

IFPI - CAMPUS TERESINA CENTRAL

**VERIFICA AI: BOT QUE USA INTELIGÊNCIA ARTIFICIAL PARA DETECÇÃO
DE FAKE NEWS NAS REDES SOCIAIS**

Teresina, PI

2025



Aquilis Alves de Melo Oliveira
Isabella dos Santos Caruso
Vítor Emanuel da Silva Rodrigues

Victor Eduardo Alves da Silva Carvalho
Abraão Lima Sousa

VERIFICA AI: BOT QUE USA INTELIGÊNCIA ARTIFICIAL PARA DETECÇÃO DE FAKE NEWS NAS REDES SOCIAIS

Relatório apresentado à 9ª FEMIC - Feira
Mineira de Iniciação Científica.

Orientação do Prof. Abraão Lima Sousa e
coorientação de Victor Eduardo Alves da Silva
Carvalho.

Teresina, PI

2025



RESUMO

A desinformação nas redes sociais representa uma ameaça global, intensificada pela inteligência artificial (IA), projetando-se como a principal preocupação para 2027. No Brasil, 57% da população não consegue identificar conteúdos falsos, conforme dados da OCDE. Diante dessa problemática, o estudo desenvolveu o *Verifica AI*, um agente de IA integrado ao Instagram. Sua função é verificar a veracidade de posts multimodais (texto, imagem e vídeo) diretamente na plataforma e classificar corretamente a desinformação com base nas sete categorias da UNESCO. O protótipo utiliza a *Graph API* da Meta para comunicação via Mensagem Direta (DM), o *Instaloader* para download do conteúdo, e o *Gemini Flash 2.5* para OCR e transcrição de áudio para padronização dos dados. Em seguida, o *Gemini* analisa o material por meio de *prompts* otimizados, comparando-o com fontes online para validação. O protótipo foi concluído com sucesso e está disponível no perfil do Instagram @verifica_ai_. Os resultados da matriz de confusão, obtidos com 200 postagens (150 verdadeiras e 50 falsas), indicaram acurácia de 97%, precisão de 98% e sensibilidade de 98%. Isso permite concluir que o sistema classifica corretamente a maioria dos casos. O protótipo alcançou todos os objetivos, com desempenho funcional robusto que supera as limitações de ferramentas como *Meta AI* e *AI Studio*, ao analisar posts de forma nativa. Em contrapartida, é evidente as limitações do uso de tecnologias gratuitas, como as restrições operacionais (limites diários de uso), a latência ocasional nas respostas e a ocorrência, ainda que rara, de falhas críticas, o que inviabiliza a escalabilidade da solução.

Palavras-chave: fake-news, desinformação, inteligência artificial.



SUMÁRIO

1 INTRODUÇÃO	5
2 JUSTIFICATIVA	6
3 OBJETIVO GERAL	7
4 METODOLOGIA	8
5 RESULTADOS E DISCUSSÕES	15
6 CONCLUSÕES OU CONSIDERAÇÕES FINAIS	22
REFERÊNCIAS	23
ANEXO	24

O surgimento da internet e a evolução dos smartphones geraram grandes impactos no mundo moderno. Como resultado, vivemos na era digital, onde a maioria das pessoas está conectada por meio de celulares e tem amplo acesso às redes sociais. Segundo o relatório da “We Are Social”, há aproximadamente 5,45 bilhões de usuários de internet no mundo em 2024, representando 67,1% da população global (WE ARE SOCIAL, 2024). No entanto, junto com essa conectividade, surge uma ameaça à sociedade: as *fake news*.

Segundo a UNESCO, “fake news” consiste na divulgação de informações falsas, uma mentira intencional usada para confundir ou manipular pessoas. Esse tipo de conteúdo é considerado altamente perigoso, pois geralmente é elaborado com estratégias persuasivas e possui alto potencial de impacto. Apesar do uso popular do termo *fake news*, a UNESCO recomenda substituí-lo por “desinformação” ou “informação incorreta”, uma vez que o termo original tem sido utilizado para desacreditar o jornalismo (UNESCO, 2018).

A desinformação é disseminada principalmente em ambientes virtuais, sobretudo nas redes sociais, onde conteúdos polêmicos tendem a viralizar rapidamente (DELMAZO, 2018). Situações como a pandemia e os períodos eleitorais demonstram que a desinformação distorce a realidade, intensifica a polarização, provoca pânico e leva a sociedade a decisões equivocadas que afetam diretamente a saúde, a segurança, as finanças e até mesmo a democracia.



2 JUSTIFICATIVA

O cenário atual torna-se ainda mais preocupante com o uso mal-intencionado de inteligência artificial, capaz de gerar conteúdos falsos, inclusive deepfakes de áudio e vídeo. O Global Risks Report 2025, do World Economic Forum, aponta a desinformação como o principal risco global até 2027, ressaltando a dificuldade crescente em identificar conteúdos produzidos por inteligência artificial (IA) (WORLD ECONOMIC FORUM, 2025). No Brasil, o problema é ainda mais alarmante: segundo a Organização para a Cooperação e Desenvolvimento Econômico (OCDE), 57% da população não sabe identificar conteúdos falsos, tornando o país o mais suscetível à desinformação no mundo (OCDE, 2024).

Diante desse contexto, torna-se indispensável o desenvolvimento de soluções tecnológicas que atuem ativamente no combate à desinformação. Assim, o objetivo deste estudo é propor uma ferramenta capaz de verificar a veracidade de conteúdos como posts e mensagens diretamente no ambiente das redes sociais. Além disso, a solução busca ser didática, confiável e ágil, fornecendo ao usuário conclusões acompanhadas de suas respectivas fontes, contribuindo ativamente no combate de desinformação nas redes sociais.



3 OBJETIVOS

3.1 OBJETIVO GERAL

Desenvolver uma ferramenta que combata ativamente a desinformação ao verificar, dentro da própria rede social, a veracidade de conteúdos publicados (posts e mensagens). A solução deve ser didática, confiável e ágil, fornecendo ao usuário uma conclusão acompanhada de fontes.

3.2 OBJETIVOS ESPECÍFICOS

- Selecionar a rede social alvo e configurar as permissões oficiais necessárias para envio e recebimento de mensagens diretas (DM);
- Desenvolver o *backend* para processar conteúdos multimodais, convertendo áudio, vídeo e imagem em texto para análise;
- Integrar LLM do *Gemini* ao algoritmo, habilitando a análise com inteligência artificial;
- Projetar e ajustar *prompts* com finalidade de identificar e classificar desinformação;
- Adotar as diretrizes da UNESCO para classificar dentre as sete categorias;
- Avaliar a confiabilidade e desempenho do protótipo através da matriz de confusão;

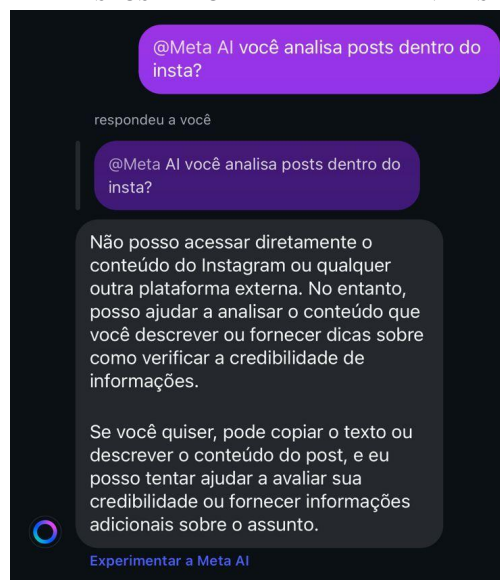
4 METODOLOGIA

O Instagram foi a plataforma escolhida para o desenvolvimento do protótipo. O problema central abordado foi a criação de um método eficaz para que os usuários possam verificar a veracidade de uma postagem diretamente na rede social, mantendo a experiência integrada à plataforma.

4.1 ANÁLISE DAS SOLUÇÕES EXISTENTES

Entre as soluções já existentes, destaca-se o *chatbot Meta AI*, oferecido pela própria Meta dentro do Instagram. Sua vantagem reside na integração nativa, que permite aos usuários interagirem diretamente com a IA, sem a necessidade de alternar entre aplicativos. No entanto, essa solução apresenta limitações, como a incapacidade de analisar imagens enviadas e de verificar diretamente as postagens do Instagram. O usuário precisa, de forma manual, transcrever o conteúdo de um post para conferir a veracidade.

IMAGEM 1 - RESPOSTA DO META AI PARA ANÁLISE DE POST



Fonte: Próprio autor, adaptado de Instagram (2025).

Uma nova funcionalidade do Instagram, que permite a criação de um agente de IA nativo por meio do *AI Studio*, foi testada. No entanto, os testes revelaram que o modelo não possui a capacidade de analisar conteúdos em postagens, além de apresentar alucinações e fornecer respostas incorretas.



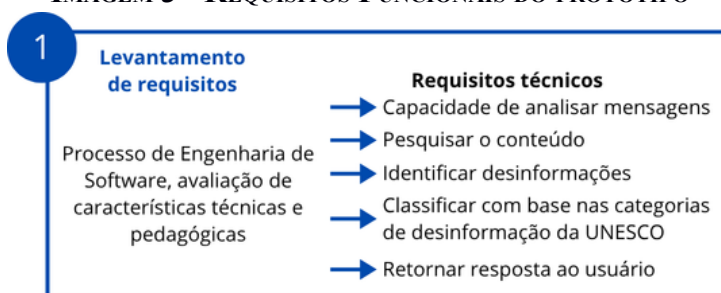
A solução mais completa encontrada é o Grok da rede social X, capaz de participar de conversas dentro da própria plataforma. Quando solicitado, o Grok analisa posts (imagem, texto e vídeo) e avalia a veracidade do conteúdo publicado. A resposta é enviada diretamente dentro do X, dando praticidade ao usuário. Inspirado nessa abordagem, os requisitos funcionais e não funcionais do projeto foram definidos.

IMAGEM 2 - RESPOSTA DO GROK NO X PARA ANÁLISE DE POST



Fonte: Próprio Autor (2025).

IMAGEM 3 - REQUISITOS FUNCIONAIS DO PROTÓTIPO



Fonte: Próprio Autor (2025).

4.2 DEFINIÇÃO DA SOLUÇÃO

Com base na análise, buscou-se desenvolver uma ferramenta que adota uma abordagem semelhante à do Grok, mas adaptada para o Instagram. A solução proposta consiste em um agente de IA integrado a um perfil da rede social.



Um agente de IA é um sistema de software que utiliza inteligência artificial para atingir objetivos e completar tarefas. Ele é capaz de demonstrar raciocínio, planejamento, memória e um nível de autonomia para tomar decisões, aprender e se adaptar. Além disso, eles processam informações multimodais, como texto, voz e vídeo, e podem executar fluxos de trabalho complexos. (GOOGLE CLOUD, 2025).,

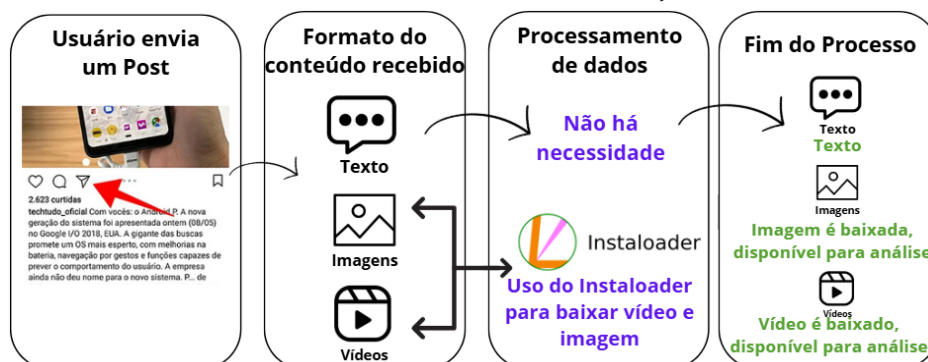
Ao se desenvolver um agente de IA, obtém-se uma solução que permite o recebimento de conteúdo diretamente por DM, o processamento de material multimodal e o retorno de uma conclusão com as respectivas fontes. Assim, o requisito central do problema é atendido, garantindo a verificação de posts dentro da rede social e a preservação da experiência do usuário.

4.3 DESENVOLVIMENTO DO PROTÓTIPO

A integração com o Instagram foi estabelecida por meio das ferramentas da *Meta for Developers*, possibilitando a recepção e o envio de mensagens diretas (DM) entre o perfil do projeto e os usuários. Após a configuração da *Graph API* com as permissões necessárias, um endpoint seguro (*webhook* HTTPS) foi implementado para receber, em tempo real, o conteúdo compartilhado pelos usuários.

Em seguida, desenvolveu o *backend*, responsável por extrair e organizar o material recebido. Para baixar imagens e vídeos dos posts, empregou-se a ferramenta *Instaloader*, viabilizando o download do conteúdo para análise.

IMAGEM 4 - FLUXO DO PROCESSO DE EXTRAÇÃO DE DADOS



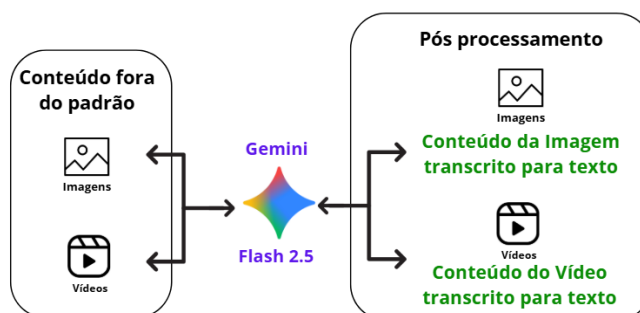
Fonte: Próprio Autor (2025).

Nesta etapa, mesmo após o processamento inicial, três tipos de conteúdo (texto, imagem e vídeo) foram identificados. O passo seguinte consistiu em padronizar todo o



material para o formato de texto. Para transcrever imagens e vídeos, foi empregada a visão computacional. As imagens foram processadas por meio de OCR (Reconhecimento Óptico de Caracteres) e o áudio dos vídeos foi transcrito, permitindo a padronização dos dados. O modelo escolhido para ser o cérebro do sistema foi *Gemini Flash 2.5*, devido à sua disponibilidade gratuita, o que assegurou a viabilidade econômica do projeto ao garantir custo zero para o desenvolvimento e para o uso pelo usuário final.

IMAGEM 5 - FLUXO DO PROCESSO DE PADRONIZAÇÃO DE DADOS



Fonte: Próprio Autor (2025).

Testes preliminares foram conduzidos para verificar a integração do *Gemini* com o perfil do Instagram. A etapa de integração do *Gemini Flash 2.5* foi confirmada como bem-sucedida por meio de uma das respostas obtidas, apresentada na Imagem 7. O algoritmo desenvolvido permite a definição de diretrizes para o *chatbot*, incluindo a forma como ele deve identificar e classificar desinformação em publicações, com o uso de *prompts* específicos para este fim.

IMAGEM 6 - PROMPT USADO NO GEMINI

```
1. CONCLUSÃO PRINCIPAL:
Diga claramente uma das opções:
- "Informação verdadeira"
- "Não verdadeira"
- "NÃO É POSSÍVEL CONCLUIR"

2. CLASSIFICAÇÃO (se for não verdadeira):
Classifique em apenas UMA das categorias:
(1) Sátira e Paródia: Sem intenção de causar dano
(2) Falsa Conexão: Manchetes não correspondem ao conteúdo
(3) Conteúdo Enganoso: Uso enganoso de informações
(4) Contexto Falso: Conteúdo genuíno com contexto falso
(5) Conteúdo Impostor: Fontes genuínas personificadas
(6) Conteúdo Manipulado: Conteúdo genuíno manipulado
(7) Conteúdo Fabricado: Totalmente falso e danoso
```

Fonte: Próprio Autor (2025).



IMAGEM 7 - TESTES DA INTEGRAÇÃO DO GEMINI AO INSTAGRAM



Fonte: Próprio Autor (2025).

O *Gemini* pode analisar o conteúdo, mas o modelo gratuito *Flash 2.5* não possui conectividade direta com a internet. Para possibilitar a verificação do conteúdo com fontes recentes, o algoritmo submete o conteúdo analisado a uma busca em múltiplas fontes relacionadas ao tema na internet, reunindo evidências para validação. Em seguida, por meio de *prompts* previamente definidos, os resultados da pesquisa são comparados com o conteúdo original, permitindo concluir se a notícia constitui ou não uma desinformação.

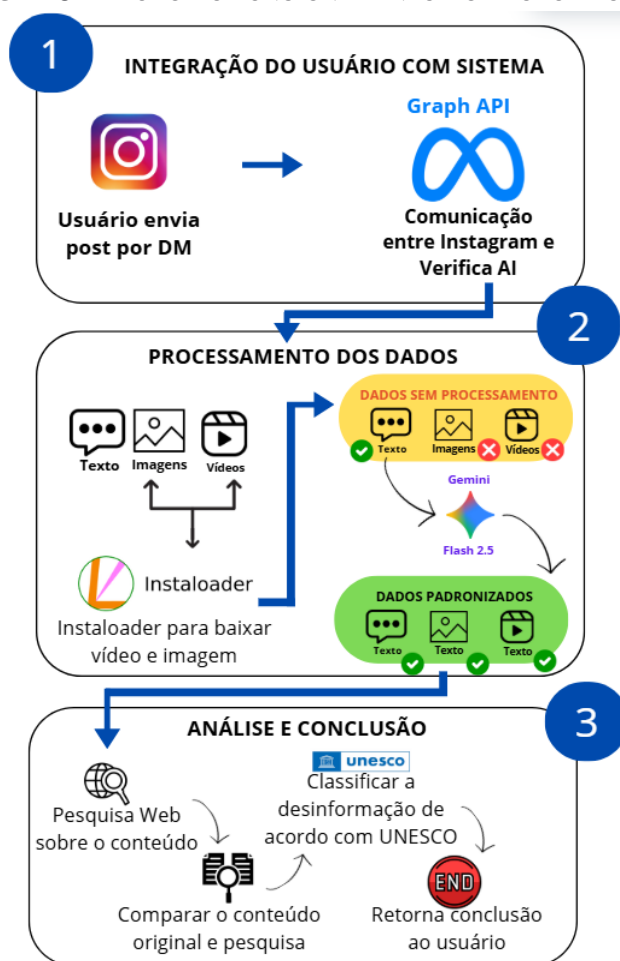
Se a probabilidade de o conteúdo ser classificado como desinformação for no mínimo de 95%, o sistema o categoriza conforme uma das sete categorias recomendadas pelo Manual da UNESCO:

1. Sátira e paródia: Conteúdo humorístico criado sem a intenção de causar danos, mas que tem o potencial de enganar;
2. Conexão falsa: Ocorre quando manchetes, imagens visuais ou legendas não correspondem ao conteúdo. Isso é frequente em manchetes do tipo "click-bait";



3. Conteúdo enganoso: Uso enganoso de informações, cortando fotos ou selecionando citações e estatísticas seletivamente, para enquadrar problemas ou indivíduos de determinadas maneiras;
4. Contexto falso: Conteúdo genuíno que é compartilhado com informações contextuais falsas, sendo visto sendo reciclado fora de seu contexto original;
5. Conteúdo impostor: Ocorre quando fontes genuínas como nomes de jornalistas ou logotipos de organizações respeitadas são apropriados indevidamente em materiais falsos;
6. Conteúdo manipulado: Quando o conteúdo genuíno (texto, fotos ou vídeos) é manipulado digitalmente para enganar;
7. Conteúdo fabricado: Conteúdo que é 100% falso, destinado a enganar e prejudicar, como sites de notícias completamente fabricados;

IMAGEM 8 - FLUXO DO FUNCIONAMENTO DO PROTÓTIPO



Fonte: Próprio Autor (2025).



Após a finalização do processo de backend, iniciou-se a criação da identidade visual. Também foi realizada nessa etapa a finalização da configuração do perfil responsável por receber as mensagens. Assim definiu-se o nome do perfil: “Verifica AI”. Após os de ajuste de *prompts* e refinamento do algoritmo, o protótipo foi finalizado com sucesso.

IMAGEM 9 - EQUIPE DESENVOLVENDO BACKEND E O DESIGN



Fonte: Próprio Autor (2025).



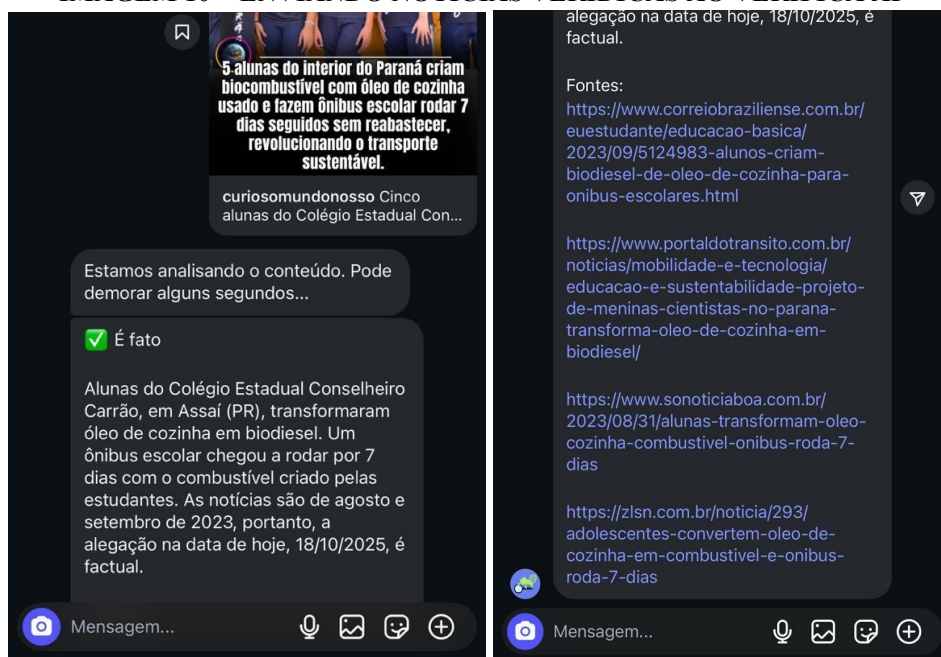
5 RESULTADOS E DISCUSSÕES

5.1 RESULTADOS

O protótipo desenvolvido está disponível aos usuários do Instagram por meio do perfil @verifica_ai_. Sua arquitetura foi pensada para que o usuário viva toda a experiência dentro da própria da própria plataforma da rede social. Ele opera de forma contínua, gratuita e sem a necessidade de instalação. O protótipo atendeu aos seguintes requisitos funcionais:

- A comunicação é nativa e realizada via mensagens diretas (DM);
- O sistema processa diferentes tipos de mídia, incluindo textos, imagens e vídeos.
- O conteúdo recebido é comparado com fontes recentes, e a veracidade da informação é verificada.
- Quando a desinformação é identificada, o caso é classificado em uma das sete categorias definidas pelo Manual da UNESCO.
- A resposta ao usuário é fornecida em um tempo ágil, com média de 30 segundos.
- As respostas são claras e objetivas, e incluem as fontes utilizadas na pesquisa para validação.

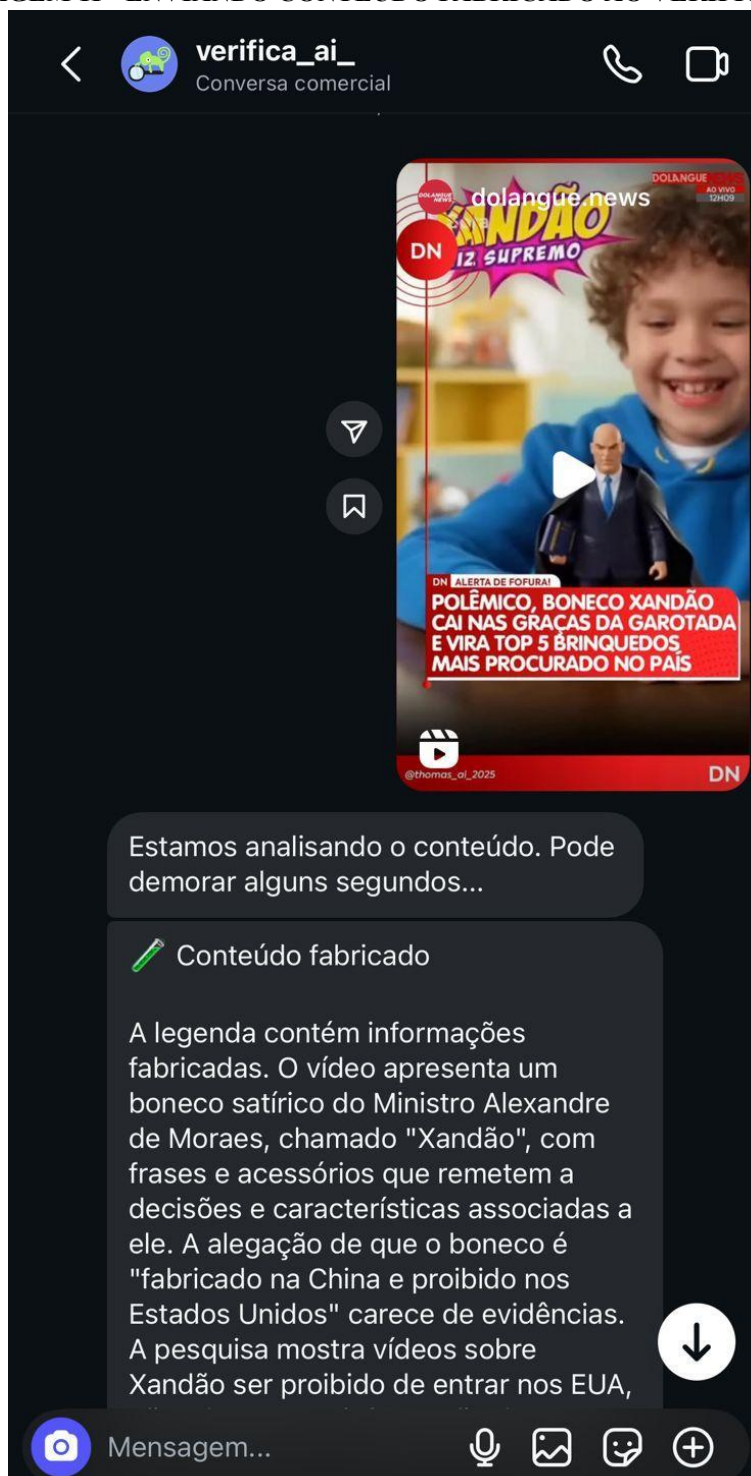
IMAGEM 10 - ENVIANDO NOTÍCIAS VERÍDICAS AO VERIFICA AI



Fonte: Próprio Autor (2025).



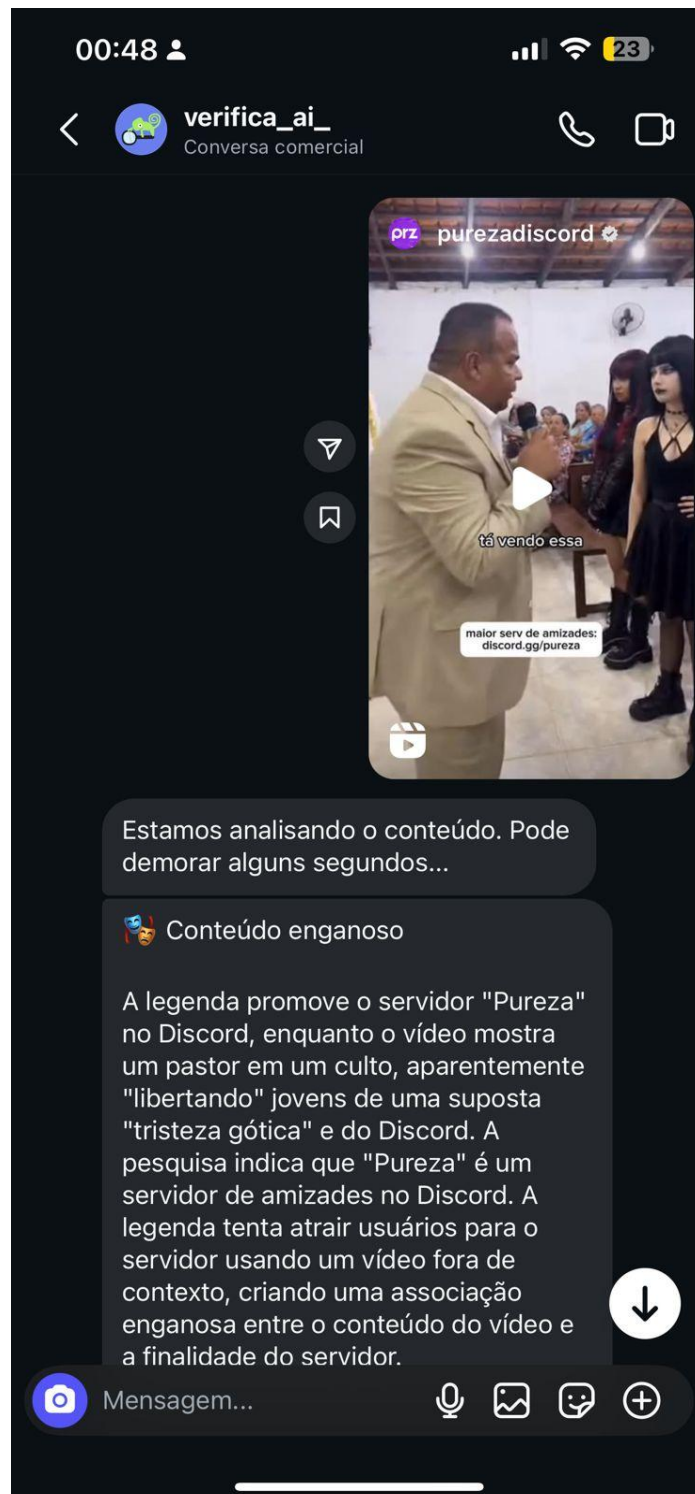
IMAGEM 11 - ENVIANDO CONTEÚDO FABRICADO AO VERIFICA AI



Fonte: Próprio Autor (2025).



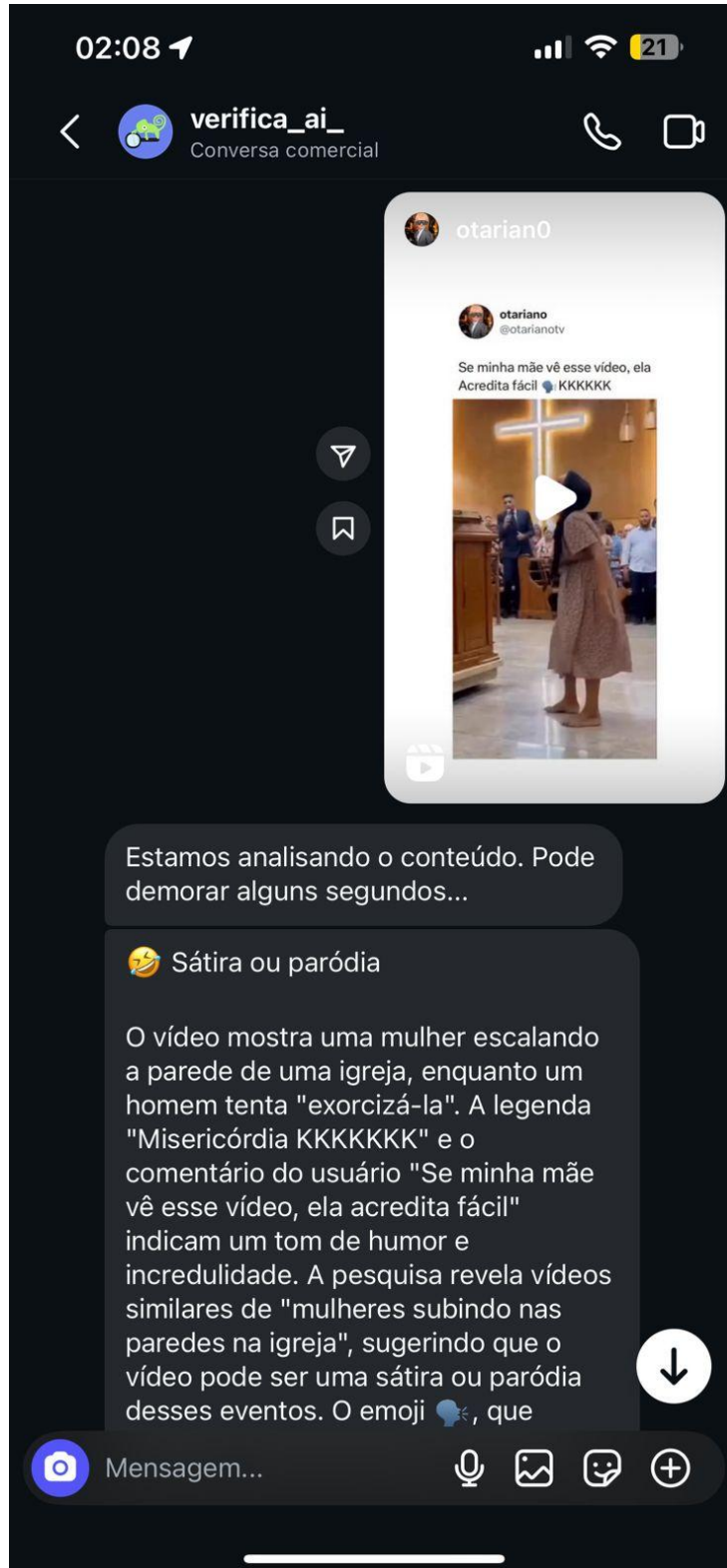
IMAGEM 12 - ENVIANDO CONTEÚDO ENGANOSO AO VERIFICA AI



Fonte: Próprio Autor (2025).



IMAGEM 13 - ENVIANDO SÁTIRA OU PARÓDIA AO VERIFICA AI



Fonte: Próprio Autor (2025).



Para avaliar o desempenho do modelo do sistema, utilizou-se a matriz de confusão. Ela é uma ferramenta tabular que resume o desempenho de um classificador, contabilizando acertos e erros por classe. Portanto, é frequentemente usada para avaliar modelos de classificação de forma objetiva identificando classes desbalanceadas. A tabela organiza os acertos e erros do modelo em verdadeiros positivos (VP), falsos positivos (FP), verdadeiros negativos (VN) e falsos negativos (FN) (GOOGLE, 2025).

A partir dela derivam-se métricas padrão:

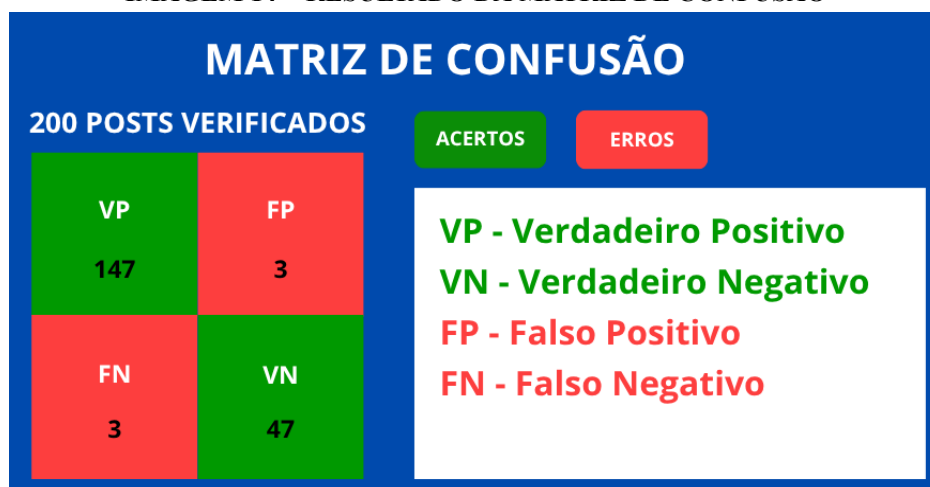
Acurácia: que expressa a proporção de classificações corretas no total, definida por $\text{Accuracy} = (\text{VP} + \text{VN}) / (\text{VP} + \text{VN} + \text{FP} + \text{FN})$;

Precisão (precision), que indica a fração de previsões positivas que de fato são positivas, $\text{Precision} = \text{VP} / (\text{VP} + \text{FP})$;

Sensibilidade ou recall, que mede a capacidade de recuperar os positivos reais, $\text{Recall} = \text{VP} / (\text{VP} + \text{FN})$.

Para cálculo da matriz de confusão, foi empregado um conjunto de 200 postagens, com temas variados, distribuídos de forma igualitária. Desse conjunto, 150 postagens eram verdadeiras e 50 falsas, contemplando os três formatos (texto, imagem e vídeo).

IMAGEM 14 - RESULTADO DA MATRIZ DE CONFUSÃO



Fonte: Próprio Autor (2025).



IMAGEM 15 - CÁLCULO DAS MÉTRICAS DA MATRIZ DE CONFUSÃO

$$\begin{aligned}\text{Acurácia: } & \frac{VP+VN}{VP+VN+FP+FN} \\ \text{Precisão: } & \frac{VP}{VP+FP} \\ \text{Recall (Sensibilidade): } & \frac{VP}{VP+FN}\end{aligned}$$

Fonte: Próprio Autor (2025).

TABELA 1 - RESULTADOS DAS MÉTRICAS DA MATRIZ DE CONFUSÃO

MÉTRICA	Cálculo	%
ACURÁCIA	Precision = $(147 + 47) / (147 + 47 + 3 + 3)$	97%
PRECISÃO	Precision = $147 / (147 + 3)$	98%
SENSIBILIDADE (RECALL)	Recall = $147 / (147 + 3)$	98%

Fonte: Próprio Autor (2025).

O modelo apresentou acurácia de 97%, precisão de 98% e sensibilidade (recall) de 98%, podendo-se concluir que classifica corretamente a grande maioria dos casos.

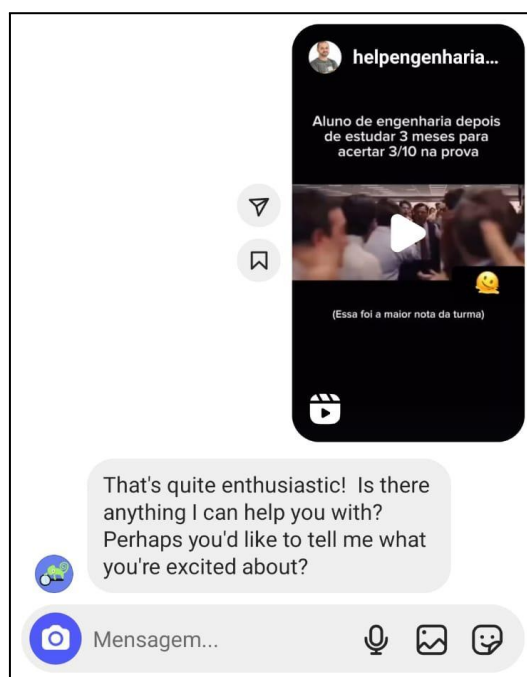
5.2 LIMITAÇÕES E DESAFIOS

A utilização de um modelo de linguagem gratuito resultou em uma vantagem financeira, contudo, isso impõe limitações ao projeto. Apesar do *Gemini Flash 2.5* ser capaz de processar diversas requisições em paralelo, ele estabelece um limite de uso, restringindo a 125 verificações por dia e 5 por minuto. Durante os testes da matriz de confusão, constatou-se uma latência na resposta de aproximadamente 10% das verificações, com tempos de resposta superiores ao dobro da média.

Houve também casos de falha crítica na resposta, correspondendo a 5% das verificações. Nesses casos, o modelo do *Gemini* ficou sobrecarregado com um grande volume de informações, e uma limitação de sua capacidade de processamento impediu a geração de uma resposta, deixando o usuário sem retorno.



Imagem 16 - Erro crítico na resposta



Fonte: Próprio Autor (2025).

A matriz de confusão indicou também a ocorrência pontual de falsos positivos e falsos negativos. Por último, em situações onde o conteúdo era totalmente fabricado, a busca na internet apresenta poucos resultados, e o sistema extrai dados de sites não oficiais para compor a resposta ao usuário, comprometendo a confiabilidade da resposta gerada.

5.3 DISCUSSÃO DOS RESULTADOS

Os resultados obtidos no desenvolvimento e avaliação do protótipo Verifica AI revelam um desempenho robusto em termos de funcionalidade, mas também expõe limitações ao uso de tecnologias gratuitas. Restrições operacionais como limites de uso diário, e apesar raramente, latência nas respostas e erros críticos, inviabilizam o uso em larga escala.

A ocorrência de falsos positivos e falsos negativos encontrados em conteúdos gerados por inteligência artificial, além da extração de dados em sites não confiáveis em casos de conteúdos totalmente fabricados, apontam a necessidade de melhorias nos *prompts* usados no modelo do *Gemini*.



6 CONCLUSÕES OU CONSIDERAÇÕES FINAIS

Este projeto alcançou todos os objetivos previamente estabelecidos, resultando na criação de uma ferramenta gratuita e integrada ao Instagram para a classificação de desinformação, em conformidade com as sete categorias da UNESCO. A solução proposta supera as limitações de ferramentas existentes, como *Meta AI* e *AI Studio*, ao analisar posts de forma nativa e contribuir para a redução da viralização de desinformação.

Os resultados obtidos demonstram alta efetividade do protótipo na detecção de desinformação em conteúdos multimodais (texto, imagem e vídeo). As métricas de desempenho alcançadas foram: 97% de acurácia, 98% de precisão e 98% de sensibilidade, em um conjunto de 200 postagens testadas.

A escalabilidade da solução, no entanto, mostra-se inviável devido à dependência do modelo gratuito *Gemini Flash 2.5*. Essa dependência limita o uso diário, causa latência nas respostas e falhas críticas. Outros desafios não relacionados ao modelo foram identificados, principalmente na classificação de conteúdos que são completamente fabricados, os quais apresentaram maior dificuldade de serem corretamente classificados.

Para projetos futuros, sugere-se o aprimoramento da versão desenvolvida neste estudo. Recomenda-se o ajuste dos *prompts* e a melhoria do algoritmo para obter maior exatidão, bem como a adoção de um modelo de IA pago para eliminar os limites de uso e melhorar a qualidade das respostas.



REFERÊNCIAS

DELMazo, Caroline; VALENTE, Jonas C.L. Fake news nas redes sociais online: propagação e reações à desinformação em busca de cliques. *Media & Jornalismo*, Lisboa, v. 18, n. 32, p. 5-23, abr. 2018. Disponível em: https://scielo.pt/scielo.php?pid=S2183-54622018000100012&script=sci_arttext. Acesso em: 18 setembro. 2025.

GOOGLE. Classification: Accuracy, recall, precision, and related metrics. Google Developers – Machine Learning Crash Course, 2025. Disponível em: <https://developers.google.com/machine-learning/crash-course/classification/accuracy-precision-recall>. Acesso em: 19 out. 2025.

GOOGLE CLOUD. What are AI agents? [S.l.]: Google Cloud, 2025. Disponível em: <https://cloud.google.com/discover/what-are-ai-agents>. Acesso em: 18 out. 2025.

IRETON, Cherilyn; POSETTI, Julie. Journalism, 'fake news' & disinformation: handbook for journalism education and training. Paris: UNESCO, 2018. Disponível em: <https://unesdoc.unesco.org/ark:/48223/pf0000265552>. Acesso em: 18 out. 2025.

OCDE (2024), “The OECD Truth Quest Survey: Methodology and findings”, OECD Digital Economy Papers, No. 369, OECD Publishing, Paris, <https://doi.org/10.1787/92a94c0f-en>.

UNESCO. Jornalismo, fake news & desinformação: manual para educação e treinamento em jornalismo. Cherilyn Ireton e Julie Posetti (org.). Brasília: UNESCO, 2018.

WE ARE SOCIAL; HOOTSUITE. Digital 2024: global overview report. [S. l.]: We Are Social; Hootsuite, 2024. Disponível em: <https://wearesocial.com/uk/blog/2024/01/digital-2024/>. Acesso em: 12 Junho. 2025.



WORLD ECONOMIC FORUM. Global risks report 2025. Genebra: World Economic Forum, 2025. Disponível em: <https://www.weforum.org/publications/global-risks-report-2025/>. Acesso em: 10 Junho. 2025.

ANEXO 1

Passo a Passo para testar o Verifica AI no Instagram:

1. Acesse o perfil do Verifica AI : https://www.instagram.com/verifica_ai/
2. Envie o conteúdo ao perfil do Instagram através de *DM (Direct Message)*. Os conteúdos aceitos são:
 - a. Post do Instagram com *Reels*, fotos ou vídeos.
 - b. Texto em formato afirmativo descrevendo o que você deseja verificar.. (ex: “A China é comunista”, “Brasil entrará na OTAN devido a ameaça de invasão americana”, “Trump morreu”) .
 - c. Enviar imagens do próprio dispositivo que contenham um texto ou imagem a ser analisada.
3. Aguarde aproximadamente 30 segundos, e o *chatbot* enviará uma resposta dentro da própria DM iniciada.

Nota 1: Como um modelo de LLM gratuito, suas respostas podem apresentar latência e aproximadamente 5% de erro. Além disso, seu propósito é especificamente verificar *fake news*, conteúdos enviados que não contenham uma afirmação ou notícia a ser analisada, podem resultar em uma resposta imprecisa.

Nota 2: Se a resposta demorar mais de 1 minuto, ocorreu um erro crítico e a resposta ao conteúdo não será enviada. Os motivos mais prováveis incluem vídeos e *Reels* muito longos, que geram muitos comandos a serem processados pela LLM gratuita, e por limitação de processamento do *Gemini* a resposta é interrompida. Meios de lidar com esse erro ainda estão sendo analisados.