

Inteligência Artificial como Tecnologia para Guerras Híbridas *Artificial Intelligence as a Technology for Hybrid Warfare*

Alessandro Pessoa da Conceição Barreto¹

Resumo

Conflitos híbridos combinam métodos convencionais e não convencionais, incluindo desinformação para influenciar a opinião pública e desestabilizar sociedades (BARBOSA, 2020). Este trabalho desenvolve uma tecnologia inédita no âmbito da defesa nacional, por meio de um MVP (mínimo produto viável) do tipo protótipo, visando à análise do ambiente informacional em conflitos híbridos modernos, especialmente no ciberespaço. O grande volume e a diversidade de dados tornam o ambiente desafiador. O presente trabalho destaca o emprego de requisitos funcionais, como Big Data, *machine learning* e bases de dados distribuídas, estabelecidos pelo Estado-Maior da Armada do Brasil (2018). Uma aplicação web serve como interface, na qual é informada uma URL para *web scraping*. Os dados coletados são armazenados em uma base de dados NoSQL e processados por uma API que utiliza uma rede neural convolucional para classificar as notícias e analisar a polaridade (sentimento) do texto. A pesquisa justifica-se pela necessidade de tecnologias baseadas em IA (inteligência artificial) para enfrentar a desinformação e criar estratégias de defesa eficazes (ESTADO-MAIOR DA ARMADA, 2018). O objetivo é fortalecer a resiliência cognitiva e promover decisões acuradas em contextos complexos. A metodologia inclui o treinamento de uma rede neural convolucional com notícias da Operação GLO (BRASIL, 2023). O MVP provou ser útil na análise do ambiente informacional e na melhoria da consciência situacional, apesar de limitações na acurácia devido ao pequeno volume de dados históricos de outras operações.

Palavras-chave: IA, Big Data, Desinformação, Guerras Híbridas, MVP.

Abstract

Hybrid conflicts combine conventional and unconventional methods, including disinformation to influence public opinion and destabilize societies (BARBOSA, 2024). This paper develops an innovative technology within the realm of national defense through a Minimum Viable Product (MVP) prototype aimed at analyzing the informational environment in modern hybrid conflicts, especially in cyberspace. The vast volume and diversity of data make the environment challenging. This work emphasizes the use of functional requirements, such as Big Data, machine learning, and distributed databases, as outlined by ESTADO-MAIOR DA ARMADA (2018). A web application serves as an interface, where a URL is provided for web scraping. The collected data is stored in a NoSQL database and processed by an API that uses a convolutional neural network to classify news and analyze the text's polarity (sentiment). The research is justified by the need for AI-based technologies to counter disinformation and create effective defense strategies (ESTADO-MAIOR DA ARMADA, 2018). The objective is to strengthen cognitive resilience and promote accurate decision-making in complex contexts. The methodology includes training a convolutional neural network with news from the GLO operation (BRASIL, 2023). The MVP has proven useful in analyzing the informational environment and enhancing situational awareness, despite limitations in accuracy due to the small volume of historical data from other operations.

Keywords: AI, Big Data, Disinformation, Hybrid Warfare, MVP

¹ Especialista em Ciência de Dados e Machine Learning - Faculdade XP Educação IGTI (XPe) - Brasil. Email: alessandro@id.uff.br.

1. INTRODUÇÃO

Daniel Kahneman, vencedor do Prêmio Nobel de Economia em 2002, revela em sua obra *Rápido e Devagar: Duas Formas de Pensar* (KAHNEMAN, 2012) que os humanos tendem a superestimar seu autocontrole e racionalidade, o que afeta a objetividade na tomada de decisões. O autor descreve dois sistemas mentais: o Sistema 1, rápido e intuitivo, e o Sistema 2, mais lento e analítico. A mente frequentemente evita o esforço do Sistema 2, tornando-se suscetível à influência de estereótipos e desinformação, o que pode prejudicar a análise crítica.

Esse fenômeno é particularmente crítico em conflitos híbridos, nos quais a desinformação explora essas vulnerabilidades cognitivas para manipular percepções e comportamentos. Além disso, Vosoughi, Roy e Aral (2018) observaram que notícias falsas alcançam mais pessoas e se disseminam mais rapidamente do que notícias verdadeiras.

Diante da pluralidade e complexidade do ciberespaço, da velocidade de propagação das informações e da natureza dos dados, o MVP tem como objetivo servir como uma ferramenta de análise do ambiente informacional, conforme os requisitos estabelecidos pelo Estado-Maior da Armada (Brasil, 2018).

2. METODOLOGIA

Adotou-se a metodologia *Design Thinking* (BROWN, 2020) para o desenvolvimento de um MVP do tipo protótipo. Dessa forma, o projeto foi estruturado em duas dimensões: no contexto do problema e no contexto da solução.

O contexto do problema contém um subdomínio de entendimento, que tem uma natureza de ideias divergentes, no qual foram compreendidas as regras de negócio e os requisitos funcionais observados no campo da Operação GLO (BRASIL, 2023) e na publicação do Estado-Maior da Armada (Brasil, 2018).

A fase seguinte, e etapa final do contexto do problema, possui a etapa de convergência, o subdomínio de definição, onde foi estabelecido um objetivo SMART (*específico, mensurável, atingível, relevante e temporal*): desenvolver um produto que pudesse ser utilizado em até quatro meses na versão beta, com inteligência artificial para análise do ambiente informacional (Figura 1).



Figura 1 – Canvas do MVP

Fonte: Autor (2024)

A fase de definição intersecta com o contexto da solução, cuja etapa inicial, chamada de desenvolvimento, é o subdomínio de convergência. Nesta fase, foi selecionado um conjunto de possíveis ferramentas e frameworks que atendiam aos requisitos funcionais do objetivo SMART. Para isso, foi criado um ambiente de experimentação com Jupyter Notebook, Docker, VirtualEnv, Python, Streamlit, MongoDB, Mongo Express, TensorFlow e Django.

Após o desenvolvimento, iniciou-se a fase de construção do MVP do tipo protótipo, seguindo uma arquitetura de software orientada à filosofia SOA (Service-Oriented Architecture – Arquitetura Orientada a Serviços). Essa arquitetura foi dividida em quatro aplicações:

1. Três aplicações foram containerizadas usando Docker, com as imagens MongoDB, Mongo Express e a aplicação desenvolvida, nomeada app_01.
2. A quarta aplicação, app_02, foi desenvolvida em um ambiente virtual local utilizando VirtualEnv, isolando-a de outros softwares.

A aplicação app_01 fornece uma interface web construída com o framework Streamlit, do Python. A aplicação app_02 recebe requisições via API, construída com Django, um framework web do Python, e executa rotinas responsáveis pela classificação. Para a classificação, utilizou-se TensorFlow para construir a rede neural convolucional, e a persistência dos dados é feita na base de dados NoSQL MongoDB, que pode operar de forma distribuída e ser gerenciada via web com o Mongo Express.

A explicação e visualização do MVP estão disponíveis no YouTube e GitHub em Pessoa (2024a) e Pessoa (2024b), respectivamente.

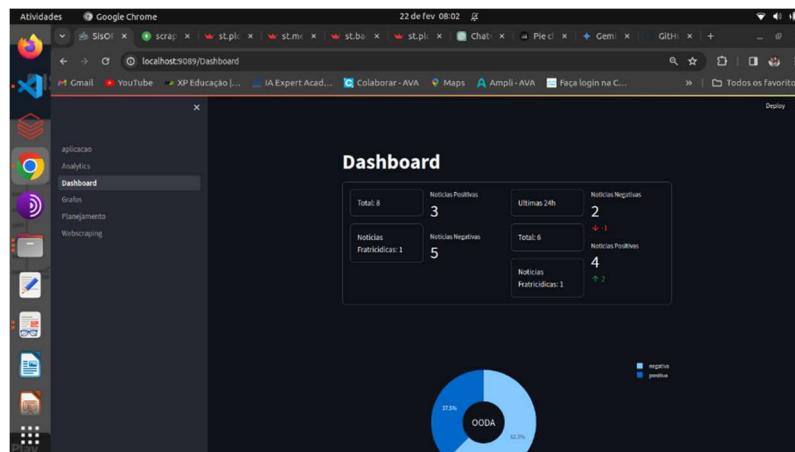


Figura 2 – Printscreen da tela de Dashboard do MVP

Fonte: Autor (2024)

2.1. Etapas do projeto

Para o desenvolvimento deste projeto, utilizou-se a metodologia CRISP-DM (*Cross Industry Standard Process for Data Mining*), que orientou a execução das atividades de forma estruturada e sistemática. O processo foi dividido nas seguintes etapas:

1. Coleta de dados via Web Scraping

Os dados foram obtidos de fontes públicas na internet, e o *web scraping* foi realizado utilizando as bibliotecas *requests* e *BeautifulSoup* do Python. A biblioteca *requests* foi usada para fazer as requisições HTTP às páginas da web, enquanto *BeautifulSoup* foi empregada para a análise e extração dos dados HTML das páginas. É importante ressaltar que, como HTML é um tipo de dado semiestruturado, podem ocorrer mudanças nas metatags das páginas, o que pode exigir a atualização do *XPath* utilizado na extração dos dados.

2. Data Understanding e Data Preparation

Na segunda etapa, que envolve a compreensão e preparação dos dados, a análise exploratória, o pré-processamento e a limpeza dos dados foram realizados no JupyterLab. Neste ambiente de experimentação, foram executados procedimentos fundamentais, incluindo a limpeza dos textos, remoção de *stopwords*, tokenização e criação de *embeddings*. Esses processos foram essenciais para preparar os dados antes da definição e construção das classes finais.

Além disso, a modelagem da rede neural convolucional para classificar a polaridade dos textos como positiva ou negativa também foi realizada no JupyterLab. Todo o processo seguiu o framework CRISP-DM (*Cross Industry Standard Process for Data Mining*), garantindo uma abordagem estruturada e sistemática para a análise e modelagem dos dados.

3. Construção das Classes e Arquitetura Orientada a Serviços

Nesta etapa, o foco foi a construção das classes, seguindo as boas práticas de design de software, com ênfase na redução de acoplamento e na criação de uma arquitetura orientada a serviços (SOA). Os modelos testados com os dados pré-processados da técnica de *web scraping* passaram por uma fase de modelagem orientada a objetos e funcional, com o objetivo de garantir uma estrutura robusta e modular.

Parte da aplicação foi implementada no app 02, responsável pela API desenvolvida com o framework Django. Esta API invoca o classificador baseado em rede neural convolucional para analisar e classificar o texto. Após a classificação, os resultados são armazenados na base de dados MongoDB. O status da operação é então retornado ao usuário via API Django.

Por sua vez, o app 01, desenvolvido com o framework Streamlit e operando em um ambiente virtual, é responsável por renderizar os dados e os resultados retornados pela API do app 02. A integração entre esses dois componentes assegura uma comunicação eficiente e a apresentação interativa e acessível dos resultados.

4. Aplicação de LLM e CNN

Nesta etapa, foi utilizada uma Large Language Model (LLM), especificamente a arquitetura BERT (Bidirectional Encoder Representations from Transformers), para a extração de características relevantes dos dados. O BERT é um modelo baseado em

Transformers, que possui a capacidade de entender o contexto bidirecional das palavras em uma frase, sendo especialmente eficaz para tarefas de Processamento de Linguagem Natural (PLN), como a extração de features.

Após a extração das features com o BERT, a Rede Neural Convolucional (CNN) foi aplicada para classificar os dados rotulados em três categorias: notícias positivas, negativas e neutras. A escolha da CNN para essa tarefa se deve à sua eficiência computacional e capacidade de extrair várias outras features de acordo com os filtros (kernels) utilizados. Em vez de ajustar todos os pesos do modelo BERT para análise de sentimentos, o que exigiria um custo computacional elevado, a CNN foi utilizada para classificar os dados de maneira mais econômica. Essa abordagem permite que o modelo aproveite as poderosas representações extraídas pelo BERT sem a necessidade de um treinamento completo, resultando em um processo de análise de sentimentos mais eficiente e menos dispendioso.

A combinação de BERT e CNN proporcionou um equilíbrio entre a precisão na extração das features e a eficiência no processamento dos dados, otimizando o uso dos recursos computacionais disponíveis, ideal para abordagem exploratória do MVP.

3. RESULTADOS E DISCUSSÃO

O Produto Mínimo Viável (MVP) do tipo protótipo demonstrou ser uma prova de conceito valiosa, desenvolvido com base nos principais requisitos estabelecidos pelo Estado-Maior da Armada (Brasil, 2018). O projeto tem como objetivo mitigar a paralisia da informação causada pelo grande volume de dados, permitindo a manutenção eficiente do ciclo OODA (Observar, Orientar, Decidir e Agir). Dessa forma, contribui para uma consciência situacional mais ágil e precisa, promovendo uma cultura militar orientada por dados.

O produto foi apresentado ao Comando Naval de Operações Especiais (CoNavOpEsp), que decidiu avançar no desenvolvimento do sistema, visando a criação de uma aplicação mais robusta e sofisticada. Em julho de 2024, o projeto encontra-se na fase de redesign da arquitetura de Machine Learning, incorporando novas abordagens e tecnologias emergentes.

A nova fase do projeto inclui testes de integração com Modelos de Linguagem de Grande Escala (LLMs), por meio de uma API baseada no LangChain, com suporte para as APIs do ChatGPT e Gemini. Além disso, estão sendo realizados testes com RAG (Retrieval-Augmented Generation), visando aprimorar a precisão da classificação de notícias.

Esse redesign busca viabilizar a possibilidade de eliminar a dependência de um parque tecnológico próprio, permitindo que os dados classificados sejam utilizados de maneira mais dinâmica e estratégica. A abordagem se baseia na utilização de uma engenharia de prompt pré-determinada, ajustada para diferentes objetivos operacionais, garantindo flexibilidade e eficiência no processamento e análise das Operações de Informação.

Futuramente, os dados classificados serão utilizados para treinar uma rede neural própria, que, inicialmente, será baseada em Transformers. Dependendo dos resultados obtidos,

o treinamento poderá envolver técnicas de fine-tuning com modelos de linguagem pré-treinados (PLMs), como BERTimbau ou modelos LLaMA 2 ou 3.

Nesse contexto, o desenvolvimento do produto segue metodologias ágeis, o que frequentemente resulta em iterações e modificações nas soluções adotadas. No entanto, para manter um pipeline de trabalho consistente, é utilizada a metodologia CRISP-DM (Cross Industry Standard Process for Data Mining) (ROBERTO, 2022).

Uma fase crítica do processo é a escolha do dataset que será utilizado no treinamento ou na recuperação dos modelos de LLMs, seja por meio de RAG (Retrieval-Augmented Generation) ou fine-tuning. O dataset está sendo construído com base em operações reais, como a GLO Lais de Guia, Operação Taquari 1 e Operação Taquari 2, entre outras que ainda serão analisadas.

Como apontado por McKinney (2017), grande parte do trabalho em ciência de dados envolve o pré-processamento e a limpeza dos dados, etapas essenciais para garantir a qualidade dos resultados analíticos. O princípio "Garbage In, Garbage Out" (GIGO), amplamente reconhecido na área, reforça que a qualidade dos resultados obtidos depende diretamente da qualidade dos dados de entrada.

Diante disso, um esforço significativo está sendo dedicado para assegurar a qualidade dos dados. Os datasets utilizados para o treinamento dos modelos estão sendo construídos com base em critérios rigorosos, incluindo o cálculo do tamanho mínimo da amostra para uma população desconhecida, seguindo padrões estatísticos. A fórmula utilizada é representada por:

$$n = \frac{Z^2 \cdot p \cdot (1 - p)}{e^2}$$

Onde:

- **n**: Tamanho mínimo da amostra, ou seja, a quantidade necessária de notícias para obter uma representação precisa das proporções de notícias positivas, negativas e neutras.
- **Z**: Valor crítico para o nível de confiança desejado, obtido da distribuição normal padrão. Para um nível de confiança de 95%, o valor adotado é 1,96.
- **p**: Proporção esperada da população. Como a população é desconhecida e não há uma estimativa prévia da distribuição das notícias em cada categoria (positivas, negativas e neutras), é prudente utilizar $p = 0,5$. Esse valor é considerado conservador, pois maximiza a variabilidade e garante que o tamanho da amostra seja suficientemente grande para refletir a população, assumindo que a probabilidade de um evento ocorrer é igual à probabilidade de não ocorrer.
- **e**: Margem de erro ou erro amostral tolerado. Neste caso, adota-se 5% (0,05), o que significa que as estimativas das proporções de notícias positivas, negativas e neutras podem variar até 5% em relação às proporções reais na população total de notícias extraídas da internet.

A Figura 3 ilustra o dataset em construção, composto por notícias mineradas do ciberespaço da região geográfica brasileira. Essas informações foram coletadas a partir de metatags HTML, elementos utilizados para fornecer metadados sobre as páginas da web, como título, descrição e palavras-chave.

Além disso, os sites foram selecionados com base em seus sitemaps, que são arquivos estruturados contendo as URLs de um site, permitindo que os motores de busca indexem seu conteúdo de maneira mais eficiente. Essa abordagem visa assegurar que o conjunto de dados utilizado para treinamento seja relevante e atualizado, refletindo com precisão as tendências informacionais do ambiente digital brasileiro.

```
datasetTotal = pd.read_csv('/content/drive/MyDrive/datasets/datasetsPolaridadeNoticias_Opinioes/noticiasSemRedesSocial_Duplicado.csv')
```

	titulo	link	snippet	site	data	descricao	imagem	source	nome_operacao	contem_redes_sociais
0	Intervenção federal no Rio de Janeiro poderá r...	https://portal.tcu.gov.br/imprensa/noticias/in...	O Tribunal de Contas da União (TCU) respondeu...	NaN	29 de jun. de 2018	[federal]	NaN	Portal TCU	Intervenção Federal	False
1	Conheça os planos elaborados pela Intervenção ...	http://www.intervencaoefederaln.gov.br/interve...	O Plano Estratégico da Intervenção Federal est...	NaN	3 de set. de 2018	[Intervenção Federal]	NaN	Gabinete de Intervenção Federal no Rio de Janeiro	Intervenção Federal	False
2	Conselhos da República e de Defesa Nacional ap...	https://www.camara.leg.br/noticias/532084-cons...	Os Conselhos da República e de Defesa Nacional...	NaN	19 de fev. de 2018	[Intervenção federal]	NaN	Portal da Câmara dos Deputados	Intervenção Federal	False
3	O que diz a Constituição Federal sobre o insti...	https://www.jusbrasil.com.br/noticias/o-que-diz-a...	O decreto de intervenção, que especificará a a...	NaN	5 de jan. de 2015	[Intervenção]	NaN	Jusbrasil	Intervenção Federal	False
4	Ministro nega MS que pretenda proibir tramita...	https://portal.stf.jus.br/noticias/ver/noticia2...	Segundo Dias Toffoli, o pedido nesse ponto não...	NaN	5 de jul. de 2018	[Intervenção federal]	NaN	Supremo Tribunal Federal	Intervenção Federal	False
1938	Operação Taquari 2 - Comando da 1ª Divisão de ...	https://www.1de.eb.mil.br/galeria-de-videos/top...	A nossa medalhista olímpica, a 3ª Sargento Nat...	NaN	28 de jun. de 2024	NaN	https://i.imgur.com/W/oSoroQSPurUlmqdefault.j...	1ª DE	Operação Taquari	False
1939	Atendimento Telefônico	https://www.rge-rs.com.br/ge/atendimento-tele...	Você pode entrar em contato com a Central de A...	NaN	NaN	NaN	NaN	RGE	Operação Taquari	False
1940	Corpo de Bombeiros Militar do RS	https://www.bombeiros.rs.gov.br/	Site CBMR5: Corpo de Bombeiros Militar do Rio...	NaN	NaN	NaN	NaN	Corpo de Bombeiros Militar do RS	Operação Taquari	False
1941	Amaggi Home	https://www.amaggi.com.br/	Logística e Operações. Administração de portos.	NaN	NaN	[Operações]	NaN	Amaggi	Operação Taquari	False
1942	Trabalha Como: Juntos criamos novos caminhos	https://www.grupoccr.com.br/grupo-ccr/trabalhe...	Naquela época eu tinha objetivos e sonhos de c...	NaN	NaN	[Operação: Operação]	NaN	Grupo CCR	Operação Taquari	False

1943 rows x 10 columns

Figura 3 – Printscreen de Jupyter notebook com notícias do dataset.

Fonte: Autor (2024)

A Figura 4 apresenta o dataset, um dataframe que contém todas as notícias, além de novas colunas que servirão como features para os modelos de sumarização, classificação e análise de dados. As novas features incluem: qtdPalavrasTexto, texto e vídeo. O texto será sumarizado caso a quantidade de palavras ultrapasse um limiar preestabelecido e será submetido à fase de classificação apenas se atender a uma quantidade mínima de palavras, também definida por um limiar específico.

Testes serão realizados para determinar o ponto de corte ideal para esses limiares, com o objetivo de otimizar os resultados. É importante destacar que, em processamento de linguagem natural (NLP), as métricas de avaliação frequentemente lidam com a complexidade da semântica textual, que não se comporta de maneira linear ou direta, como ocorre em outras abordagens de inteligência artificial. A extração de significado em textos é uma tarefa complexa e requer técnicas que vão além das métricas tradicionais utilizadas em dados numéricos.

	titulo	link	snippet	site	data	descricao	imagem	source	nome_operacao	contem_redes_sociais
0	Intervenção federal no Rio de Janeiro poderá r...	https://portal.tcu.gov.br/imprensa/noticias/n...	O Tribunal de Contas da União (TCU) respondeu ...	NaN	29 de jun. de 2018	[Federal]	NaN	Portal TCU	Intervenção Federal	False
1	Conheça os planos elaborados pela intervenção...	http://www.intervencao-federal.gov.br/interve...	O Plano Estratégico da Intervenção Federal est...	NaN	3 de set. de 2018	[Intervenção Federal]	NaN	Cabinete de Intervenção Federal no Rio de Janeiro	Intervenção Federal	False
2	Conselhos da República e de Defesa Nacional ap...	https://www.camara.leg.br/noticias/532084-cons...	Os Conselhos da República e de Defesa Nacional...	NaN	19 de fev. de 2018	[Intervenção Federal]	NaN	Portal da Câmara dos Deputados	Intervenção Federal	False
3	O que diz a Constituição Federal sobre o insti...	https://www.jusbrasil.com.br/noticias/o-que-d...	O decreto de intervenção, que especifica a a...	NaN	5 de jan. de 2015	[Intervenção]	NaN	Jusbrasil	Intervenção Federal	False
4	Ministro nega MS que pretendia proibir trami...	https://portal.stf.jus.br/noticias/verNoticiaD...	Segundo Dias Toffi, o pedido nesse ponto não...	NaN	5 de jul. de 2018	[Intervenção Federal]	NaN	Supremo Tribunal Federal	Intervenção Federal	False
1938	Operação Taquari 2 - Comando da 1ª Divisão de ...	https://www.1de.ri.mil.br/galeria-de-videos/op...	A nossa medalhista olímpica, a 3ª Sargento Nat...	NaN	28 de jun. de 2024	NaN	https://i.yimg.com/W/e/Source/SP/Unit/default.j...	1ª DE	Operação Taquari	False
1939	Atendimento Telefônico	https://www.rge-rs.com.br/ge/atendimento-tele...	Você pode entrar em contato com a Central de A...	NaN	NaN	NaN	NaN	RGE	Operação Taquari	False
1940	Corpo de Bombeiros Militar do RS	https://www.bombeiros.rs.gov.br/	Site CBMRS: Corpo de Bombeiros Militar do Rio ...	NaN	NaN	NaN	NaN	Corpo de Bombeiros Militar do RS	Operação Taquari	False
1941	Amaggi Home	https://www.amaggi.com.br/	Logística e Operações. Administração de portos...	NaN	NaN	[Operações]	NaN	Amaggi	Operação Taquari	False
1942	Trabalhe Conosco: Juntos criamos novos caminhos	https://www.grupoccr.com.br/grupo-ccr/trabalhe...	Naquela época eu tinha objetivos e sonhos de c...	NaN	NaN	[Operação/Operação]	NaN	Grupo CCR	Operação Taquari	False

Figura 4 – Printscreen de Jupyter notebook com notícias do dataset.

Fonte: Autor (2024)

Paralelamente à construção do dataset, foi empregada a técnica de sumarização de texto, uma área fundamental dentro da Geração de Linguagem Natural (NLG). Essa abordagem se torna especialmente relevante porque os Modelos de Linguagem de Grande Escala (LLMs), como ChatGPT, Gemini e outros, apresentam limitações no número de *tokens* que podem processar em suas entradas.

Na sumarização de textos, é possível adotar duas abordagens principais: sumarização extrativa e sumarização abstrativa. Para mitigar problemas relacionados a vieses algorítmicos, optou-se pela sumarização extrativa, que consiste na seleção direta das sentenças mais relevantes do texto, de acordo com o modelo matemático subjacente a cada técnica.

Entre as abordagens disponíveis, escolheu-se a técnica de Latent Semantic Analysis (LSA), ou Análise Semântica Latente, priorizando critérios de explicabilidade e efetividade. A LSA utiliza a Decomposição de Valor Singular (SVD), uma técnica de álgebra linear que facilita a compreensão do modelo, pois permite decompor a matriz de dados em três componentes: uma matriz de termos, uma matriz diagonal de valores singulares e uma matriz de documentos.

A redução de dimensionalidade ocorre ao manter apenas os principais componentes (os valores singulares mais significativos), descartando aqueles de menor importância. Esse processo auxilia na captura das relações semânticas mais relevantes. Ao transformar os dados em um espaço semântico latente, a LSA representa os documentos e os termos em um espaço vetorial de menor dimensão. Nesse novo espaço, palavras com significados semelhantes estão mais próximas umas das outras, facilitando a identificação das sentenças mais significativas para a sumarização.

Como mostrado na Figura 5, as áreas destacadas em amarelo representam as partes do texto selecionadas pelo processo de sumarização extrativa.



Figura 5 – Resumo extrativo utilizando técnica LSA e markdown para realçar o resumo.

Fonte: Autor (2024)

Outras técnicas de sumarização, como TextRank, KLSummarizer e Luhn, estão sendo testadas e serão avaliadas por meio de métricas amplamente utilizadas em Processamento de Linguagem Natural (NLP), tais como ROUGE, BLEU e similaridade do cosseno. A interface do usuário, ilustrada na Figura 6, exibe o campo “Conteúdo da Notícia”, previamente preenchido com o texto extraído por web scraping. Esse conteúdo pode ser submetido à classificação por inteligência artificial por meio do botão “Classificar com IA”. Ao acioná-lo, uma requisição é enviada via API RESTful para uma aplicação Django, que processa a classificação da notícia, retorna sua polaridade e preenche automaticamente o campo “Polaridade da Notícia - IA”.

Para a classificação de polaridade, está sendo utilizado o modelo bert-base-multilingual-uncased-sentiment (NLPTOWN), uma LLM baseada na arquitetura BERT (Bidirectional Encoder Representations from Transformers) em paralelo a rede neural convolucional com embedding do BERT.

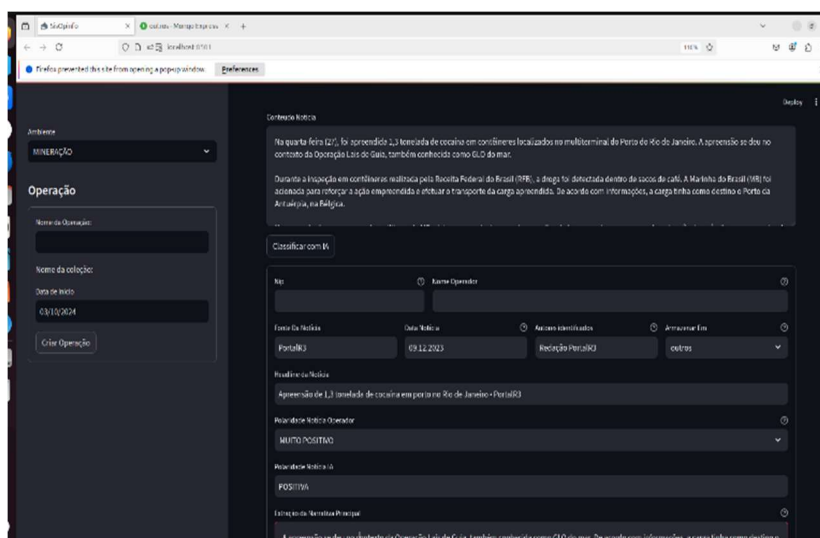


Figura 6 – Interface de Operação do usuário

Fonte: Autor (2024)

O modelo da NLPTOWN modelo foi ajustado para realizar classificação de sentimentos em diversas línguas. No entanto, o português não estava contemplado entre os idiomas para os quais o modelo foi treinado. Para contornar essa limitação, o fluxo de processamento inclui a

tradução automática do texto sumarizado para o inglês, utilizando a API do Google no Python. Após a tradução, o texto em inglês é submetido à classificação de sentimentos, uma vez que esse idioma está entre aqueles para os quais o modelo foi otimizado.

O fine-tuning do modelo foi realizado com o objetivo de adaptá-lo à tarefa de análise de sentimentos em múltiplos idiomas, utilizando um conjunto de dados específico com avaliações rotuladas. Durante esse processo, as camadas e pesos pré-treinados do BERT foram preservados, enquanto novas camadas foram ajustadas para a tarefa específica de classificação de sentimentos, garantindo que o modelo fosse capaz de associar textos às suas respectivas polaridades.

Além disso, estão sendo conduzidos testes com as boas práticas de visualização de dados (CNA, 2016) na interface de análise do ambiente informacional, denominada DASHBOARD. Essa interface encontra-se em fase de aperfeiçoamento e validação, conforme ilustrado na Figura 7. Os gráficos disponibilizados oferecem suporte às funcionalidades de drill down e drill up, permitindo uma exploração mais aprofundada e flexível dos dados.

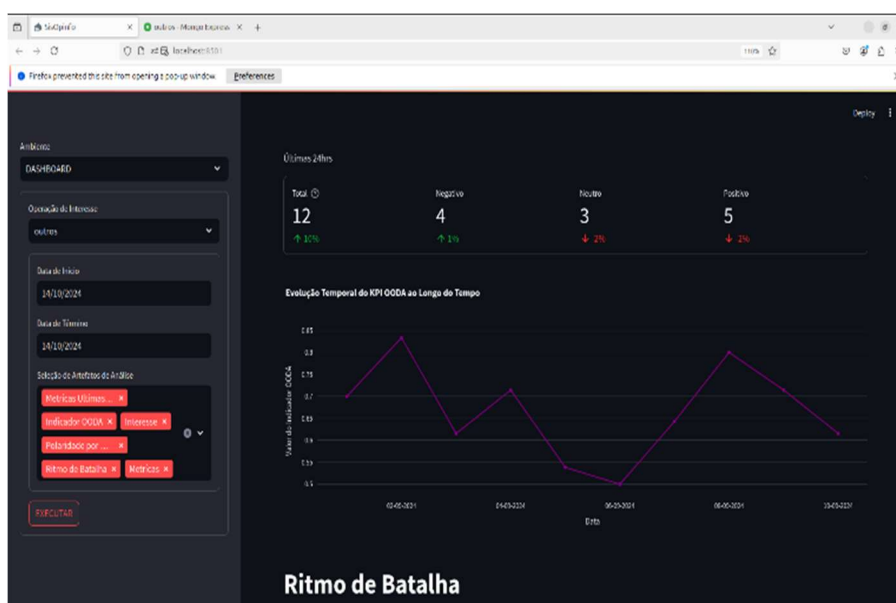


Figura 7 – Tela de DASHBOARD e métricas da operação

Fonte: Autor (2024)

Nesta interface, diversos tipos de gráficos estão disponíveis para análise do ambiente informacional, como demonstrado nas "Métricas das Últimas 24 Horas" (Figura 07). Esses gráficos evidenciam o volume de notícias conforme a polaridade, destacando a capacidade analítica do ambiente informacional.

Uma ferramenta valiosa é o “Indicador OODA”, que possibilita a análise do KPI OODA ao longo do tempo. Esse indicador foi criado para mensurar a frequência das atividades, fundamentando-se no princípio de que, assim como em qualquer operação militar, é essencial que o Estado-Maior mantenha seu ciclo OODA (Observar-Orientar-Decidir-Agir) em funcionamento mais rápido que o do adversário (ESTADO-MAIOR DA ARMADA, 2018).

Outros gráficos, como os ilustrados nas Figuras 8 e 9, podem ser utilizados para obter insights sobre quais mídias publicam mais notícias com polaridade positiva, negativa ou neutra, fornecendo dados valiosos para a elaboração de uma matriz SWOT, por exemplo.

Além disso, estão sendo conduzidos testes para identificar quais sentenças estão mais presentes nas notícias, com o objetivo de compreender as possíveis narrativas utilizadas e desenvolver contramedidas eficazes para operações psicológicas e de comunicação social.



Figura 8 – Gráfico representando média x polaridade x quantidade de notícias
Fonte: Autor (2024)

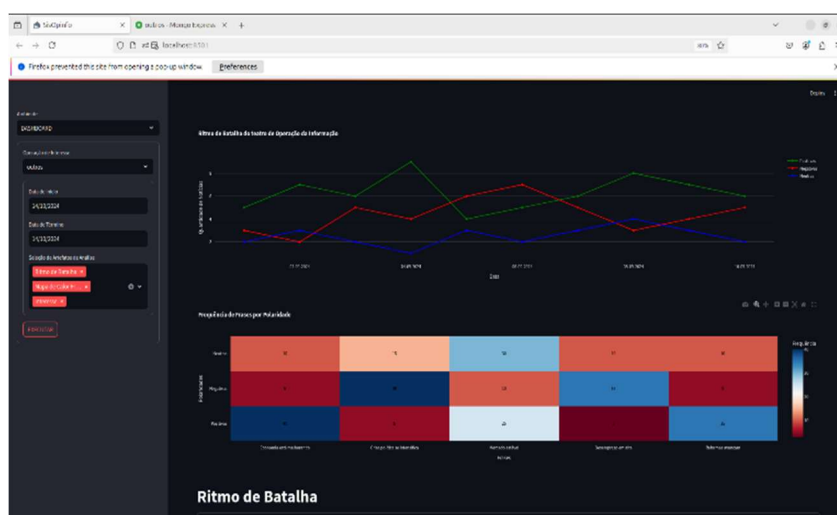


Figura 9 – Gráfico com o ritmo de batalha conforme evolução temporal da operação e mapa de calor com sentenças mais utilizadas nas notícias conforme polaridade
Fonte: Autor (2024)

4. CONSIDERAÇÕES FINAIS

O projeto proporciona ganhos significativos para a gestão do conhecimento, ao utilizar tecnologias de Big Data, armazenamento de dados estruturados e não estruturados e inteligência artificial. Essas tecnologias operacionalizam a pirâmide DIKW (IBM, 2024) e permitem auditoria e governança de dados, possibilitando uma rápida consciência informacional do ciberespaço e favorecendo uma tomada de decisão mais assertiva, baseada em dados históricos.

Além disso, a solução garante a possibilidade de auditoria, pois armazena e documenta todos os dados e raciais utilizados. Dessa forma, assegura-se que as decisões são informadas e rastreáveis, otimizando a gestão da informação e a resposta a situações complexas.

O MVP (Produto Mínimo Viável), do tipo protótipo, demonstra ser uma prova de conceito eficaz para mitigar a paralisia da informação, utilizando o ciclo OODA para aprimorar a consciência situacional no ambiente informacional militar. O sistema desenvolvido encontra-se em fase de testes, aplicando modelos de Linguagem de Grande Escala (LLMs), com técnicas como fine-tuning e RAG (Retrieval-Augmented Generation) para a classificação de notícias.

Os resultados preliminares indicam que a abordagem de sumarização extrativa via LSA (Latent Semantic Analysis), com redução de dimensionalidade para entrada na LLM, ajuda a evitar vieses inerentes à sumarização abstrativa. A tradução para o inglês permitiu o uso do modelo bert-base-multilingual-uncased-sentiment na classificação de sentimentos, embora melhorias no treinamento específico para o português ainda sejam necessárias.

O dashboard interativo demonstrou ser uma ferramenta valiosa para análise de dados, fornecendo métricas como o KPI OODA e insights sobre narrativas midiáticas.

Como contribuições e Limitações pode-se destacar:

- Na perspectiva teórica, este estudo contribui para as áreas de ciência de dados e gestão do conhecimento, aplicadas ao contexto militar. Ele demonstra como a integração de modelos de machine learning e LLMs pode aprimorar a análise de dados em tempo real.
- A utilização de técnicas de sumarização, como a LSA, e de classificação de sentimentos via modelos pré-treinados, oferece uma abordagem inovadora para lidar com o volume crescente de dados no ciberespaço.
- Na perspectiva pragmática, o uso de modelos de NLP (Processamento de Linguagem Natural) alinhados ao ciclo OODA evidencia como tecnologias emergentes podem ser integradas a práticas militares tradicionais, o que abre espaço para pesquisas interdisciplinares futuras.

O protótipo demonstra ser uma solução viável para a gestão de grandes volumes de dados em operações militares, com aplicações práticas em operações psicológicas e de comunicação social. Além disso, o sistema proporciona capacidade de auditoria, essencial para revisões e justificativas das decisões tomadas.

No entanto, algumas limitações do projeto devem ser destacadas:

1. Dependência de um modelo pré-treinado em inglês para classificação de sentimentos, o que adiciona uma etapa de tradução e pode introduzir ruídos nos resultados;
2. A sumarização extrativa pode não capturar nuances textuais que técnicas abstrativas conseguiriam.

Diante disso, sugere-se que pesquisas futuras explorem: (i) o treinamento de redes neurais baseadas em Transformers diretamente em português, eliminando a necessidade de tradução (exemplo: BERTimbau); (ii) o desenvolvimento de uma base de dados própria e robusta para operações militares específicas; e (iii) a inclusão de novas dimensões de análise, como sentimentos em mídias visuais e de áudio.

5. REFERÊNCIAS

BARBOSA, A. H. B. Tese apresentada como requisito parcial para a conclusão do Curso de Política e Estratégia Marítimas da Escola de Guerra Naval. 2024. Disponível em: <https://www.marinha.mil.br/egn/sites/www.marinha.mil.br/egn/files/C-PEM010%20-%20CMG%20%28FN%29%20ALEXANDRE%20HENRIQUE%20BATISTA%20BARBOSA%20%20A%20DESINFORMA%C3%87%C3%83O%20COMO%20FERRAMENTA%20DA%20GUERRA%20H%C3%84BRIDA.pdf>. Acesso em: 12 jul. 2024, às 12h50.

BRASIL. Decreto nº 11.765, de 15 de novembro de 2023. Autoriza o emprego das Forças Armadas para a Garantia da Lei e da Ordem em portos e aeroportos. Brasília, DF: [s.n.], 2023.

BRASIL. Estado-Maior da Armada. EMA-335. Doutrina de Operações de Informação. Brasília: EMA, 2018.

BROWN, T. Design Thinking: uma metodologia poderosa para decretar o fim das velhas ideias. Edição comemorativa 10 anos. Rio de Janeiro: Alta Books, 2020.

CNA, C. N. Storytelling com dados: um guia sobre visualização de dados para profissionais de negócios. 1. ed.. ed. Rio de Janeiro: Editora Alta Books, 2016.

IBM. From data to knowledge: knowledge management with AI. 2024. Disponível em: <https://developer.ibm.com/articles/ba-data-becomes-knowledge-1/>. Acesso em: 14 out. 2024.

NLPTOWN. *bert-base-multilingual-uncased-sentiment*. Disponível em: <https://huggingface.co/nlptown/bert-base-multilingual-uncased-sentiment>. Acesso em: 14 out. 2024.

PESSOA, A. *Pitch MBA Ciência de Dados - Sistema Operação da Informação Aplicativo*. 2024a. Disponível em: <https://youtu.be/yyDdlp2uo4M?si=RSwlcYhEi-fHAXGT>. Acesso em: 3 ago. 2024.

PESSOA, A. *Projeto Aplicado em Sistemas Operacionais e Informacionais* [repositório GitHub]. 2024b. Disponível em: <https://github.com/alessandropessoa/projetoAplicadoSisOpInfo>. Acesso em: 3 ago. 2024.

ROBERTO, C. *CRISP-DM: as 6 etapas da metodologia do futuro*. Disponível em: <https://blog.mbauspesalq.com/2022/04/12/crisp-dm-as-6-etapas-da-metodologia-do-futuro/>. Acesso em: 3 ago. 2024.