoptadas para avaliar o modelo do desempenho foram:

- **Acurácia:** proporção total de previsões;
- **Precisão:** taxa de acertos nas previsões positivas;
- **Revocação:** taxa de recuperação das amostras positivas
- **F1-score:** média harmônica entre precisão e revocação

Os resultados obtidos demonstraram um desempenho satisfatório para um produto mínimo viável (MVP), permitindo classificar anúncios com boa precisão e resposta rápida, essenciais para a aplicação em tempo real via extensão de navegador.

4.4 Armazenamento do modelo

Após o treinamento, é necessário efectuar a serialização. Para tal, o *joblib*, que é uma biblioteca do *Python*, permite que o modelo seja carregado e posteriormente utilizado no *backend* da aplicação. Com isso, a extensão do navegador poderá enviar o texto de um anúncio para o *backend*, que por sua vez retornará uma percentagem de credibilidade baseada na classificação feita pelo modelo.

5 Capítulo V - Desenvolvimento da solução proposta

Conforme abordado no capítulo I, a metodologia seleccionada para o desenvolvimento da solução proposta é uma adaptação da metodologia de desenvolvimento waterfall cujas fases, bem como os seus artefactos serão detalhados a seguir.

5.1 Elicitação dos requisitos

A elicitação de requisitos é uma etapa fundamental no desenvolvimento de *software*, sendo responsável pela identificação e documentação das necessidades e expectativas dos *stakeholders* em relação ao sistema a ser desenvolvido. De acordo com Sommerville (2011), a elicitação dos requisitos envolve actividades como colecta, análise e documentacao dos requisitos. Para este trabalho, a elicitação foi realizada com base em uma abordagem documental e analítica, utilizando informações obtidas em fontes secundárias, incluindo: revisão de literatura e estudos sobre as práticas de credibilidade e segurança em plataformas de emprego online.

Com base no levantamento documental realizado, foram identificados os seguintes problemas principais:

- falta de mecanismo de verificação de anúncios em tempo real;
- dificuldade de diferenciação entre anúncios legítimos e fraudulentos por parte dos utilizadores do LinkedIn e o consequente
- risco de exposição a fraudes e roubo de informações pessoais ao interagir com anúncios suspeitos.

Pelo que, essencialmente, a solução deve permitir que os utilizadores do linkedIn consigam de forma rápida e fluida ter informações sobre a legitimidade dos anúncios de emprego, ajudando-lhes na tomada de decisão e gestão até de expectativas.

5.2 Análise da solução

5.2.1 Stackholders da solução

A parte mais importante de qualquer projecto de *software* são os *stakeholders*. *Stakeholders* são o conjunto de partes interessadas na solução. Para a solução proposta, não existe nenhum grupo a não ser o utilizador do LinkedIn. Este é o único

utilizador desta plataforma, pois ela será desenvolvida especificamente para o LinkedIn.

Tabela 3: Stakeholders e os benefícios da solução proposta

Stakeholders	Benefícios
Utilizador do <i>LinkedIn</i>	Através desta solução, os utilizadores do <i>LinkedIn</i> poderão ter a capacidade de ter uma análise de credibilidade das oportunidades que lhes são apresentadas, reduzindo assim o número de fraudes que tem resultado em impactos financeiros, emocionais, entre outros.

5.2.2 Requisitos do sistema

Para Pressman (2016), requisitos de sistema são uma descrição detalhada das funções, características e restrições que um sistema deve possuir para atender às expectativas do utilizador e às metas do negócio.

- **Priorização dos requisitos**

Karl Wiegers (2013), discute sobre técnicas de priorização classificando os requisitos em essencial, importante e desejável.

- **Essencial:** é o requisito sem o qual o sistema não pode funcionar ou que inviabiliza o objectivo principal do projecto. Esses requisitos são obrigatórios e sua ausência impede o funcionamento do sistema.
- **Importante:** um requisito que melhora significativamente o sistema, mas cuja ausência não inviabiliza o funcionamento básico. Ele é implementado caso os recursos e o tempo permitam.
- **Desejável:** requisito que agrega valor adicional, mas cuja ausência não prejudica o sistema nem as expectativas dos *stakeholders*. Estes são normalmente considerados opcionais.

- **Requisitos funcionais**

Para Pressman (2016), requisitos funcionais são especificações que descrevem as interações entre o sistema e seus utilizadores ou entre diferentes partes do sistema.

Tabela 4: Descrição dos requisitos funcionais

ID	Requisito	Descrição	
RF01	Seleção de texto	O utilizador deve poder seleccionar o texto do anúncio de vaga desejado para a extensão	Essencial
RF02	Visualização dos resultados	Após o processamento de texto e posterior análise, os utilizadores devem ter acesso aos resultados relativamente ao nível de credibilidade	Essencial
RF03	Apresentação de um relatório detalhado dos resultados	Os utilizadores precisam ter detalhes sobre os resultados do modelo.	Desejável

- **Requisitos não funcionais**

Sommerville (2011) define esses requisitos como restrições que afectam a forma como o sistema deve operar, como limitações de *hardware*, requisitos de desempenho ou padrões organizacionais.

Tabela 5: Descrição dos requisitos não funcionais

ID	Requisito	Descrição	
RNF01	Precisão	O intervalo de confiança do algoritmo de ML, deve ser de pelo menos 92%	Essencial
RNF02	Facilidade de uso	Para evitar complexidade aos utilizadores, o sistema deve ser fácil de usar	Desejável

RNF03	Tempo de resposta	O utilizador deve ter a resposta do sistema em menos de 60s, garantindo uma boa experiência de utilizador	Importante
RNF04	Compatibilidade	A extensão deve ser compatível com as versões mais utilizadas do <i>Google Chrome</i>	Importante

- **Restrições de desenvolvimento**

A secção de restrições do desenvolvimento consiste em abordar aspectos relacionados a implementação da solução no que diz respeito às tecnologias utilizadas no desenvolvimento da solução.

Quanto a linguagem de programação, foram utilizadas as seguintes:

- *Python*: É a linguagem de programação principal para o desenvolvimento do projecto, incluindo a análise de dados, desenvolvimento do modelo de aprendizagem de máquina e da API.
- *Javascript*: É a linguagem fulcral para o desenvolvimento da extensão a ser utilizadas do lado do cliente

No que diz respeito ao ambiente de desenvolvimento, estas são as ferramentas:

- *Jupyter Notebook*: ambiente para exploração inicial do *dataset*
- *Visual Studio Code*: editor de código leve para a codificação da API e da extensão,

5.2.3 Modelagem da solução proposta

A modelagem da solução proposta segue anotações UML, que segundo Fowler (2004) é uma coleção de diagramas para auxiliar no projecto e na documentação de sistemas de *software*. Fowler acrescenta dizendo que é um tipo de anotação amplamente usada em sistemas orientados a objectos, mas também pode ser adaptado para outras abordagens.

5.2.4 Modelos de casos de uso

Os modelos de casos de uso são utilizados para capturar os requisitos funcionais e organizar as interações em cenários compreensíveis. Jacobson e al. (1992), define casos de uso como uma especificação de um conjunto de acções realizadas por um sistema que produz um resultado observável de valor para um ator específico.

Os componentes de um caso de uso são

- **Ator:** utilizador ou outro sistema que interage com o sistema principal
- **Objectivo:** Resultado esperado da interação
- **Fluxo principal:** Sequência típica de passos para alcançar o objectivo
- **Fluxos alternativos:** Cenários que desviam do fluxo principal devido a variações ou erros.

5.2.5 Casos de uso

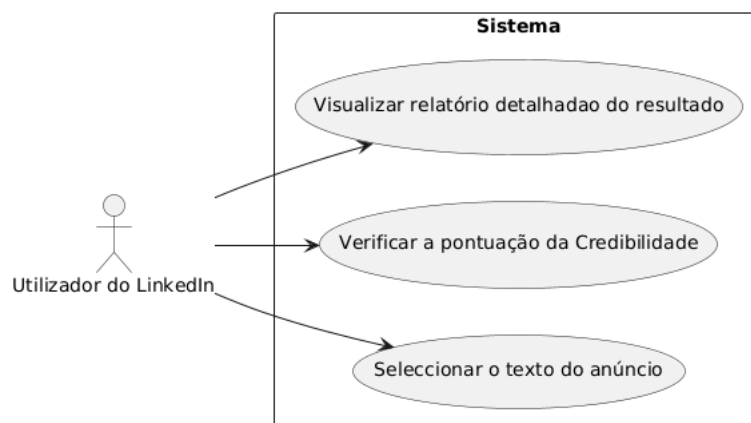


Figura 11: Diagrama de casos de uso

5.3 Projecto

A arquitectura de um projecto de software descreve a estrutura de alto nível dos seus componentes, as suas inter-relações e os princípios que governam o seu *design* e evolução (Bass, Clements & Kazman, 2012). Compreender a arquitectura é crucial para o desenvolvimento, manutenção e escalabilidade de qualquer sistema complexo. No contexto, deste projecto, e inspirados na estrutura modular e nos mecanismos de comunicação observados nas extensões do Google Chrome (conforme ilustrado na figura 1), podemos delinear uma arquitectura que promove a separação de responsabilidades e a comunicação eficiente entre diversos módulos.

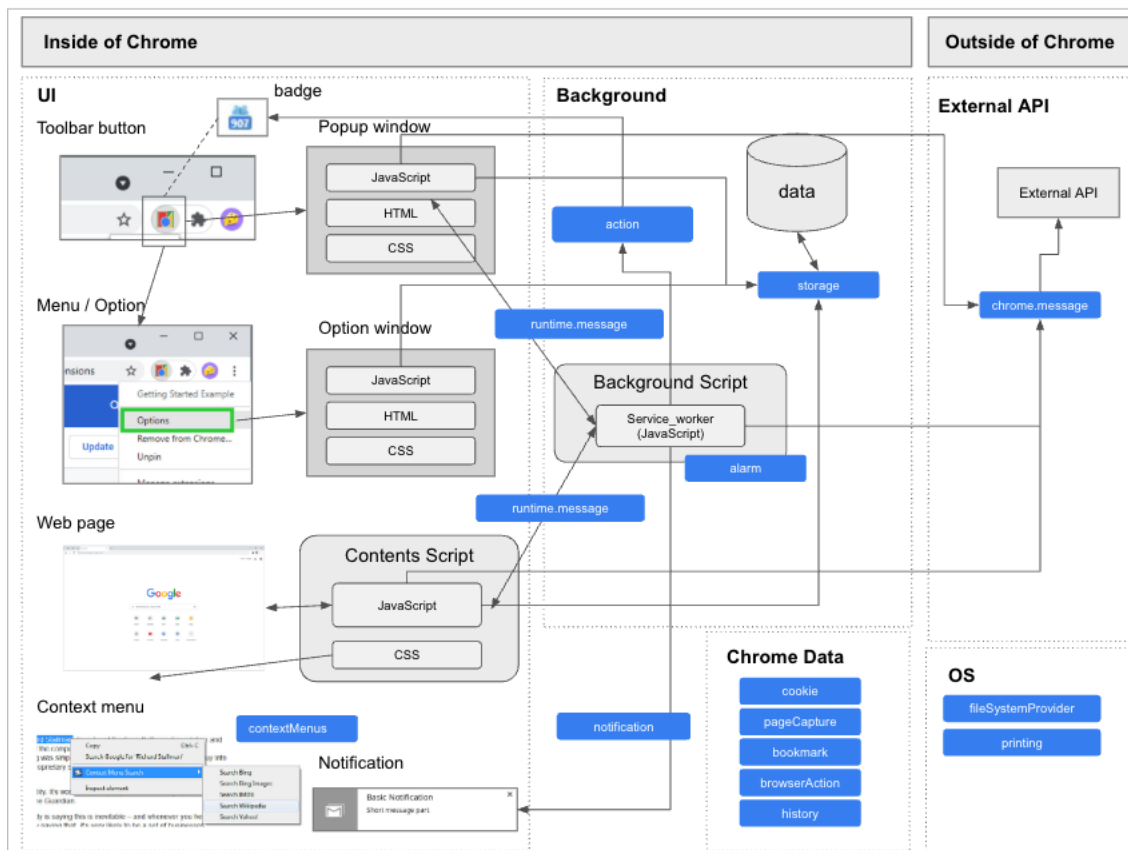


Figura 12: Arquitetura de uma extensão do google chrome (Fonte: Publicação de Yoshi no Medium, 2022)

A figura 12, divide a arquitetura da extensão em duas partes fundamentais: a parte interna e a externa. Abaixo iremos descrever o que cada uma delas de facto contém:

1. Dentro do Google Chrome

Aqui encontramos os principais elementos que compõem a experiência do utilizador e a lógica interna da extensão.

A. Interface do utilizador (UI - User Interface)

Esta camada representa os elementos visuais com os quais o utilizador interage directamente. Vemos exemplos como:

- **Toolbar button:** um ícone na barra de ferramentas do Chrome que permite acesso rápido às funcionalidades da extensão.
- **Badge:** um pequeno indicador visual sobre o ícone do *toolbar button*, usado para exibir informações ou notificações.
- **Popup window:** uma pequena janela que surge quando o utilizador interage com o toolbar button, geralmente construída através de tecnologias como **HTML**, **CSS** para a estrutura e estilo, e **JavaScript** para interatividade.

- **Menu / Option:** indica um menu de opções acessíveis através de algum elemento da UI.
- **Option window:** uma página dedicada para configurar as opções da extensão, também construída com HTML, CSS e JavaScript.
- **Web page:** a própria página *web* que o utilizador está a visualizar o Chrome.
- **Context menu:** o menu que aparece quando o utilizador clica com o botão direito numa página *web*, onde a extensão pode adicionar as suas próprias opções.
- **Notification:** mensagens que o Chrome exibe ao utilizador fora do contexto da página *web*, informando sobre eventos relevantes.

B. Background

Esta camada lida com a lógica da extensão que não está directamente ligada a uma janela ou página específica. É composta por:

- **Background Script** (*Service Worker* - JavaScript): um script JavaScript que é executado em segundo plano. Ele pode realizar tarefas como monitorar eventos do navegador, gerir dados e comunicar com outras partes da extensão. A menção a *service worker* indica uma abordagem moderna para executar scripts em background de forma eficiente.
- **Alarm:** uma funcionada que permite ao script de background agendar tarefas para serem executadas em momentos específicos.
- **Data:** representa os dados que a extensão pode precisar armazenar e gerir.
- **Storage:** a API do Chrome que permite à extensão armazenar dados localmente (como configurações do utilizador). As setas indicam que o *Background Script* pode armazenar e aceder esses dados.
- **Action:** refere-se às acções que o Background Script pode realizar em resposta a eventos ou agendamentos.
- **Runtime.message:** um mecanismo de comunicação que permite que diferentes partes da extensão, como **Popup window**, **Option window** e **Contents Script**, troquem mensagens com o Background Script

C. Content Script

Este script JavaScript é injectado nas páginas *web* que o utilizador visita. Ele pode interagir com o conteúdo da página (DOM), ler informações e até mesmo modificá-lo. Geralmente é acompanhado por CSS para aplicar estilos específicos. Na página injectada. A comunicação com o Background Script também ocorre através de *runtime.message*.

D. Chrome Data

Este representa os dados internos do navegador Chrome que podem ser acedidos e manipulados pelas extensões através de APIs específicas.

2. Fora do Chrome

Esta secção descreve como o Google Chrome e especificamente a extensão, interage com o mundo exterior:

- **External API:** representa serviços ou aplicações externas com as quais a extensão pode precisar comunicar. A comunicação é feita através da API *chrome.message*, indicando que o Chrome fornece uma ponte para essa interação.
- **OS (Sistema operativo):** o sistema operativo subjacente também desempenha um papel fundamental:
 - **fileSystemProvider:** uma API que permite que extensões interajam com o sistema de ficheiros do utilizador (com as devidas permissões).
 - **Printing:** funcionalidade de impressão do sistema operativo que o Chrome utiliza.

5.4 Codificação do protótipo

A solução proposta visa auxiliar os utilizadores na identificação de anúncios de emprego potencialmente falso na plataforma LinkedIn, utilizando-se de uma abordagem baseada em aprendizado de máquina. Para isso, foi desenvolvida uma extensão de navegador, especificamente para o Google Chrome, que permite a leitura do conteúdo textual dos anúncios e a avaliação da sua credibilidade, retornando ao utilizador um uma percentagem de credibilidade.

A extensão pode ser executada directamente no navegador do utilizador, sem a necessidade de instalação de um *software* adicional ou dependências externas, tornando-se uma solução prática e de rápida adoção, especialmente num contexto em que a maioria das interações com redes sociais e plataformas de emprego ocorre via *web*.

Do ponto de vista técnico, a extensão foi desenvolvida com base no **Manifesto versão 3** (*Manifest v3*), a especificação actual para extensões do navegador Google Chrome. Os principais *scripts* são escritos em JavaScript e a interface utiliza HTML e CSS para a parte visual dos componentes gráficos.

5.4.1 Estrutura do protótipo

Para ilustrar como a estrutura do código da solução proposta se alinha com a arquitectura da extensão do Chrome, a seguinte tabela mapeia os componentes arquitectónicos identificados na Figura 12 com os ficheiros de código para a solução proposta:

Tabela 6: Estrutura do protótipo

Componente arquitectónico (figura 12)	Ficheiros específicos do projecto	Descrição e relação com a arquitectura
UI (Popup window)	- <i>popup.html</i> (estrutura) - <i>popup.js</i> (lógica) - <i>style.css</i> (estilos)	Apresenta a interface visual ao utilizador, permitindo interação e o acionamento da funcionalidade de colecta de dados do LinkedIn. O <i>style.css</i> garante a apresentação visual adequada
Contents Script	<i>content.js</i>	Responsável por ser injectado nas páginas do LinkedIn e implementar a lógica específica para

		recolher o conteúdo desejado.
Background Script	<i>Background.js</i>	Embora esteja planeada para o futuro, ele representa a lógica em segundo plano (notificações, eventos, etc.) que pode complementar a colecta de dados ou outras funcionalidades da extensão.
Manifest File (implícito)	<i>Manifest.json</i>	Define as permissões necessárias para aceder às páginas do LinkedIn e para executar o <code>contente.js</code> . Também declaram o <code>popup.html</code> e o <code>background.js</code> essencialmente configurando a arquitectura da extensão para o Chrome

5.4.2 Testes do protótipo

Para realização dos testes da extensão desenvolvida, utilizou-se o navegador Google Chrome, em um computador portátil de marca Apple, modelo MacBook Pro, com o processador Apple M1 Pro 3.2 GHz, memória de 16GB e sistema operativo macOS Sonoma 14.6.1. O teste consistiu em navegar pela plataforma LinkedIn e aceder a anúncios de emprego e ativar a extensão a partir da barra de extensões do navegador.

Durante o uso, a extensão permite que o utilizador seleccione livremente o conteúdo textual de um anúncio de emprego visualizado no LinkedIn. Ao accionar a

funcionalidade principal da extensão, o texto previamente destacado é capturado e submetido a uma rotina de análise.

Essa rotina envolve o envio do conteúdo para uma API responsável por processar e avaliar a credibilidade da publicação com base em critérios predefinidos. O resultado da análise é apresentado de forma clara e intuitiva, por meio de uma percentagem que expressa o nível de confiabilidade identificado no texto.

A seguir, apresentam-se algumas imagens ilustrativas que espelham o funcionamento da extensão no ambiente real de navegação:

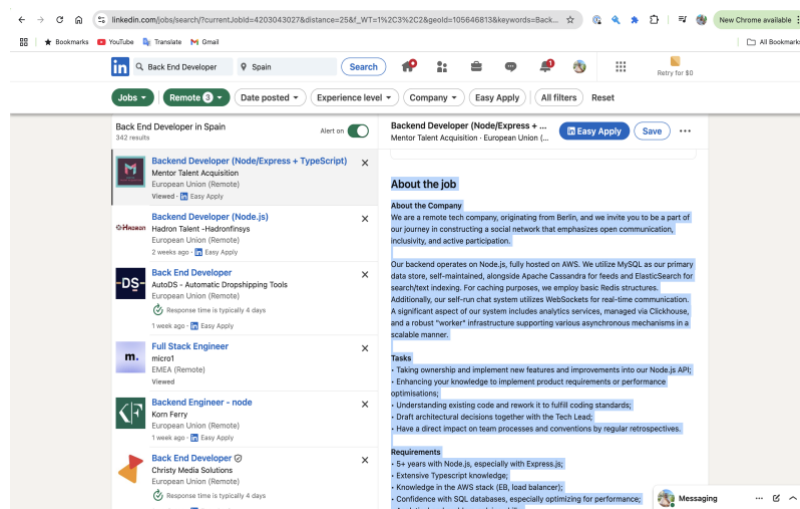


Figura 13: Selecção do texto a ser analisado

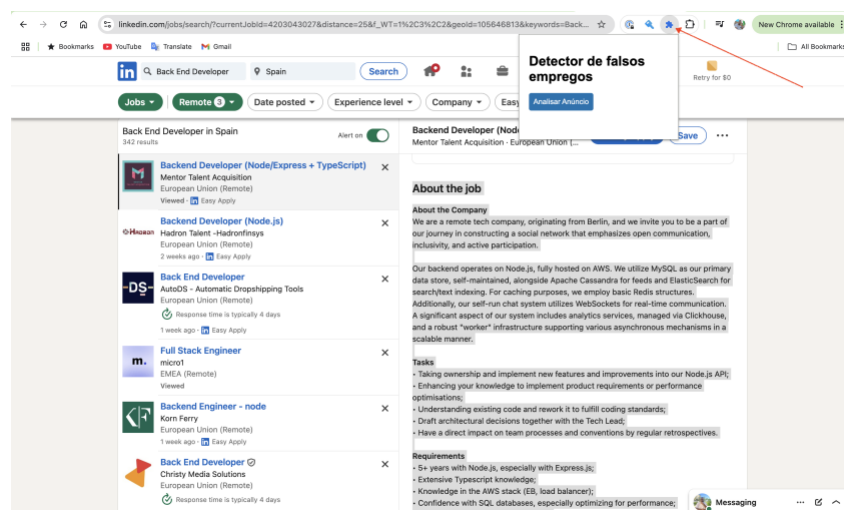


Figura 14: a extensão do navegador, antes da realização da análise.

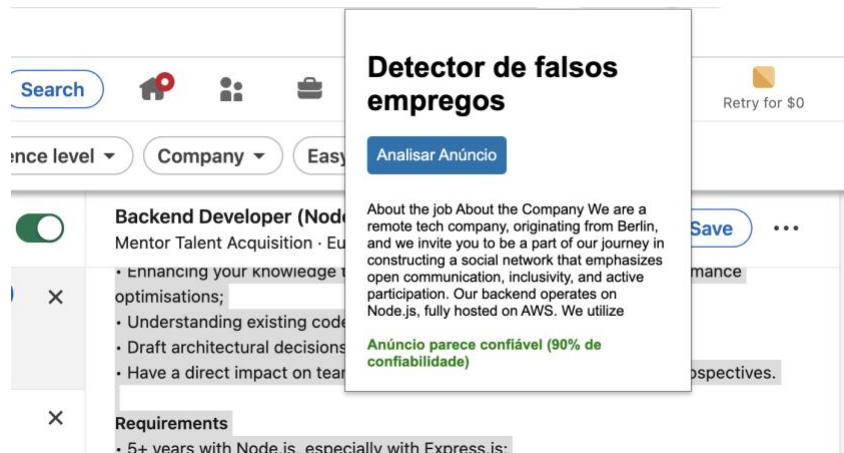


Figura 15: Resultado da análise de credibilidade de anúncio confiável.

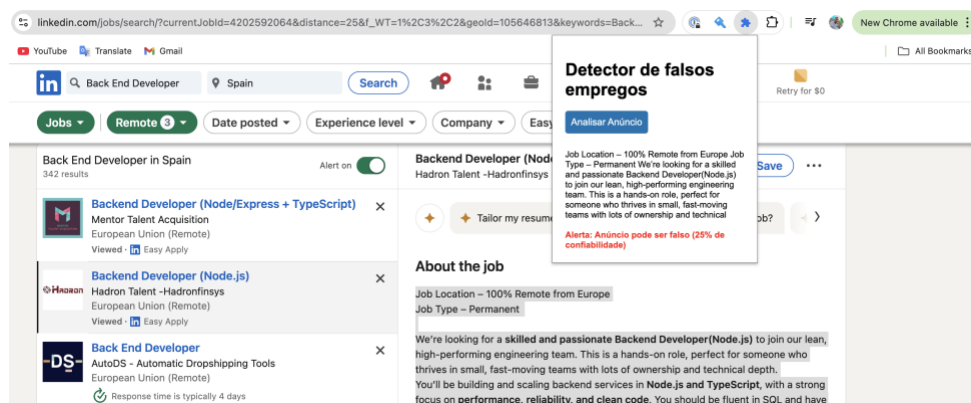


Figura 16: Resultado de análise de credibilidade de um anúncio não confiável.

6 Capítulo VI – Teste do algoritmo

O objectivo principal deste capítulo é identificar um modelo de aprendizagem de máquina capaz de distinguir anúncios de emprego legítimos dos falsos, e para tal, foram realizados diversos testes com diferentes algoritmos. Essa etapa foi crucial para avaliar o desempenho de cada abordagem, entender os desafios da base de dados e refinar a solução proposta para alcançar resultados mais robustos.

6.1 Contextualização do *Dataset*

O *dataset* utilizado contém aproximadamente 18 mil descrições de vagas, das quais cerca de 800 são falsas, representando cerca de 4.4% do total. Ou seja, temos um cenário altamente desequilibrado, em que a classe majoritária representa anúncios legítimos, mas, mais adiante iremos implementar mecanismos para lidar com essa situação.

6.2 Metodologia

O pipeline de testes seguiu as seguintes etapas:

1. Pré-processamento textual

Conforme detalhado no capítulo IV, especificamente em 4.2, neste passo foi aplicada vectorização por meio do método TF-IDF para transformar o conteúdo textual dos anúncios em vectores números compreensíveis pelos algoritmos.

2. Divisão dos dados

Para os testes, o *dataset* foi dividido em 80% para o treino e 20% para o teste, garantindo uma avaliação eficaz da capacidade de generalização do modelo sobre dados não vistos durante o treinamento. Esta prática é amplamente adoptada na literatura especializada em aprendizagem de máquina. Segundo Géron (2019), uma divisão 80/20 é frequentemente utilizada, pois oferece boa quantidade de dados para o treinamento sem comprometer a confiabilidade da proporcione um equilíbrio adequado entre aprendizado e validação.



Figura 17: Divisão dos dados para teste (fonte: Guido Mascia)

3. Avaliação dos modelos

Conforme descrito no capítulo 4, especificamente em 4.3, para avaliar o desempenho do modelo, foram utilizadas as seguintes métricas:

- **Precisão (*Precision*):** mede a proporção de verdadeiros positivos em relação ao total de predições positivas feitas pelo modelo.
- **Revocação (*Recall*):** mede a proporção de verdadeiros positivos identificados em relação ao total de exemplos realmente positivos.
- **F1-Score:** é a média harmônica entre precisão e revocação, proporcionando uma métrica mais equilibrada.
- **Acurácia (*Accuracy*):** mede a proporção de predições correctas no total de predições feitas, porém pode ser enganosa em casos de desequilíbrio de classes.

6.3 Algoritmos avaliados e resultados

6.3.1 Regressão logística (modo padrão)

Inicialmente, utilizou-se o algoritmo de regressão logística com configurações padrão, ou seja, sem ajustes específicos para lidar com o desequilíbrio dos dados. Neste modo, o modelo considera cada classe igualmente importante, sem levar em conta a quantidade desigual de exemplos

Resultados da classificação

```
○ Delfims-MacBook:project delfimuqueio$ python3 fake_job_posting_detection.py
Matplotlib is building the font cache; this may take a moment.
● Acurácia: 0.9677

■ Relatório de Classificação:
      precision    recall  f1-score   support

     0       0.97       1.00       0.98       2895
     1       1.00       0.31       0.47        142

 accuracy          0.97          0.97          0.96          3037
 macro avg          0.98          0.65          0.73          3037
 weighted avg       0.97          0.97          0.96          3037
```

Figura 18: Resultados da classificação (regressão logística padrão)

Tabela 7: Resultados da classificação (regressão logística padrão)

Classe	Precisão (Precision)	Revocação (Recall)	F1-Score
Falsa (0)	97%	100%	98%
Legítima (1)	100%	31%	47%

Acurácia geral do modelo: 96%

Apesar da alta acurácia, o modelo praticamente ignorava a classe 1 (legítimos), obtendo um *recall* de apenas 31%. Isso significa que ele acertava quase todos os falsos, mas errava a maioria dos anúncios legítimos.

6.3.2 Regressão logística (modo balanceado)

Para resolver o problema encontrado ao usar o algoritmo de regressão logística no modo padrão, o modelo foi configurado com a opção `class_weight='balance'`. Esse parâmetro faz com que o algoritmo dê maior peso à classe minoritária durante o treinamento, compensando a sua baixa frequência.

Relatório de classificação

```
■ Relatório de Classificação:
      precision    recall  f1-score   support

     0       0.99       0.96       0.98       2895
     1       0.52       0.80       0.63        142

 accuracy          0.96          0.96          0.96          3037
 macro avg          0.76          0.88          0.80          3037
 weighted avg       0.97          0.96          0.96          3037
```

Figura 19: Resultados da classificação (regressão logística com balanceamento)

Tabela 8: Resultados da classificação (regressão logística com balanceamento)

Classe	Precisão (<i>Precision</i>)	Revocação (<i>Recall</i>)	F1-Score
Falsa (0)	99%	96%	98%
Legítima (1)	52%	80%	63%

Acurácia geral do modelo: 96%

Com essa alteração, o modelo conseguiu identificar muito melhor os anúncios legítimos, alcançando um *recall* de 0.80 (80%) e um *f1-score* de 0.63. Essa mudança garantiu um comportamento mais equilibrado entre as duas classes, sem comprometer o desempenho geral.

7 Capítulo VII – Discussão de resultados

Tendo em vista o alcance do objectivo geral deste trabalho que é o de desenvolver uma extensão de navegador para análise de credibilidade de anúncios de emprego no LinkedIn utilizando aprendizagem de máquina e processamento de linguagem natural (NLP), foram realizadas a revisão de literatura, pesquisa documental e partilha de um formulário ao público. Estas etapas foram fundamentais para a compreensão do problema, definição dos requisitos da solução e validação da abordagem adoptada. Como produto deste percurso, foi desenvolvido um modelo de classificação de credibilidade baseado em aprendizagem de máquina e, posteriormente, implementado um protótipo funcional da extensão para testes práticos.

7.1 Revisão de literatura

A revisão de literatura focou-se no levantamento de conceitos essenciais relacionados à detecção de fraudes em ambientes online, com ênfase no contexto dos anúncios de emprego. Através da consulta a fontes especializadas e estudos recentes, foi possível confirmar que o número de esquemas fraudulentos envolvendo ofertas de trabalho em plataformas digitais vem crescendo consideravelmente, impulsionado pela popularidade dos canais de recrutamento online.

Os estudos analisados evidenciam que técnicas de processamento de linguagem natural, quando associadas a modelos de aprendizagem de máquina supervisionado, podem desempenhar um papel decisivo na identificação de padrões comuns em anúncios de emprego fraudulentos, como o uso de linguagem excessivamente vaga, promessas irreais, exigências atípicas de pagamento antecipado ou ausência de informações verificáveis sobre a empresa contratante.

Ficou igualmente claro que, embora existam mecanismos de segurança implementados pelas próprias plataformas, como *LinkedIn*, esses sistemas nem sempre conseguem impedir a publicação de anúncios fraudulentos em tempo hábil. Assim, soluções complementares como extensões de navegador capazes de fornecer análise adicional e independente dos conteúdos visualizados, mostram-se de grande utilidade prática.

Além disso, observou-se que a adopção de modelos de aprendizagem de máquina oferece vantagens significativas em relação às abordagens tradicionais, por permitir a evolução contínua da ferramenta em resposta às mudanças nas estratégias

utilizadas por fraudadores. A capacidade de aprender a partir de novos exemplos torna o sistema mais resiliente e adaptável a novas tentativas de fraude.

Em síntese, a revisão de literatura sustentou a relevância proposta desenvolvida neste trabalho, ao demonstrar tanto a gravidade do problema abordado quanto a viabilidade e vantagem do uso de técnicas modernas de inteligência artificial para enfrentá-lo.

7.2 Proposta de solução

A análise dos dados recolhidos através do formulário, aliados à fundamentação teórica apresentada nos capítulos anteriores, permitiu validar a relevância, viabilidade e adequação da solução proposta para detecção de anúncios fraudulentos no *LinkedIn*. Os resultados discutidos a seguir consolidam os principais pontos que fundamentaram as decisões de projecto.

A proposta de solução apresentada neste trabalho visa combater a presença de falsos anúncios de emprego no *LinkedIn*, um problema que, conforme discutido anteriormente, pode trazer prejuízos consideráveis tanto para candidatos quanto para empresas. Para garantir maior confiabilidade e minimizar vieses, adotou-se a triangulação metodológica, que combina revisão da literatura e coleta de dados empíricos junto ao público-alvo. Esta abordagem permitiu uma visão mais ampla e fundamentada sobre o problema e contribuiu para a definição de decisões estratégicas da solução proposta.

A solução projetada está dividida em duas dimensões complementares: o modelo de aprendizado de máquina, responsável pela análise da credibilidade dos anúncios, e o sistema de interação com o utilizador, materializado em uma extensão de navegador. A escolha do Google Chrome como navegador prioritário para o desenvolvimento da extensão foi baseada em dados robustos. Tanto a literatura como a pesquisa realizada junto ao público indicam que o Chrome lidera em participação de mercado e é o navegador mais utilizado para acesso ao *LinkedIn*, o que maximiza o potencial de adoção da ferramenta.

O funcionamento da extensão foi pensado a partir da análise combinada das preferências do público-alvo e das limitações técnicas existentes. Embora a maioria dos utilizadores tenha manifestado preferência por filtros nativos ou mecanismos internos no *LinkedIn*, a revisão da literatura revelou que estas abordagens são

limitadas em termos de transparência, controle do utilizador e capacidade de adaptação. Além disso, as restrições legais e técnicas impostas pelas políticas do *LinkedIn* tornam inviável uma integração direta com suas *APIs* para auditoria de anúncios. Diante deste cenário, a extensão de navegador se mostrou a alternativa mais viável, equilibrando autonomia do utilizador, transparência na avaliação e viabilidade técnica. O fluxo de funcionamento proposto, onde o utilizador seleciona o texto, a extensão o extrai e envia para um *backend* que retorna a pontuação de credibilidade, foi estruturado para garantir uma experiência de uso intuitiva e eficiente. O *dataset* selecionado para treinar o modelo de machine learning, o *Fake Job Posting Prediction* do *Kaggle*, fornece uma base sólida, composta por cerca de 18.000 anúncios rotulados, permitindo um treinamento eficaz do modelo. Para a escolha do algoritmo de aprendizagem de máquina mais adequado, foram considerados fatores técnicos, as características da solução e a experiência do utilizador. A análise comparativa entre Naive Bayes, Regressão Logística, *Random Forest* e SVM indicou que a Regressão Logística apresenta o melhor equilíbrio entre desempenho, leveza, interpretação e facilidade de integração, aspectos fundamentais para uma solução embutida em uma extensão de navegador.

Assim, a solução proposta neste trabalho é fruto de um processo criterioso que combina evidências da literatura, dados coletados diretamente do público-alvo e análise técnica. Ao desenvolver uma extensão para o Chrome com um modelo leve e eficiente de classificação, pretende-se oferecer uma ferramenta prática e confiável para auxiliar utilizadores do *LinkedIn* a avaliarem a credibilidade de anúncios de emprego, contribuindo, assim, para um ambiente de recrutamento *online* mais seguro e transparente.

8 Capítulo VIII – Considerações finais

8.1 Conclusões

A detecção de anúncios de emprego fraudulentos é uma questão crescente, especialmente em plataformas como o *LinkedIn*, onde a busca por oportunidades profissionais é constante e os utilizadores frequentemente se deparam com anúncios enganosos. Durante o desenvolvimento deste trabalho, foi possível identificar que a solução proposta, que integra aprendizagem de máquina e processamento de linguagem natural (PLN), apresenta um potencial significativo para resolver essa problemática. A utilização da tecnologia de aprendizagem de máquina, através do algoritmo de regressão logística, permitiu uma análise eficiente da credibilidade dos anúncios de emprego, proporcionando uma forma rápida e precisa de identificar conteúdos suspeitos.

Através do desenvolvimento de uma extensão de navegador, foi possível implementar uma ferramenta prática que pode ser facilmente integrada ao fluxo de uso diário dos utilizadores do *LinkedIn*. A extensão, ao processar e avaliar a credibilidade dos anúncios, oferece uma camada adicional de segurança e confiança, essencial para os candidatos que buscam informações fidedignas sobre as vagas de emprego. Além disso, o trabalho demonstrou que a utilização de técnicas de NLP, como a análise de texto, é eficaz para extrair padrões e características de anúncios que indicam a sua autenticidade ou, em contrapartida, a sua falsidade.

Durante a pesquisa, foi possível constatar que a ausência de uma ferramenta automatizada para a verificação da veracidade dos anúncios de emprego é uma dor real para muitos utilizadores da plataforma, que frequentemente se veem enganados por ofertas fraudulentas. Esta ausência de confiança nos anúncios compromete a integridade da busca por emprego, prejudicando a experiência do utilizador e a credibilidade da própria plataforma.

Com o protótipo desenvolvido, que demonstra como a tecnologia de aprendizagem de máquina pode ser aplicada na detecção de anúncios fraudulentos, foi possível alcançar os objectivos propostos inicialmente. O sistema desenvolvido mostra-se eficaz na melhoria da segurança na navegação dos utilizadores em plataformas de anúncios de emprego, e atende à necessidade de um mecanismo que possa filtrar e alertar sobre fraudes em tempo real. Dessa forma, o objectivo geral deste trabalho foi

alcançado com sucesso, evidenciando o potencial de tecnologias emergentes na resolução de problemas reais em contextos de alta interação online.

8.2 Recomendações

Com a aplicação da solução proposta para a detecção de anúncios de emprego fraudulentos, espera-se que se possa proporcionar um ambiente mais seguro e confiável para os utilizadores do *LinkedIn*, protegendo-os de fraudes e promovendo a transparência nas ofertas de trabalho.

Recomenda-se que o sistema desenvolvido seja amplamente divulgado e integrado em outras plataformas de emprego, para que mais pessoas possam beneficiar desta ferramenta de verificação em tempo real, especialmente em contextos onde os anúncios fraudulentos são uma preocupação crescente.

Além disso, espera-se que, com a evolução e sucesso da solução, sejam implementadas novas funcionalidades que tornem o sistema ainda mais eficiente na detecção de fraudes. A incorporação de algoritmos mais avançados de aprendizagem profunda (*deep learning*) e de técnicas de inteligência artificial poderia melhorar a capacidade do sistema em identificar padrões de fraude mais complexos, aumentando a precisão da análise e ampliando a abrangência da detecção. Também seria benéfico que o sistema fosse adaptado para outras plataformas de redes sociais e sites de empregos, tornando-o uma ferramenta universal de combate a fraudes em anúncios de emprego.

Recomenda-se que sejam incorporados novos *datasets*, que sejam ainda mais recentes e variados, para treinar o modelo com exemplos mais atualizados e representativos de diferentes tipos de fraudes. Aliado a isso, seria igualmente importante desenvolver um mecanismo para que os próprios utilizadores da extensão possam reportar avaliações correctas ou incorrectas, permitindo o refinamento contínuo do modelo por meio de *learning with human feedback*.

Finalmente, recomenda-se que, se realizem melhorias na *interface* tornando a interação com a extensão mais intuitiva e amigável, permitindo que o utilizador

obtenha rapidamente a avaliação do anúncio e aceda a informações complementares sobre a decisão da ferramenta.

Referências bibliográficas

Bibliografia

- [1] **ABAGNALE, F. W.** (2019). *Scam Me If You Can: Simple Strategies to Outsmart Today's Rip-off Artists*. New York: Portfolio/Penguin.
- [2] **AKERLOF, G. A.; SHILLER, R. J.** (2015). *Phishing for Phools: The Economics of Manipulation and Deception*. Princeton: Princeton University Press.
- [3] **ANDERSON, D. J.** (2010). *Kanban: Successful Evolutionary Change for Your Technology Business*. Blue Hole Press.
- [4] **B2B HOUSE.** (2023). Top 2023 LinkedIn statistics for marketers. Disponível em: <https://www.charleagency.com/articles/linkedin-statistics/#:~:text=Surprisingly%2C%20while%20mobile%20usage%20dominates,traffic%2C%20accommodating%20users'%20preferences>. Acesso em: 16 nov. 2024.
- [5] **BECK, K.; BEEDLE, M.** (2001). Manifesto for Agile Software Development.
- [6] **BOGDANOVIC, T.** (2023). Best job search sites in 2023: LinkedIn, Indeed, Glassdoor and more. *Forbes*. Disponível em: <https://www.forbes.com/advisor/business/best-job-search-sites/>. Acesso em: 26 set. 2024.
- [7] **BROWSERSTACK.** (2024). Browser Market Share Analysis for 2024. Disponível em: <https://www.browserstack.com>. Acesso em: 17 nov. 2024.
- [8] **BURGES, C. J. C.** (1998). A tutorial on support vector machines for pattern recognition. *Data Mining and Knowledge Discovery*, 2(2), 121-167.
- [9] **CAMACHO-COLLADOS, J.; PILEHVAR, M. T.** (2018). From word to sense embeddings: A survey on vector representations of meaning. *Journal of Artificial Intelligence Research*, 63, 743-788.

- [10] **CHAKRABARTI, A.; SINGH, P.** (2020). *Advances in Web Technologies*. New York: Springer.
- [11] **COCKBURN, A.** (2006). *Agile Software Development: The Cooperative Game*. Addison-Wesley.
- [12] **CORTES, C.; VAPNIK, V.** (1995). Support-vector networks. *Machine Learning*, 20(3), 273-297.
- [13] **CRISTIANINI, N.; SHAWE-TAYLOR, J.** (2000). *An Introduction to Support Vector Machines and Other Kernel-Based Learning Methods*. Cambridge: Cambridge University Press.
- [14] **FAROOQ, F.** (2021). How LinkedIn became a game-changer for job seekers. *Forbes*. Disponível em: <https://www.forbes.com>. Acesso em: 16 nov. 2024.
- [15] **FEDERAL TRADE COMMISSION – FTC.** (2023). *Consumer Sentinel Network Data Book 2023*. Washington, D.C.: FTC. Disponível em: <https://www.ftc.gov>. Acesso em: 15 maio 2025.
- [16] **FENG, S.; XIE, L.** (2021). Detecting and preventing fraudulent job postings in online platforms. *Computers in Human Behavior*, 108, 106310.
- [17] **FOWLER, M.** (2004). *UML Distilled: A Brief Guide to the Standard Object Modeling Language*. 3. ed. Addison-Wesley.
- [18] **FRANK, A.** (2019). *Scam Me If You Can: Simple Strategies to Outsmart Today's Rip-off Artists*. New York: Portfolio/Penguin.
- [19] **GLOBAL RECRUITER.** (2024). *How many applications per job? Survey reveals UK average*. Disponível em: <https://www.theglobalrecruiter.com>. Acesso em: 02 de Novembro de 2024.
- [20] **GOODFELLOW, I.; BENGIO, Y.; COURVILLE, A.** (2016). *Deep Learning*. MIT Press.

- [21] **GOOGLE DEVELOPERS.** (2024). Chrome Extension Documentation. Disponível em: <https://developer.chrome.com>. Acesso em: 17 nov. 2024.
- [22] **HADNAGY, C.** (2018). *Social Engineering: The Science of Human Hacking*. 2. ed. Indianapolis: Wiley.
- [23] **HIGHSMITH, J.** (2002). *Agile Project Management: Creating Innovative Products*. Addison-Wesley.
- [24] **INDURKHYA, N.; DAMERAU, F. J.** (2010). *Handbook of Natural Language Processing*. 2. ed. Boca Raton: Chapman and Hall/CRC.
- [25] **INSTITUTO NACIONAL DE ESTATÍSTICA – INE.** (2023). *Inquérito ao emprego 2023*. Maputo: INE. Disponível em: <http://www.ine.gov.mz>. Acesso em: 15 maio 2025.
- [26] **INTERNATIONAL LABOUR ORGANIZATION.** (2020). The role of digital platforms in the labour market. Disponível em: <https://www.ilo.org>. Acesso em: 16 nov. 2024.
- [27] **INTUITION SOFTWARE.** (2023). Recruitment 2023: Analysis of applicants on job boards. *Intuition Software*. Disponível em: <https://www.intuition-software.com/recruitment-analysis-2023>. Acesso em: 26 set. 2024.
- [28] **JORNAL NOTÍCIAS.** (2022). *Burlas em nome de empregos preocupam autoridades*. Maputo, 23 nov. 2022. Disponível em: <https://www.jornalnoticias.co.mz>. Acesso em: 15 maio 2025.
- [29] **JURAFSKY, D.; MARTIN, J. H.** (2021). *Speech and Language Processing*. Draft 3rd ed. Stanford University. Disponível em: <https://web.stanford.edu/~jurafsky/slp3/>. Acesso em: 21 abr. 2025.
- [30] **KUMAR, M.** (2021). How Indeed revolutionized global hiring with advanced features. New York: Recruitment Insights.

- [31] **LAKATOS, E. M.; MARCONI, M. A.** (2017). *Fundamentos de metodologia científica*. 7. ed. São Paulo: Atlas.
- [32] **LAUDON, K. C.; LAUDON, J. P.** (2020). *Sistemas de Informação Gerenciais: Gerenciando a Empresa Digital*. 15. ed. São Paulo: Pearson.
- [33] **LEFFER, L.** (2022). Don't Get Scammed by Fake Job Postings on Sites Like LinkedIn and Indeed. *Gizmodo*, 23 dez. 2022. Disponível em: <https://www.gizmodo.com.au>. Acesso em: 1 out. 2024.
- [34] **LINKEDIN.** (2022). Transparency report. *LinkedIn*. Disponível em: <https://linkedin.com/transparency>. Acesso em: 16 nov. 2024.
- [35] **MANNING, C. D.; RAGHAVAN, P.; SCHÜTZE, H.** (2008). *Introduction to Information Retrieval*. Cambridge: Cambridge University Press.
- [36] **McKINSEY & COMPANY.** (2017). *Jobs lost, jobs gained: What the future of work will mean for jobs, skills, and wages*. Disponível em: <https://www.mckinsey.com/featured-insights/future-of-work/jobs-lost-jobs-gained>. Acesso em: 15 maio 2025.
- [37] **MITCHELL, T. M.** (1997). *Machine Learning*. McGraw Hill.
- [38] **MITNICK, K. D.; SIMON, W. L.** (2002). *The Art of Deception: Controlling the Human Element of Security*. Indianapolis: Wiley Publishing, Inc.
- [39] **MOLLY, M.** (2021). How LinkedIn fights fraud: Algorithms and manual review. Disponível em: <https://www.mollymworld.com/linkedin-fraud-detection>. Acesso em: 16 nov. 2024.
- [40] **MOZILLA.** (2024). *Developing Extensions for Firefox*. Disponível em: <https://developer.mozilla.org/en-US/docs/Mozilla/Add-ons/WebExtensions>. Acesso em: 15 maio 2025.
- [41] **MÜLLER, J.; OLIVEIRA, A.** (2020). Combatendo fraudes em anúncios de emprego: Uma análise do impacto das tecnologias de detecção automática.

Proceedings of the 2020 International Conference on Cybersecurity and Privacy Protection.

- [42] **OBERLO.** (2021). *How Many Job Openings Are There in the US?* Disponível em: <https://www.oberlo.com/statistics/number-of-job-openings-us>. Acesso em: 15 maio 2025.
- [43] **PIETERSON, M.** (2019). *Web Development and Design Foundations*. 2^a ed. Wiley.
- [44] **POPPENDIECK, M.; POPPENDIECK, T.** (2003). *Lean Software Development: An Agile Toolkit*. Addison-Wesley.
- [45] **RAJARAMAN, A.; ULLMAN, J. D.** (2011). *Mining of Massive Datasets*. Cambridge: Cambridge University Press.
- [46] **RUMELHART, D. E.; HINTON, G. E.; WILLIAMS, R. J.** (1986). Learning representations by back-propagating errors. *Nature*, 323(6088), 533-536.
- [47] **RUSSELL, S.; NORVIG, P.** (2003). *Artificial Intelligence: A Modern Approach*. 2. ed. Prentice Hall.
- [48] **RUSSELL, S.; NORVIG, P.** (2010). *Artificial Intelligence: A Modern Approach*. 3. ed. Upper Saddle River: Prentice Hall.
- [49] **SCHWABER, K.; BEEDLE, M.** (2002). *Agile Software Development with Scrum*. Prentice Hall.
- [50] **SINGER, P. W.; FRIEDMAN, A.** (2014). *Cybersecurity and Cyberwar: What Everyone Needs to Know*. New York: Oxford University Press.
- [51] **SOCIAL MEDIA TODAY.** (2023). How LinkedIn moderates job postings to reduce fraud. Disponível em: <https://www.socialmediatoday.com/linkedin-job-fraud>. Acesso em: 16 nov. 2024.

- [52] **SOMERS, S.** (2023). Beware of LinkedIn Scams When Searching for Jobs Online. *AARP*. Disponível em: <https://www.aarp.org>. Acesso em: 1 out. 2024.
- [53] **TANENBAUM, A. S.** (2003). *Computer Networks*. 5ª ed. Pearson Education.
- [54] **TECHZIM.** (2022). *Online scams on the rise in Southern Africa: Job seekers beware*. Disponível em: <https://www.techzim.co.zw>. Acesso em: 15 maio 2025.
- [55] **TURBAN, E. et al.** (2018). *Information Technology for Management: On-Demand Strategies for Performance, Growth and Sustainability*. 11. ed. Hoboken: Wiley
- [56] **WEBB, D.** (2017). *Web Design: The Complete Guide to Web Design & Development*.
- [57] **WORLDMETRICS.** (2024). Browser Usage Statistics: Market Data Report 2024. Disponível em: <https://worldmetrics.org>. Acesso em: 17 nov. 2024.
- [58] **YUSOFF, M. et al.** (2022). *Fake job advertisements: A new challenge in e-recruitment*. *International Journal of Business and Society*, v. 23, n. 1, p. 112–128.
- [59] **ZHAO, Y.; WANG, Q.** (2019). The rise of fake job postings: A systematic analysis of fraudulent job advertisements. *Journal of Internet and Digital Media Studies*, 10(4), 1-18.

Outras bibliografias consultadas

- [1] Elias, Richaldo (2019). *Detecção e Classificação de Doenças em Plantas Agrícolas com Recurso a Inteligência Artificial*. Maputo: UEM.
- [2] Macamo, Mónica (2022). *Desenvolvimento de uma plataforma de partilha de processos clínicos baseada em Blockchain; caso de estudo: Sistema Nacional de Saúde*. Maputo: UEM.
- [3] Mahanjane, Erick (2019). *Desenvolvimento de um sistema de apoio ao processo educativo de pessoas com deficiência auditiva; caso de estudo: Escola de Educacao Especial nº 1 da Cidade de Maputo*. Maputo: UEM.

Anexos

Anexo 1: Formulário

De onde é que responde a este inquérito?

33 responses

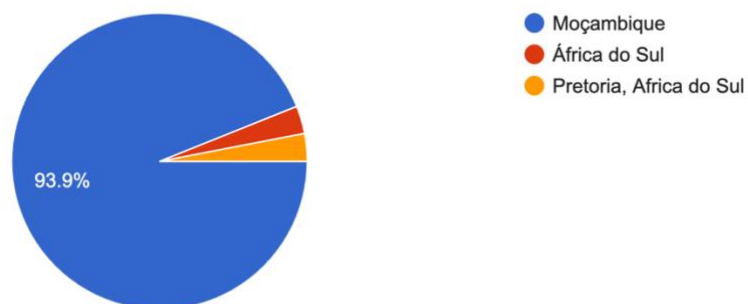


Figura A1 - 1: Formulário - respostas a pergunta 1

Alguma vez acedeu ao LinkedIn para procurar emprego?

33 responses

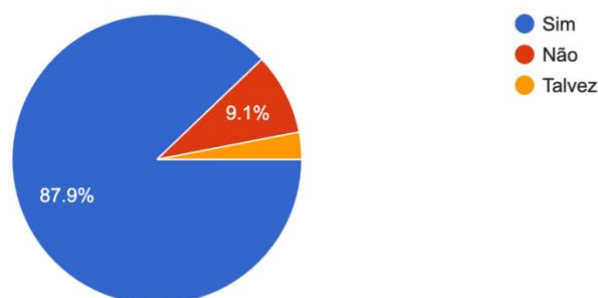


Figura A1 - 2: Formulário - respostas a pergunta 2

Qual é o navegador que você usa com mais frequência para procurar vagas de emprego no LinkedIn?

33 responses

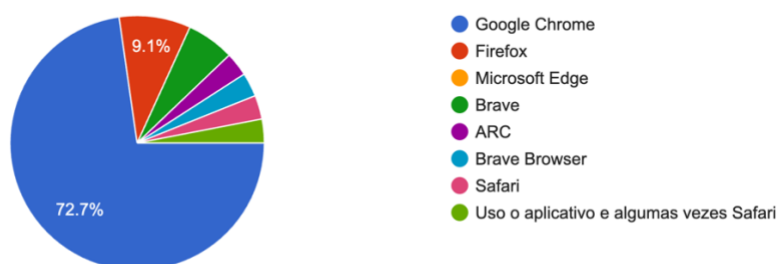


Figura A1 - 3: Formulário - respostas a pergunta 3

Quão bom seria para você ter uma forma "automática" de verificar a credibilidade de anúncios de emprego no LinkedIn?

33 responses

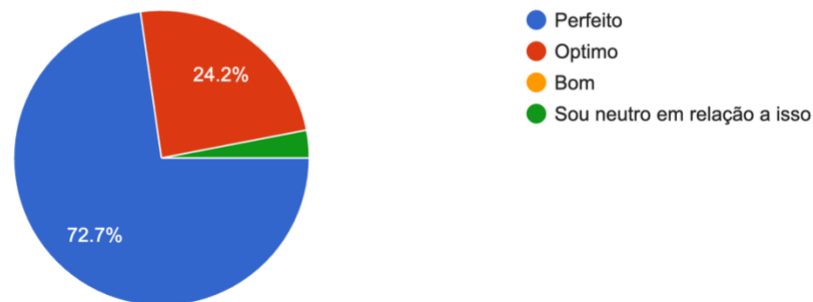


Figura A1 - 4: Formulário - respostas a pergunta 4

De que forma acha que a verificação de credibilidade de anúncios de emprego no LinkedIn deve funcionar?

33 responses

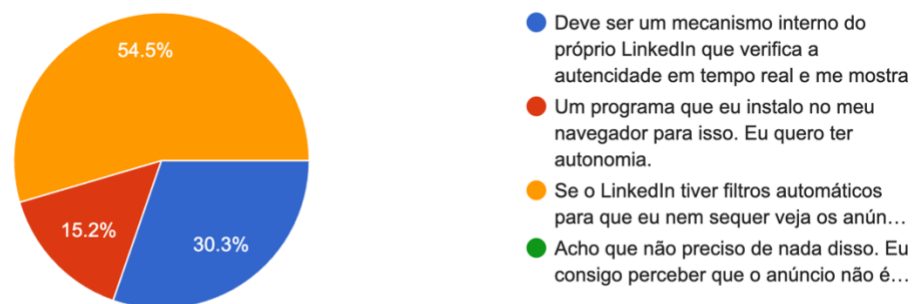


Figura A1 - 5: Formulário - respostas a pergunta 5

Se existisse uma extensão de navegador que ajuda a identificar anúncios de emprego falsos no LinkedIn, você preferiria interagir diretamente com a extensão ou que ela funcione automaticamente?

33 responses

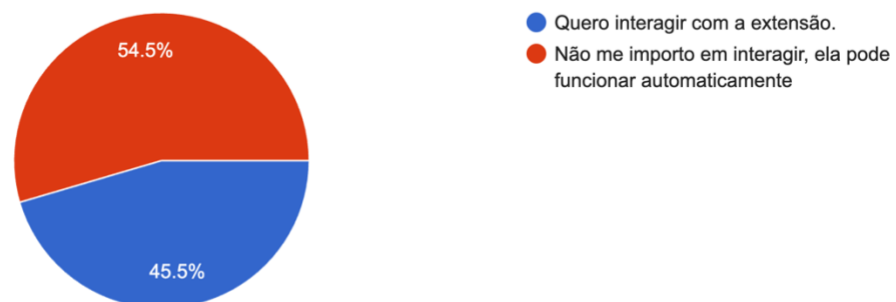


Figura A1 - 6: Formulário - resposta a pergunta 6

Você já teve experiências negativas com anúncios de emprego falsos no LinkedIn?

33 responses

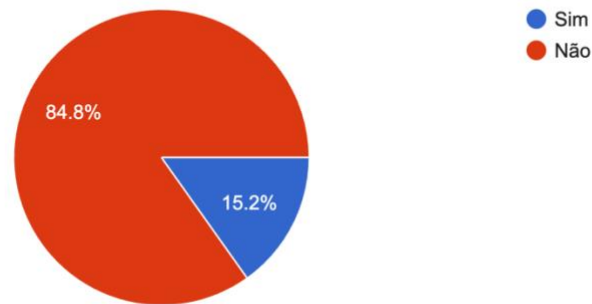


Figura A1 - 7: Formulário - resposta a pergunta 7

Se alguma vez teve uma experiência negativa no LinkedIn, pode contar como foi?

10 responses

N/A

Candidatar-me e nunca ser notificado

Era uma burla que consegui identificar indo para a página oficial da empresa identificada na vaga

Nunca tive

Não.

Ghost jobs

Todos empregos anunciados nunca se quer cheguei a uma entrevista. Ou seja, nem cheguei a saber se os anuncios eram verdadeiros ou nao. Entao posso assumir que eram falsos por causa disso.

Actually, já não me recordo se foi no LinkedIn but mesmo no LinkedIn existem(supostas vagas de emprego) mas no final são pessoas que dizem que tem empresa e que te querem para trabalhar sem pagar e te prometem acções se o projecto prosperar

Figura A1 - 8: Formulário - resposta a pergunta 8