



**UNIVERSIDADE FEDERAL DO CARIRI
CENTRO DE CIÊNCIAS E TECNOLOGIA
DEPARTAMENTO DE CIÊNCIA DA COMPUTAÇÃO
CURSO DE GRADUAÇÃO EM CIÊNCIA DA COMPUTAÇÃO**

CARLOS EDUARDO TAVARES DO NASCIMENTO

**UFCHAT: UMA PROPOSTA DE ASSISTENTE DE INTELIGÊNCIA ARTIFICIAL
PARA APOIO INFORMACIONAL NA UFCA**

JUAZEIRO DO NORTE - CE

2026

CARLOS EDUARDO TAVARES DO NASCIMENTO

UFCHAT: UMA PROPOSTA DE ASSISTENTE DE INTELIGÊNCIA ARTIFICIAL PARA
APOIO INFORMACIONAL NA UFCA

Trabalho de Conclusão de Curso apresentado ao
Curso de Graduação em Ciência da Computação
do Centro de Ciências e Tecnologia da Universi-
dade Federal do Cariri, como requisito parcial
à obtenção do grau de bacharel em Ciência da
Computação.

Orientadora: Prof. Dra. Luana Batista
da Cruz

JUAZEIRO DO NORTE - CE

2026

Dados Internacionais de Catalogação na Publicação
Universidade Federal do Cariri
Sistema de Bibliotecas

N244u Nascimento, Carlos Eduardo Tavares do.

UFChat: uma proposta de assistente de inteligência artificial para apoio
informacional na UFCA / Carlos Eduardo Tavares do Nascimento. – 2026.
75 f. : il. color - Inclui bibliografia: f. 56-59.

Trabalho de Conclusão de Curso (Graduação em Ciência da Computação) – Centro de
Ciências e Tecnologia, Departamento de Ciência da Computação, Universidade
Federal do Cariri, Juazeiro do Norte, CE, 2026.

Orientadora: Prof. Dra. Luana Batista da Cruz.

1. Inteligência Artificial. 2. IA Generativa. 3. Retrieval Augmented Generation (RAG).
4. Large Language Models (LLMs). 5. Chatbots. I. Cruz, Luana Batista da (Orient.). II.
Universidade Federal do Cariri. III. Título.

CDD 006.3

CARLOS EDUARDO TAVARES DO NASCIMENTO

UFCHAT: UMA PROPOSTA DE ASSISTENTE DE INTELIGÊNCIA ARTIFICIAL PARA
APOIO INFORMACIONAL NA UFCA

Trabalho de Conclusão de Curso apresentado ao
Curso de Graduação em Ciência da Computação
do Centro de Ciências e Tecnologia da Universi-
dade Federal do Cariri, como requisito parcial
à obtenção do grau de bacharel em Ciência da
Computação.

Aprovada em: 27/03/2026.

BANCA EXAMINADORA

Prof. Dra. Luana Batista da Cruz (Orientadora)
Universidade Federal do Cariri (UFCA)

Prof. Dr. João Otávio Bandeira Diniz
Instituto Federal do Maranhão (IFMA)

Prof. Dr. Jayr Alencar Pereira
Universidade Federal do Cariri (UFCA)

À minha família, por sua capacidade de acreditar em mim e investir em mim. Em especial, ao meu irmão por me orientar com conselhos e dicas de construção científica, à minha mãe, sempre presente com carinho em minha vida e a João, sempre me apoiando e me motivando, com muito amor, a melhorar.

AGRADECIMENTOS

À professora Luana Batista por me apoiar e me orientar em meu projeto de conclusão de curso.

Aos meus colegas da universidade que são também meus amigos, e me apoiaram com testes, conteúdo e conselhos para meu projeto.

À Universidade Federal do Cariri, a qual sempre foi muito acolhedora e confortável, além de bastante apoiadora com projetos inovadores de seus alunos.

“Assim como uma mãe protege com a própria vida o seu filho, o seu único filho, assim deve o coração expandir-se sem limites para abraçar todos os seres vivos.”

(Siddharta Gautama)

RESUMO

No ambiente acadêmico, a dispersão e o volume de informações dificultam o acesso rápido a dados essenciais por parte de alunos, técnicos e professores, como contatos institucionais, links e documentos. A ausência de um canal centralizado compromete a eficiência da comunicação e sobrecarrega os agentes da universidade com dúvidas recorrentes. Este trabalho propõe a utilização de modelos de *Large Language Models* (LLMs) como assistentes conversacionais capazes de atuar como recepcionistas virtuais, oferecendo respostas ágeis e contextualizadas às demandas informativas da instituição. Para avaliar a receptividade da comunidade acadêmica, foi aplicado um questionário sobre o uso prévio de ferramentas de Inteligência Artificial (IA) e as funcionalidades que tornariam um assistente virtual útil para as atividades da UFCA. A partir dessas indicações, é apresentado o UFChat, um modelo LLM baseado no GPT-4.1 mini, que utiliza um mecanismo de Roteador Semântico para classificar e direcionar as requisições do usuário, seguido da integração com um Banco de Dados Vetorial por meio da tecnologia *Retrieval Augmented Generation* (RAG), responsável pela busca de informações e melhoramento de respostas. O sistema também incorpora técnicas de *Web Scraping* para coleta de dados em fontes oficiais, permitindo a geração de respostas mais precisas e alinhadas ao domínio institucional. Os testes envolvendo a comunidade acadêmica da UFCA indicam que o UFChat produz respostas mais eficazes e preferíveis em comparação a outros assistentes comerciais, como ChatGPT e Gemini, além de apresentar a flexibilidade de funcionamento via WhatsApp, contribuindo para a melhoria da comunicação institucional.

Palavras-chave: Inteligência Artificial. *Large Language Model*. Chatbots. UFChat.

ABSTRACT

In the academic environment, the dispersion and volume of information hinder quick access to essential data by students, staff, and professors, such as institutional contacts, links, and documents. The absence of a centralized channel compromises communication efficiency and overloads university agents with recurring questions. This work proposes the use of *Large Language Models* (LLMs) as conversational assistants capable of acting as virtual receptionists, offering fast and contextualized responses to the institution's informational demands. To evaluate the receptivity of the academic community, a questionnaire was applied regarding the prior use of Artificial Intelligence (AI) tools and the functionalities that would make a virtual assistant useful for UFCA activities. Based on these indications, UFChat is presented, an LLM model based on GPT-4.1 mini, which uses a Semantic Router mechanism to classify and route user requests, followed by integration with a Vector Database through *Retrieval Augmented Generation* (RAG) technology, responsible for information retrieval and response improvement. The system also incorporates *Web Scraping* techniques for data collection from official sources, allowing the generation of more precise and institutionally aligned responses. Tests involving the UFCA academic community indicate that UFChat produces more effective and preferable responses compared to other commercial assistants, such as ChatGPT and Gemini, in addition to presenting the flexibility of operation via WhatsApp, contributing to the improvement of institutional communication.

Keywords: Artificial Intelligence. Large Language Model. Chatbots. UFChat.

LISTA DE FIGURAS

Figura 1 – Organograma institucional da UFCA.	18
Figura 2 – Representação da proximidade de vetores baseado em seus valores semânticos.	19
Figura 3 – Transformação de uma sentença em <i>tokens</i> e, em seguida, em vetores de <i>embeddings</i>	20
Figura 4 – Exemplo de Roteador Semântico.	22
Figura 5 – Exemplo de vetores de <i>embeddings</i> no espaço vetorial.	24
Figura 6 – Distribuição das respostas à pergunta 5 do questionário.	31
Figura 7 – Distribuição das respostas à pergunta 6 do questionário.	32
Figura 8 – Etapas do Método Proposto.	33
Figura 9 – Esquema de integração com o WhatsApp.	43
Figura 10 – Exemplos de diálogos com o UFChat.	57
Figura 11 – Exemplo de diálogo com o UFChat.	58

LISTA DE TABELAS

Tabela 1 – Questionário aplicado para levantamento de necessidades da comunidade acadêmica.	30
Tabela 2 – Categorias possíveis e seus respectivos contextos pré-definidos.	34
Tabela 3 – Fluxo para classificação e resumo do Modelo de Classificação e Intenção. .	35
Tabela 4 – Fluxo para classificação e resumo do Modelo de Classificação e Intenção. .	36
Tabela 5 – Fluxo para classificação e resumo do Modelo de Classificação e Intenção. .	37
Tabela 6 – Frequência de atualização do Banco de Dados Vetorial.	38
Tabela 7 – Fluxo para geração de resposta pelo Modelo Final.	39
Tabela 8 – Respostas dos modelos para a pergunta: Qual o endereço da PRPI UFCA? .	42
Tabela 9 – Exemplo de interação registrada como Dúvidas Gerais.	45
Tabela 10 – Exemplo de interação registrada como Dúvidas Gerais.	45
Tabela 11 – Continuação de exemplo de interação registrada como Dúvidas Gerais. . .	46
Tabela 12 – Exemplo de interação registrada como Link Edital.	47
Tabela 13 – Exemplo de interação registrada como Novidades ou Notícias.	48
Tabela 14 – Exemplo de interação sobre Cardápio do RU.	49
Tabela 15 – Exemplo de interação registrada como Requisição Perigosa.	49
Tabela 16 – Outro exemplo registrado como Requisição Perigosa.	50
Tabela 17 – Exemplo de solicitação sobre um projeto inexistente na UFCA.	50
Tabela 18 – Exemplo com campo <i>system</i> padrão.	51
Tabela 19 – Resultados do RAGAS usando o campo <i>system</i> padrão do UFChat.	52
Tabela 20 – Exemplo com campo <i>system</i> provisório.	53
Tabela 21 – Resultados do RAGAS usando o campo <i>system</i> provisório.	53
Tabela 22 – Distribuição de votos por assistente no teste A/B.	54
Tabela 23 – Maior similaridade do cosseno entre assistentes por pergunta.	56
Tabela 24 – Formulário aplicado para avaliação de respostas do UFChat.	66

LISTA DE ABREVIATURAS E SIGLAS

HTML	<i>Hypertext Markup Language</i>
IA	Inteligência Artificial
IFSC	<i>Instituto Federal de Santa Catarina</i>
PLN	Processamento de Linguagem Natural
RAG	<i>Retrieval Augmented Generation</i>
SIGAA	Sistema Integrado de Gestão de Atividades Acadêmicas
UFC	Universidade Federal do Ceará
UFCA	Universidade Federal do Cariri
UFPI	Universidade Federal do Piauí
UNESC	Centro Universitário do Espírito Santo
Unimib	Universidade de Milano Bicocca
VPS	<i>Virtual Private Server</i>

SUMÁRIO

1	INTRODUÇÃO	14
1.1	Objetivo Geral	15
1.2	Objetivos Específicos	15
1.3	Organização do Trabalho	15
2	FUNDAMENTAÇÃO TEÓRICA	17
2.1	Universidade Federal do Cariri (UFCA)	17
2.2	IA Generativa	18
2.2.1	<i>Tokens e Embeddings</i>	<i>18</i>
2.2.2	<i>Transformers</i>	<i>20</i>
2.2.3	<i>Large Language Models (LLMs)</i>	<i>20</i>
2.3	Roteador Semântico	21
2.4	<i>Prompt Chaining</i>	<i>22</i>
2.5	Banco de Dados Vetorial e Retrieval Augmented Generation (RAG) . . .	23
2.5.1	<i>Similaridade do Cosseno</i>	<i>23</i>
2.5.2	<i>Obtenção de Dados por meio de Web Scraping</i>	<i>24</i>
2.6	<i>Chatbots</i>	<i>25</i>
2.6.1	<i>Memória por meio de Conversation Buffer Memory</i>	<i>25</i>
2.7	Métodos de Avaliação	26
2.7.1	<i>Avaliação RAGAS</i>	<i>26</i>
2.7.2	<i>Teste A/B</i>	<i>27</i>
3	TRABALHOS RELACIONADOS	28
4	MATERIAIS E MÉTODO PROPOSTO	30
4.1	Materiais	30
4.1.1	<i>Levantamento de Necessidades da Comunidade Acadêmica</i>	<i>30</i>
4.2	Método Proposto	32
4.3	Roteador Semântico	33
4.3.1	<i>Memória do Assistente</i>	<i>36</i>
4.4	Coleta de Informações	37
4.5	Produção de Respostas	38
4.6	Avaliação do Assistente	40

4.6.1	<i>Avaliação Quantitativa</i>	40
4.6.2	<i>Avaliação Qualitativa</i>	41
4.6.3	<i>Integração do Método Proposto ao WhatsApp</i>	42
5	RESULTADOS E DISCUSSÕES	44
5.1	Configuração Experimental	44
5.2	Resultados Experimentais	44
5.2.1	<i>Dúvidas Gerais</i>	45
5.2.2	<i>Links de Editais</i>	46
5.2.3	<i>Novidades ou Notícias</i>	47
5.2.4	<i>Cardápio do RU</i>	48
5.2.5	<i>Requisição Perigosa</i>	49
5.2.6	<i>Mitigação de Alucinação</i>	50
5.2.7	<i>Avaliação Quantitativa</i>	51
5.2.8	<i>Avaliação Qualitativa</i>	54
5.2.8.1	<i>Similaridade entre Respostas dos Assistentes</i>	55
5.3	Integração com o WhatsApp	56
5.4	Impactos Importantes e Limitações	59
6	CONCLUSÕES E TRABALHOS FUTUROS	61
	REFERÊNCIAS	62
	APÊNDICES	66
	APÊNDICE A – Apêndice A: Formulário de Avaliação de Assistente Virtual	66
	ANEXOS	81

1 INTRODUÇÃO

No ambiente universitário, a disseminação diária de informações para alunos, técnicos e professores nem sempre é clara ou de fácil acesso. Na Universidade Federal do Cariri (UFCA), com cinco campi em cidades distintas e diversos órgãos vinculados à instituição, a gestão e o acesso a informações tornam-se desafiadores (Universidade Federal do Cariri, 2025b).

A ausência de um canal único e inteligente de comunicação leva a problemas recorrentes, como informações erradas, mal distribuídas ou confusas que acabam gerando retrabalho administrativo, atrasos em processos e congestionamento dos canais de atendimento. Em um contexto de sobrecarga de secretarias e pró-reitorias, a comunidade acadêmica frequentemente enfrenta longos tempos de espera para obter informações simples, como prazos de matrícula e documentos necessários para processos administrativos (OLIVEIRA, 2024).

Esse cenário não é exclusivo da UFCA. Em universidades brasileiras e internacionais (ZHANG *et al.*, 2016; JEZLER *et al.*, 2025; OLIVEIRA, 2024), observa-se um aumento expressivo no volume de demandas informacionais, ao mesmo tempo em que os canais tradicionais (atendimento presencial, telefone e e-mail) não acompanham a velocidade exigida pela comunidade acadêmica. Pesquisas indicam que o tempo médio de resposta de canais institucionais pode variar entre dias e semanas (ZHANG *et al.*, 2016), comprometendo a experiência dos usuários e o bom andamento das atividades acadêmicas.

Algumas instituições já buscaram mitigar esse problema com o uso de assistentes virtuais. Um exemplo é o assistente dedicado ao atendimento de estudantes da Universidade de Milano Bicocca (Unimib) (ANTICO *et al.*, 2024), que atua como assistente geral dos alunos, fornecendo informações sobre a universidade, materiais de estudo e até preços de refeitórios. Outro caso é o modelo implementado no Centro Universitário do Espírito Santo (UNESC) (JEZLER *et al.*, 2025), que reforça a viabilidade dessa abordagem.

Os resultados observados em diversas universidades indicam que a utilização da Inteligência Artificial (IA) Generativa na comunicação institucional contribui para reduzir os tempos de resposta, aumentar a satisfação dos usuários e otimizar a organização interna da informação (HSAIN; HOUSNI, 2024). Além disso, as aplicações de IA não se resumem apenas a atendimentos, a Universidade Federal do Piauí (UFPI), por exemplo, implementou um *chatbot* para auxiliar no ensino de programação (IFPI – Instituto Federal do Piauí, 2025). Isso demonstra o amplo potencial da IA Generativa em contribuir para diferentes áreas do ambiente acadêmico.

Portanto, com os avanços recentes em sistemas baseados em LLM, como o ChatGPT,

torna-se viável a criação de assistentes de IA, autônomos e capazes de responder a dúvidas institucionais de forma contextualizada, contínua e adaptável. Diferentemente dos *chatbots* tradicionais, baseados em respostas rígidas e fluxos pré-definidos, os LLMs permitem interações mais naturais e flexíveis, reduzindo ambiguidades e aumentando a assertividade no atendimento.

1.1 Objetivo Geral

Desenvolver o UFChat, um assistente de IA capaz de auxiliar a comunidade acadêmica da UFCA no acesso a informações institucionais, oferecendo respostas rápidas, confiáveis e centralizadas.

1.2 Objetivos Específicos

- Buscar na literatura métodos, tecnologias e aplicações utilizadas na construção de assistentes de IA modernos;
- Mapear as dúvidas mais recorrentes entre alunos, técnicos e professores da UFCA para orientar o desenvolvimento do assistente;
- Estudar e desenvolver um Banco de Dados Vetorial com informações oficiais e documentos disponibilizados pela UFCA;
- Avaliar a eficácia do método proposto utilizando métricas de desempenho para modelos baseados em LLMs e testes com usuários reais;
- Comparar o desempenho do método proposto com abordagens já consolidadas na literatura;
- Disponibilizar o *chatbot* em plataformas de fácil acesso, como o WhatsApp, para atendimento à comunidade acadêmica.

1.3 Organização do Trabalho

Os demais capítulos deste trabalho foram organizados em:

- Capítulo 2 – Trabalhos Relacionados: apresenta trabalhos realizados com caráter similar ao que é proposto por este trabalho.
- Capítulo 3 – Fundamentação Teórica: discute os conceitos relacionados a IA Generativa, *chatbots* e ambientes educacionais.
- Capítulo 4 – Método Proposto: descreve as etapas do UFChat e as ferramentas utilizadas.
- Capítulo 5 – Resultados e Discussão: apresenta os resultados obtidos nos testes realizados

e a análise da eficácia do sistema.

- Capítulo 6 – Conclusão: sintetiza as contribuições do trabalho e aponta possíveis melhorias.

2 FUNDAMENTAÇÃO TEÓRICA

Neste capítulo são apresentados os fundamentos conceituais e tecnológicos que orientam a construção do UFChat. A ideia central é estabelecer uma base sólida que justifique tanto a necessidade do assistente quanto as escolhas técnicas feitas para sua implementação. Para isso, parte-se da caracterização do ambiente organizacional da UFCA, o qual fornece o contexto do problema, e, em seguida, desenvolvem-se os conceitos de IA Generativa e suas tecnologias associadas.

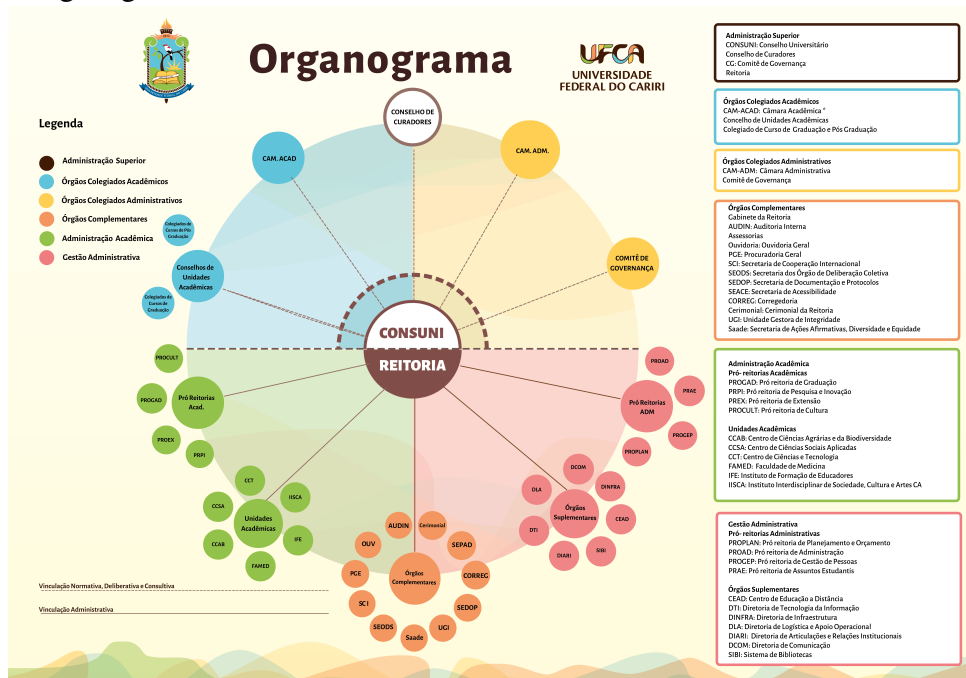
2.1 Universidade Federal do Cariri (UFCA)

A UFCA foi fundada em 2013, a partir do desmembramento da Universidade Federal do Ceará (UFC), e desde então tem se consolidado como uma instituição fundamental para a região do Cariri cearense. Atualmente, a UFCA conta com cinco campi, distribuídos nas cidades de Juazeiro do Norte, Barbalha, Crato, Brejo Santo e Icó, atendendo a milhares de estudantes e articulando atividades de ensino, pesquisa e extensão (Universidade Federal do Cariri, 2025a).

Sua estrutura organizacional é composta por dezenas de órgãos e setores administrativos, como representado no organograma institucional (Figura 1). Cada um desses órgãos possui atribuições específicas, variando desde a gestão acadêmica até a extensão universitária, o que, embora fundamental para a eficiência administrativa, frequentemente gera confusões para estudantes e servidores. Por exemplo, é comum que haja dúvidas sobre quais demandas devem ser direcionadas a determinadas pró-reitorias, e, muitas vezes, a resolução dessas dúvidas leva dias.

Segundo dados do Censo da Educação Superior (CARIRI, 2019), em 2023 a UFCA registrou 1109 ingressos e 3354 matrículas. Esse número expressivo, aliado à multiplicidade de órgãos, evidencia um problema, que é a dispersão da informação. Atualmente, os dados de interesse da comunidade acadêmica encontram-se fragmentados entre o site institucional, o Sistema Integrado de Gestão de Atividades Acadêmicas (SIGAA), sistemas de pró-reitorias, páginas específicas e até e-mails institucionais. Tal fragmentação dificulta a busca por informações e gera sobrecarga no atendimento humano. Nesse sentido, um canal unificado de comunicação, ágil e centralizado, torna-se não apenas desejável, mas necessário. É aqui que surge o papel de sistemas de automação baseados em IA.

Figura 1 – Organograma institucional da UFCA.



Fonte: <<https://www.ufca.edu.br/instituicao/administrativo/estrutura-organizacional/>>

2.2 IA Generativa

A IA Generativa representa a vertente mais inovadora da IA nos últimos anos. De acordo com a plataforma Statista, o mercado desta tecnologia apresenta taxa de crescimento estimada em 37% ao ano (STATISTA, 2025), reflexo do seu impacto em múltiplos setores. Sua principal característica é a capacidade de compreender e gerar novos conteúdos como textos, imagens e até códigos, de forma natural, aproximando-se do modo humano de produção de informação.

O funcionamento desse tipo de IA é viabilizado pela arquitetura *Transformer*, que processa informações em vetores numéricos, chamados *embeddings*, representando o significado de palavras e frases (UDDIN *et al.*, 2025). Essa arquitetura abriu caminho para sistemas amplamente difundidos, como o ChatGPT e o Gemini, que tornaram a tecnologia acessível ao grande público e exemplificaram seu potencial para tarefas de atendimento, ensino e pesquisa.

2.2.1 Tokens e Embeddings

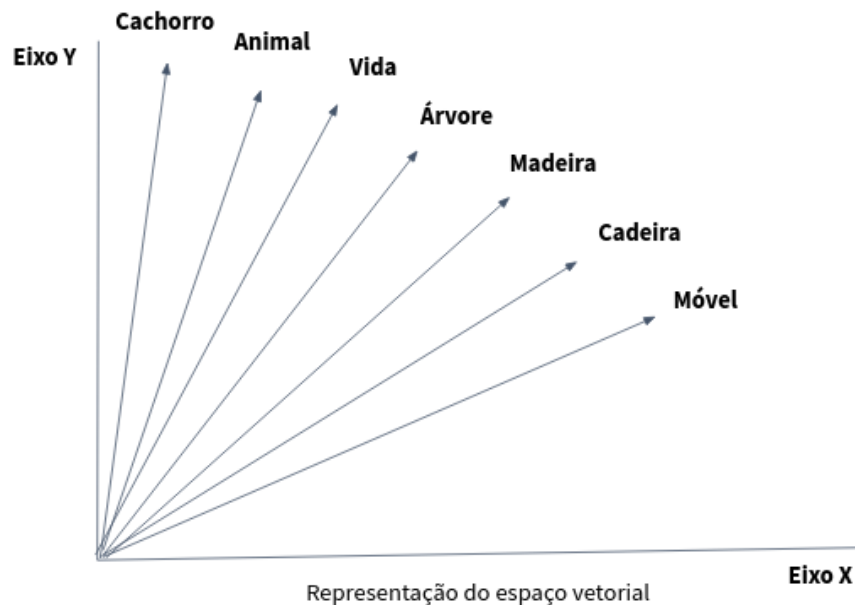
O primeiro passo para que o LLM processe a linguagem natural é a *tokenização*. O *tokenizer* decompõe um texto em unidades básicas, como palavras ou pontuações, e em seguida essas unidades são transformadas em *tokens*. Os *tokens* representam estas unidades como

números a partir de um vocabulário próprio do *tokenizer*, onde cada *token* representa unicamente aquele fragmento textual (TEKGUL, 2023).

Entretanto, esses números ainda são dados crus. Para que se tornem representações úteis, os *tokens* são convertidos em *embeddings*, vetores de alta dimensionalidade que capturam o significado semântico dos *tokens* (LOPES, 2024). Estes vetores de *embeddings*, por representarem o sentido das palavras, possuem posições próximas no espaço vetorial quando o sentido entre eles é similar. Por exemplo: “gato” e “cachorro” possuem maior proximidade entre si do que “gato” e “cadeira”.

Na Figura 2, é apresentado um exemplo ilustrativo de vetores de *embeddings*. Nesse contexto, palavras que possuem relações semânticas entre si estão mais próximas no espaço vetorial, enquanto aquelas com menor relação permanecem mais distantes. Esta aproximação pode ser utilizada em buscas a partir de cálculos espaciais com vetores, como a Similaridade do Cosseno, que será explicada nas seções seguintes.

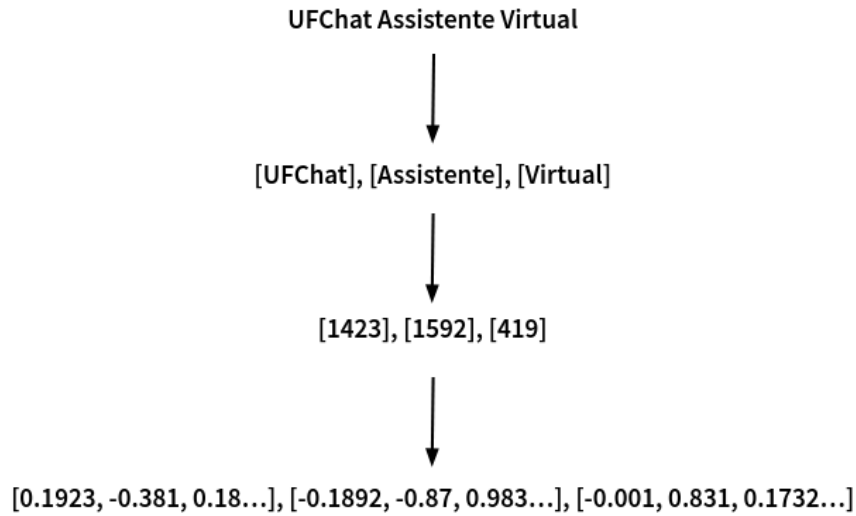
Figura 2 – Representação da proximidade de vetores baseado em seus valores semânticos.



Fonte: Elaborado pelo autor.

Com base nesse conceito, a Figura 3 apresenta um exemplo prático do processo, em que a frase “UFChat Assistente Virtual” é inicialmente separada em palavras e depois convertida em números (*tokens*) a partir de um dicionário próprio do programa *Tokenizer*. Posteriormente, esses *tokens* são transformados em vetores de *embeddings*, tornando a informação interpretável pelo LLM.

Figura 3 – Transformação de uma sentença em *tokens* e, em seguida, em vetores de *embeddings*.



Fonte: Elaborado pelo autor.

2.2.2 Transformers

O salto tecnológico ocorreu em 2017, quando Vaswani et al. apresentaram o modelo *Transformer* no artigo *Attention is All You Need* (VASWANI *et al.*, 2017). Diferentemente de modelos anteriores, como os *Seq2Seq*, que são modelos que transformam um texto em outro texto, os *Transformers* introduziram o mecanismo de atenção (*attention*), que atribui pesos a cada *token* em relação a todos os outros da sequência. Dessa forma, o modelo consegue capturar dependências de longo alcance, fundamentais para preservar o sentido em textos longos ou complexos.

Esse mecanismo tornou possível não apenas traduções automáticas mais precisas, mas também a síntese de textos, a resposta a perguntas e a criação de narrativas completas. Em outras palavras, o *Transformer* é o alicerce da moderna IA Generativa.

2.2.3 Large Language Models (LLMs)

Os *Transformers*, quando expandidos em escala, deram origem aos LLMs. Esses modelos são treinados em quantidades massivas de dados, provenientes de repositórios digitais como o *Common Crawl*, textos acadêmicos, obras literárias e até comentários em redes sociais (Stryker, Cole, 2023; NIEMEYER, 2025). Essa amplitude de dados garante aos LLMs uma capacidade versátil de responder a diferentes tipos de questionamento.

Na prática, essa capacidade se manifesta por meio da interação com o usuário, na

qual o modelo recebe uma entrada e gera uma resposta em linguagem natural. Nesse contexto, a forma como a entrada é estruturada influencia diretamente o comportamento do modelo. Em aplicações baseadas em LLMs, é comum a utilização de campos como o *system*, responsável por definir diretrizes comportamentais, restrições e o papel do assistente na interação e o *user*, que contém a requisição do usuário, permitindo maior controle e consistência nas respostas (SAVAGE, 2024).

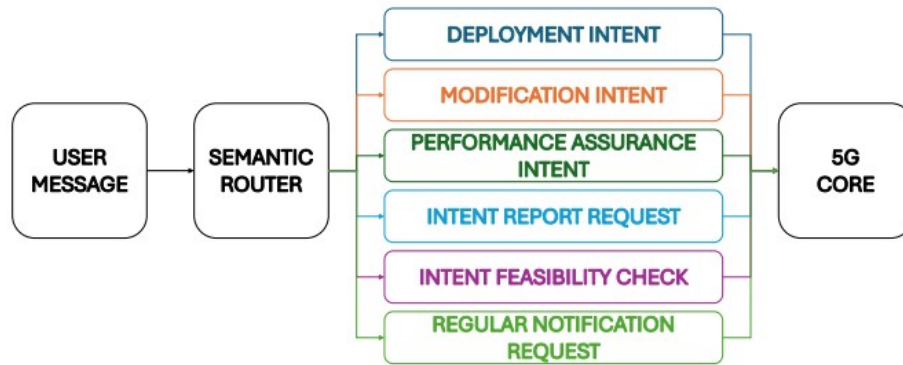
Entretanto, essa mesma versatilidade traz limitações. Os modelos podem incorrer em erros conhecidos como alucinações, quando geram informações falsas ou imprecisas, além de serem sensíveis à forma como o usuário formula a pergunta (XU *et al.*, 2024). Dessa forma, torna-se necessário o uso de estratégias adicionais para controlar, direcionar e enriquecer a interação, como o Roteador Semântico, o *Prompt Chaining* e os Bancos de Dados Vetoriais.

2.3 Roteador Semântico

Os LLMs não operam como sistemas monolíticos. Em arquiteturas contemporâneas, como a do Gemini, diferentes modelos especializados podem ser acionados conforme a natureza da tarefa (ANIL *et al.*, 2023). Para que essa seleção ocorra de maneira eficiente, emprega-se o Roteador Semântico.

Esse mecanismo, conforme exemplo ilustrado na Figura 4, realiza a análise da mensagem do usuário, interpretando seu contexto por meio de técnicas de Processamento de Linguagem Natural (PLN), e a direciona ao modelo de LLM ou ao Banco de Dados Vetorial mais apropriado (MANIAS *et al.*, 2024). De modo análogo a um sistema de atendimento telefônico, no qual diferentes setores são responsáveis por demandas específicas, o roteamento semântico estrutura a comunicação, previne sobrecargas nos modelos principais e encaminha solicitações de menor complexidade para modelos mais leves e economicamente viáveis.

Figura 4 – Exemplo de Roteador Semântico.



Fonte: (MANIAS *et al.*, 2024).

2.4 Prompt Chaining

Após o processo de roteamento, torna-se frequentemente necessário transformar ou enriquecer o *prompt* antes de seu envio ao modelo de LLM final. Essa tarefa é desempenhada pela técnica conhecida como *Prompt Chaining*, a qual consiste em reestruturar a entrada do usuário de modo a minimizar ambiguidades, corrigir inconsistências e incorporar informações contextualmente relevantes (FRANÇA *et al.*, 2025).

Por exemplo, em um sistema de geração de mídia digital, abrangendo imagens e vídeos, podem coexistir modelos especializados para cada modalidade. Assim, de acordo com a natureza da solicitação do usuário, caso esta envolva a criação de imagens geradas por IA, o *Prompt Chaining* realiza a reformulação do pedido, adicionando instruções complementares, como a correção ortográfica da requisição, a inclusão de restrições para evitar elementos potencialmente inadequados e outras diretrizes destinadas a otimizar a qualidade e a segurança do resultado gerado.

Na aplicação do UFChat, essa etapa é particularmente útil para integrar o modelo com informações institucionais armazenadas em arquivos do sistema. O *Prompt Chaining* reestrutura o *prompt* do modelo adicionando novas diretrizes pré-definidas e informações. O algoritmo de *Prompt Chaining* recupera informações no banco de informações do sistema, o Banco de Dados Vetorial, a partir de um algoritmo de busca utilizando a requisição do usuário. Com a informação recuperada e a requisição do usuário, um novo *prompt* é produzido para ser entregue ao modelo de LLM final.

2.5 Banco de Dados Vetorial e *Retrieval Augmented Generation* (RAG)

Conforme citado na subseção 2.4, as informações que o modelo de LLM utiliza são mantidas em Bancos de Dados Vetoriais, projetados para armazenar *embeddings*. Diferentemente dos bancos relacionais, os Bancos Vetoriais possuem maior eficiência em gerenciar vetores e realizam buscas baseadas em similaridade vetorial, utilizando cálculos como a Similaridade do Cosseno (subseção 2.5.1) para buscar informações mais próximas ao contexto da consulta (TAIPALUS, 2024).

O *Retrieval Augmented Generation* (RAG) consiste em incrementar a resposta de uma IA a partir de uma fonte de informações, buscando informações externas por meio de busca de similaridade semântica em um Banco de Dados. No contexto de LLMs, informações textuais podem ser transformadas em vetores e armazenadas em Bancos de Dados Vetoriais. A partir da consulta do usuário, os conteúdos mais relevantes são identificados e retornados. Essas informações são então encaminhadas para o *Prompt Chaining*, onde são adicionadas ao *prompt* final antes de serem submetidas ao modelo de LLM final. Isso promove a confiabilidade do sistema, permitindo ao UFChat oferecer respostas com informações sempre alinhadas à realidade institucional (LEWIS *et al.*, 2020).

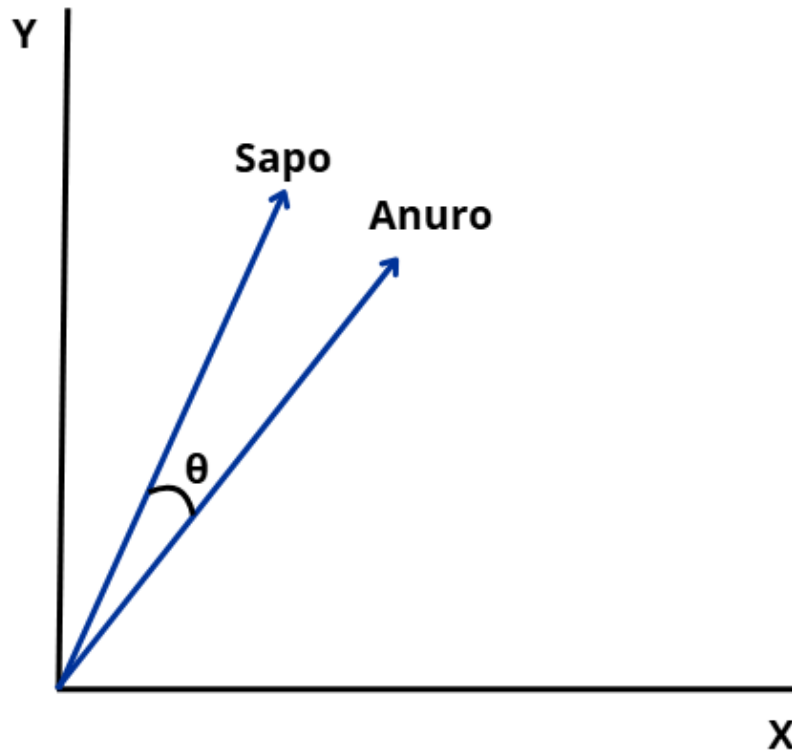
2.5.1 Similaridade do Cosseno

A busca de informações em um Banco de Dados Vetorial pode ser feita a partir do cálculo da Similaridade do Cosseno. Calculando o cosseno do ângulo formado por dois vetores, é possível medir sua proximidade semântica, pois quanto mais próximos semanticamente, menor é o ângulo entre eles e, consequentemente, maior o cosseno. Como mencionado anteriormente, os vetores de *embeddings* possuem alta proximidade quando sua similaridade semântica também é alta. Portanto, isso pode ser aproveitado, uma vez que a pergunta de um usuário geralmente compartilha alta semelhança semântica com sua resposta.

Na Figura 5, dois vetores semelhantes estão próximos um do outro, portanto o cosseno do ângulo formado entre esses dois vetores é maior que o de dois vetores semanticamente diferentes.

O cálculo para identificar o cosseno do ângulo entre estes vetores é descrito pela Equação 2.1. Valores mais próximos de 1 indicam maior similaridade semântica entre os vetores, enquanto valores próximos de 0 indicam baixa relação. Essa medida pode ser utilizada para

Figura 5 – Exemplo de vetores de *embeddings* no espaço vetorial.



Elaborado pelo autor.

identificar os documentos mais próximos da consulta do usuário.

$$\text{simil}(A, B) = \frac{A \cdot B}{\|A\| \|B\|} = \frac{\sum_{i=1}^n A_i B_i}{\sqrt{\sum_{i=1}^n A_i^2} \sqrt{\sum_{i=1}^n B_i^2}} \quad (2.1)$$

2.5.2 Obtenção de Dados por meio de Web Scraping

Para que o Banco de Dados Vetorial possua o conhecimento atualizado, é necessário que o sistema possa realizar a busca de informações de maneira autônoma. Uma das formas mais comuns de adquirir uma base de conhecimento atualizada para LLMs é a partir de *Web Scraping*. Esta técnica consiste em utilizar programas que vasculham sites para obter informações, extraindo textos, imagens, links e outros dados importantes que possam ser adicionados ao conhecimento do sistema. O *Common Crawl*, um dos maiores arquivos digitais de *Web Scraping*, é um dos principais fornecedores de informações para grandes modelos de LLM (Mozilla Foundation, 2024).

O *Web Scraping* funciona vasculhando o código que constrói os sites. Este código, construído em *Hypertext Markup Language* (HTML), contém os dados e informações que

serão extraídos pelo programa. Porém, esta prática exige que os sites permitam *Web Scraping*, geralmente informado em links do próprio site (Columbia University Mailman School of Public Health, 2025).

Dessa forma, o *Web Scraping* contribui para manter o Banco de Dados Vetorial atualizado com informações relevantes.

2.6 Chatbots

A convergência dos LLMs, Roteadores Semânticos, *Prompt Chaining* e RAG dá origem a IAs Generativas conversacionais mais robustas e especializadas. Esses componentes, quando integrados, compõem a arquitetura de um *chatbot*, um agente inteligente capaz de compreender intenções, recuperar informações relevantes e gerar respostas contextualizadas.

O Roteador Semântico atua como uma ferramenta de triagem de solicitações, identificando o tipo de pergunta e direcionando-a para o contexto mais adequado. Em seguida, o *Prompt Chaining* organiza e complementa a entrada do usuário, estruturando as instruções e preparando a consulta para as etapas subsequentes. Nesse processo, o RAG busca conteúdo relevante em um Banco de Dados Vetorial com base na solicitação do usuário, recuperando dados externos atualizados, como documentos institucionais, editais ou informações disponíveis no site. As informações recuperadas são então integradas ao *prompt* final, construída por meio do *Prompt Chaining* e submetida ao modelo LLM, permitindo a geração de respostas contextualizadas, precisas e claras, semelhantes a um diálogo humano (CALDARINI *et al.*, 2022).

No caso do UFChat, essa integração permite que o *chatbot* interaja de maneira fluida com a comunidade acadêmica, oferecendo respostas oficiais, confiáveis e adaptadas ao contexto de cada solicitação, ao mesmo tempo em que mantém um tom conversacional e acessível.

2.6.1 Memória por meio de Conversation Buffer Memory

Por mais que o *chatbot* seja capaz de realizar tarefas complexas e gerar textos similares à comunicação humana, ainda é necessário que esses assistentes sejam capazes de compreender o contexto da conversa. Quando um usuário realiza uma requisição à um LLM, cada requisição é tratada como uma nova entrada, sem referências à conversa anterior. Isto gera uma desconexão entre as requisições, que pode fazer com que o assistente não seja capaz de completar sua tarefa ou entender toda a necessidade ou problema do usuário.

Para lidar com essa limitação de ausência de contexto entre requisições, é necessário implementar um mecanismo de memória para LLMs. Dessa forma, o sistema passa a ser capaz de armazenar o histórico de diálogo e permitir que o assistente faça referência a tópicos ou informações mencionadas anteriormente na conversa. Na biblioteca LangChain, projetada para estender funcionalidades de sistemas de IA, existem componentes que possibilitam o armazenamento das interações entre usuário e assistente. Nesse processo, é possível definir quais mensagens serão mantidas no histórico, identificando se pertencem ao usuário ou ao assistente, bem como a quantidade de mensagens armazenadas, permitindo que o *prompt* final utilize esse histórico para fornecer mais contexto ao modelo (BRIGGS; INGHAM, 2025; TEAM, 2025).

2.7 Métodos de Avaliação

A avaliação de assistentes baseados em LLM pode ser realizada por meio de abordagens quantitativas e qualitativas. A avaliação quantitativa pode utilizar métricas específicas para mensurar o desempenho do RAG, como as fornecidas pelo *framework* RAGAS. Já a avaliação qualitativa pode ser conduzida por meio de testes A/B com usuários.

2.7.1 Avaliação RAGAS

Para avaliar de forma rigorosa um *chatbot* baseado em RAG, é essencial utilizar métricas que mensurem tanto a qualidade da resposta quanto a eficácia da etapa de recuperação de documentos. A *Hugging Face* publicou o *framework* RAGAS, que reúne métricas específicas para esse tipo de sistema (ES *et al.*, 2024; Ragas Team, 2024). Entre as principais métricas utilizadas, destacam-se:

- **Faithfulness**: avalia o quanto a resposta gerada é fiel às evidências presentes nos documentos recuperados. Essa métrica verifica se as afirmações contidas na resposta podem ser sustentadas pelo conteúdo dos documentos utilizados como contexto, ajudando a identificar possíveis alucinações do modelo.
- **Answer Correctness**: mede o grau de correção da resposta gerada em relação a uma resposta de referência previamente definida. Essa métrica considera se a resposta fornecida pelo modelo está semanticamente alinhada com a resposta esperada.
- **Context Precision**: verifica a relevância dos documentos recuperados pelo mecanismo de busca do RAG. Essa métrica analisa se os trechos de contexto fornecidos ao modelo

realmente contribuem para responder à pergunta realizada pelo usuário.

- **Context Recall:** mede se o mecanismo de busca do RAG conseguiu recuperar documentos suficientes e relevantes para produzir a resposta correta. Em outras palavras, avalia se informações importantes presentes na base de conhecimento foram efetivamente recuperadas.

Em conjunto, essas métricas permitem avaliar de forma abrangente a qualidade do pipeline RAG, considerando tanto a qualidade das respostas geradas quanto a eficácia do mecanismo de recuperação de informações. Seus valores podem variar entre 0 e 1, sendo melhor quanto mais próximo a 1.

2.7.2 *Teste A/B*

Além da avaliação do RAG, também é necessário analisar a experiência dos usuários com o assistente. Para isso, pode-se utilizar pesquisas de teste A/B para comparar a qualidade das respostas geradas pelo assistente com aquelas produzidas por outros modelos, conforme realizado em (FRANÇA *et al.*, 2025). Por meio desse tipo de avaliação também é possível analisar a satisfação dos usuários, identificar críticas e verificar se o assistente cumpre adequadamente seu objetivo.

No contexto de assistentes baseados em LLM, o teste A/B permite avaliar se as respostas geradas pelo sistema proposto são comparáveis ou superiores às produzidas por outros assistentes já consolidados (Oracle, 2024). Esse tipo de avaliação também possibilita coletar percepções dos usuários sobre a qualidade das respostas, incluindo satisfação, críticas e sugestões, bem como verificar se o assistente cumpre adequadamente o seu propósito.

Assim, as respostas geradas pelo UFChat e por outros assistentes, como o Google Gemini e o ChatGPT, são apresentadas aos participantes por meio de uma pesquisa. Com base nessa comparação, os usuários indicam quais respostas consideram mais adequadas ao contexto acadêmico. Dessa forma, torna-se possível mensurar a preferência dos usuários e avaliar o desempenho do assistente proposto em relação a outros amplamente utilizados.

3 TRABALHOS RELACIONADOS

Modelos baseados em LLM têm sido amplamente explorados como assistentes conversacionais desde sua popularização. Um exemplo é o estudo conduzido pela Unimib (ANTICO *et al.*, 2024), no qual o ChatGPT foi utilizado como suporte a estudantes de psicologia. O assistente era capaz de responder a dúvidas acadêmicas, fornecer informações sobre a universidade e a cidade, além de orientar sobre custos de estadia, alimentação e horários. Apesar dos resultados positivos, o trabalho evidencia limitações ligadas a alucinações e à possibilidade de transmitir informações incorretas.

No contexto brasileiro, iniciativas semelhantes também têm surgido. Em (SANTOS *et al.*, 2024), por exemplo, foi desenvolvido um modelo de atendimento voltado para uma Secretaria Acadêmica do *Instituto Federal de Santa Catarina* (IFSC). O foco do projeto está em auxiliar ingressantes diante da complexidade dos processos universitários e da presença de termos muitas vezes desconhecidos, como egresso, ingresso e grade curricular. Nessa perspectiva, um LLM atua como mediador acessível, traduzindo burocracias institucionais em linguagem clara e compreensível.

Outro estudo relevante é apresentado por (JEZLER *et al.*, 2025), no qual o UNESC aplicou *chatbots* para otimizar a comunicação com alunos. Os autores destacam que as dificuldades recorrentes de compreensão por parte dos estudantes acabam sobrecarregando secretarias e setores administrativos, problema que pode ser mitigado com soluções automatizadas. Esse caso reforça que a adoção de LLMs não se restringe ao atendimento básico de dúvidas, mas pode contribuir diretamente para a redução da carga de trabalho institucional.

A aplicação de LLMs também vem sendo investigada em contextos distintos do atendimento acadêmico. Na Universidade de Siegen, Alemanha (ABU-RASHEED *et al.*, 2024), foi proposto um modelo que intermedeia a comunicação entre alunos e mentores. O assistente organiza as perguntas enviadas pelos estudantes em questionários otimizados, que são analisados pelos mentores. Para enriquecer as respostas, o sistema integra informações em um *Knowledge Graph*, resultando em planos de estudo personalizados. A avaliação qualitativa apontou boa aceitação tanto pelos alunos quanto pelos mentores, com média mínima de quatro em seis dimensões analisadas.

Apesar dos avanços, nota-se desafios recorrentes nas obras citadas similares aos problemas apontados por (WILLIAMS; HUCKLE, 2024), que se destacam: LLMs tendem a gerar informações falsas (conhecidas por alucinações), possuem limitação de conhecimento,

carecem de memória entre requisições e necessitam de uma requisição bem detalhada para produção de uma resposta satisfatória. Além disso, algumas obras apresentam apenas avaliações qualitativas dos assistentes, com ausência de avaliações quantitativas.

Nesse contexto, destacam-se propostas como a de (FRANÇA *et al.*, 2025), que explora o uso de LLMs por assistentes comunitários de saúde para a prevenção de doenças crônicas não transmissíveis em crianças com menos de mil dias de vida. O estudo se diferencia pela aplicação de técnicas como *Prompt Chaining*, *Prompt Engineering* e métricas quantitativas, incluindo a Similaridade do Cosseno. Complementarmente, (GUAN *et al.*, 2025), desenvolvido por pesquisadores da Microsoft, propõe uma visão integrada do processo avaliativo de assistentes baseados em LLM, abordando desde o planejamento das métricas e critérios de desempenho até a mensuração dos resultados obtidos por cada módulo do sistema.

Em conjunto, essas abordagens fornecem uma base metodológica consistente para este trabalho, orientando tanto o desenvolvimento quanto a avaliação do UFChat, especialmente no que diz respeito à adoção de estratégias quantitativas e interpretáveis para análise de sua eficácia no contexto acadêmico.

4 MATERIAIS E MÉTODO PROPOSTO

Neste capítulo, apresenta-se a estrutura de funcionamento do UFChat, destacando como suas funcionalidades se interconectam para cumprir o propósito estabelecido.

4.1 Materiais

Antes de descrever o método proposto em detalhes, foi realizada um levantamento de necessidades da comunidade acadêmica para identificar as principais demandas acadêmicas e institucionais que nortearam o desenvolvimento do método.

4.1.1 *Levantamento de Necessidades da Comunidade Acadêmica*

Para garantir que o assistente atenda às reais demandas da comunidade acadêmica, foi aplicado um questionário enviado ao canal institucional de comunicação da UFCA (dialogol.todos@UFCA.edu.br). O instrumento buscou identificar tanto o perfil dos respondentes quanto suas expectativas em relação ao uso de ferramentas de IA no ambiente acadêmico. As questões encontram-se resumidas na Tabela 1.

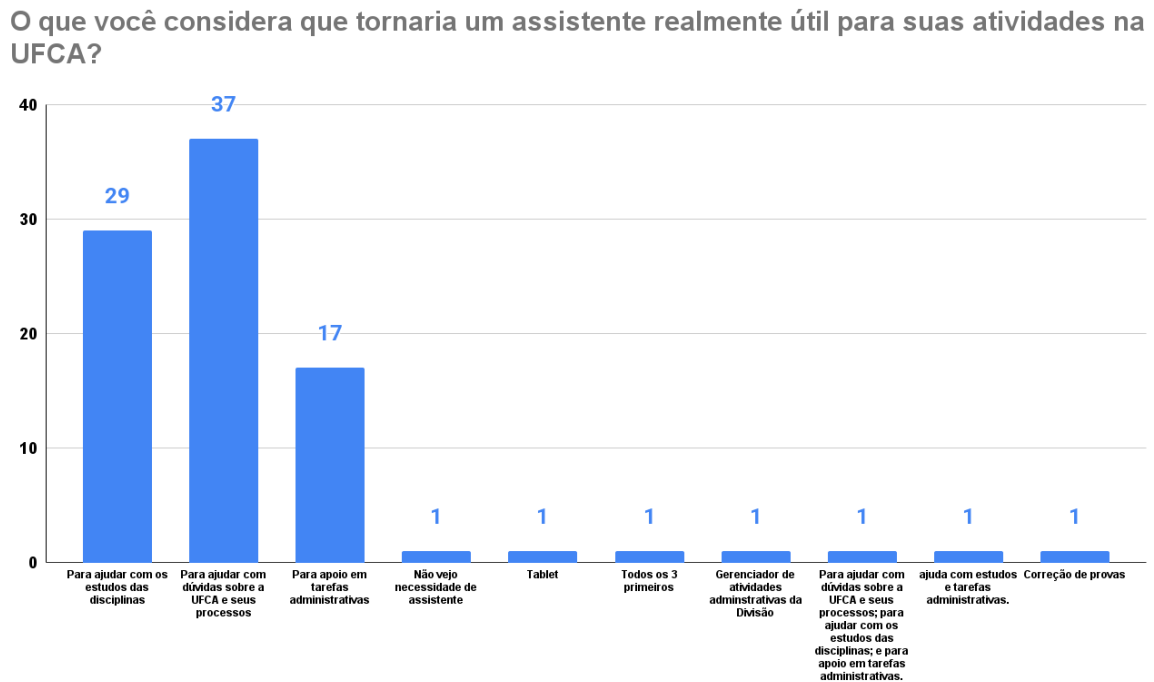
Tabela 1 – Questionário aplicado para levantamento de necessidades da comunidade acadêmica.

Nº	Pergunta
1	Identificação
2	Qual curso ou setor você pertence?
3	Você já utilizou alguma ferramenta baseada em IA?
4	Se sim, quais ferramentas de IA você já usou?
5	O que você considera que tornaria um assistente realmente útil para suas atividades na UFCA?
6	Como você prefere que seja o acesso ao assistente?
7	Gostaria de deixar alguma sugestão?

Fonte: Elaborado pelo autor.

A pesquisa levantou respostas de 90 respondentes. Dentre eles, 63 eram alunos de graduação, 2 de pós graduação, 13 técnicos e 12 professores. Entre as perguntas, destaca-se a de número 5, considerada essencial para a definição do escopo do assistente. Os resultados dessa questão estão representados na Figura 6.

Figura 6 – Distribuição das respostas à pergunta 5 do questionário.

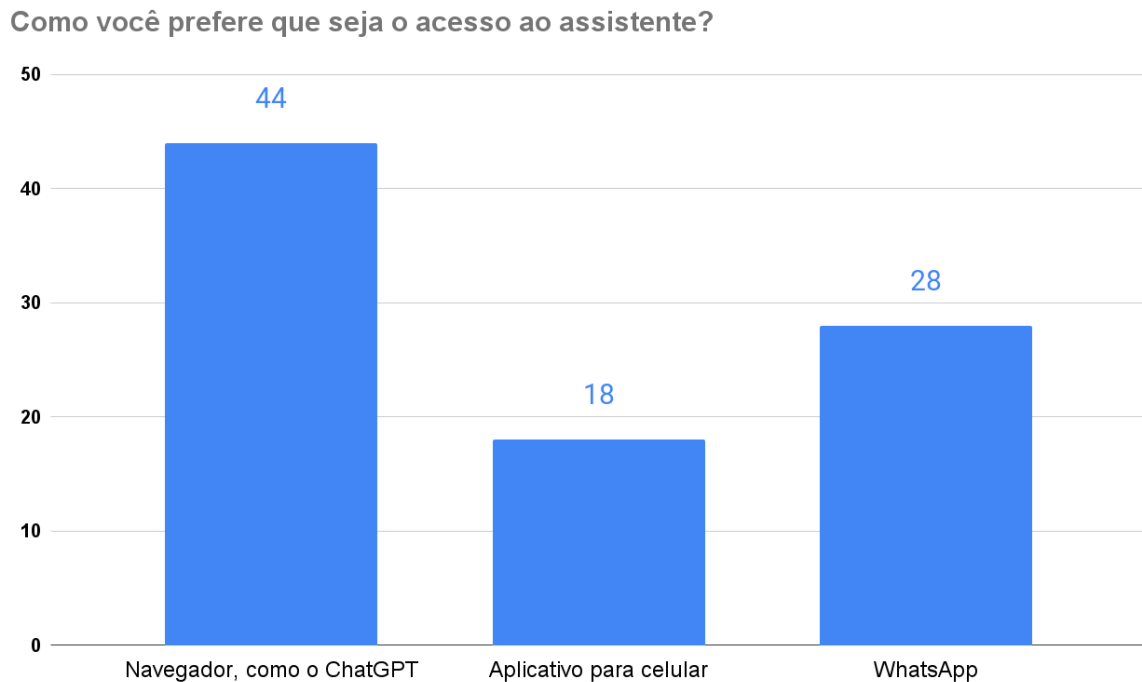


Fonte: Elaborado pelo autor.

Os dados revelam que 37 (41,1%) dos participantes apontaram como prioridade a função de “*ajudar com dúvidas sobre a UFCA e seus processos*”, enquanto 32,2% destacaram a utilidade de um assistente voltado para “*apoio aos estudos das disciplinas*”. Dessa forma, o escopo do modelo será orientado principalmente para o atendimento de dúvidas institucionais e acadêmicas, com ênfase no suporte às necessidades cotidianas dos estudantes.

É necessário também definir a alocação do UFChat e sua interface. Conforme as respostas da pergunta 6, visíveis na Figura 7, 44 (48,9%) de 90 respondentes preferem a interface similar ao ChatGPT, em navegador. 18 (20%) preferem que o UFChat seja adaptado em um aplicativo móvel, e 28(31,1%) preferem que seja pelo WhatsApp. Considerando que o WhatsApp possa ser utilizado em navegador e em aplicativo de celular, além de ser a rede social mais utilizada no Brasil em 2025 (DOURADO, 2025; MLABS, 2025), esta opção será a mais benéfica para todos, e será a escolhida para o UFChat.

Figura 7 – Distribuição das respostas à pergunta 6 do questionário.



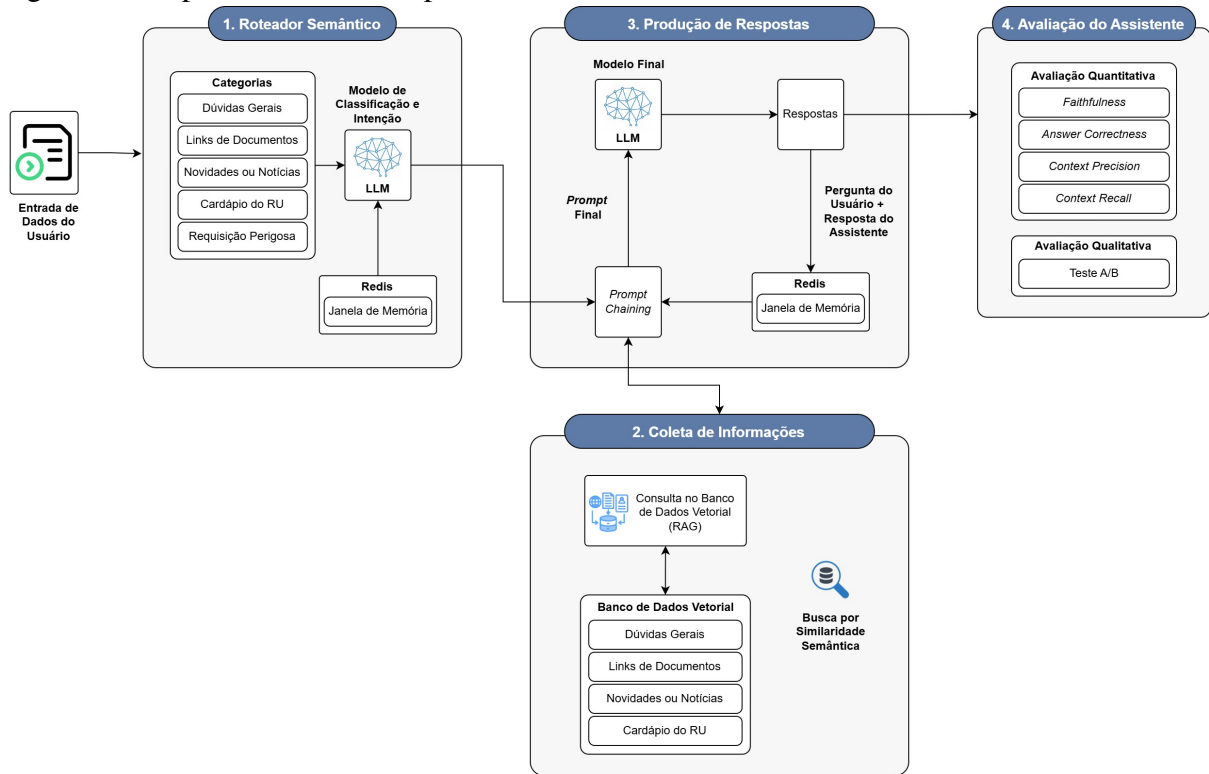
Fonte: Elaborado pelo autor.

4.2 Método Proposto

Definido o objetivo do UFChat, segue-se para a apresentação da estrutura do assistente, conforme ilustrado na Figura 8.

Inicialmente, a requisição do usuário é enviada ao servidor do assistente. A partir desse ponto, o método proposto executa quatro etapas principais. Na etapa de Roteador Semântico, um modelo de LLM referido por Modelo de Classificação e Intenção analisa a entrada do usuário, classifica em uma das categorias definidas e gera um resumo da requisição identificando sua intenção. Após essa etapa, as informações são encaminhadas ao *Prompt Chaining*. Nesse processo, o mecanismo RAG entra em operação, realizando a etapa de Coleta de Informações por meio de uma busca por similaridade semântica em Bancos de Dados Vetoriais, com o objetivo de recuperar conteúdos relevantes relacionados ao resumo da requisição do usuário. Posteriormente, na etapa de Produção de Respostas, utiliza-se o *Prompt Chaining* para construir o *prompt* final, combinando a pergunta do usuário, as informações recuperadas e o histórico recente da conversa, que é então enviado ao modelo de LLM final referido por Modelo Final para gerar a resposta. Por fim, na etapa de Avaliação do Assistente, são aplicados indicadores que permitem avaliar a

Figura 8 – Etapas do Método Proposto.



Fonte: Elaborado pelo autor.

qualidade das respostas geradas pelo assistente.

4.3 Roteador Semântico

O processo inicia com o envio de uma mensagem pelo usuário ao servidor do UFChat, que receberá a mensagem e a encaminhará para a primeira etapa do método, denominada Roteador Semântico.

O Roteador Semântico é um mecanismo responsável por realizar uma triagem a partir do contexto da requisição do usuário. No momento em que a mensagem do usuário chega a esta etapa, ela é interpretada pelo Modelo de Classificação e Intenção, que fornece uma categoria a partir de uma lista bem definida. As categorias possíveis estão apresentadas na Tabela 2 junto de seus contextos pré-definidos. Estes contextos servem como explicações ao modelo para ajudá-lo a escolher melhor a categoria.

Tabela 2 – Categorias possíveis e seus respectivos contextos pré-definidos.

Categoria	Contexto Pré-Definido
Dúvidas Gerais e Contatos	O usuário deseja informações gerais da Universidade, dos cursos, do SIGAA, das suas partes constituintes como proreitorias ou cotidianas.
Links de Editais	O usuário deseja documentos ou links de documentos como editais, convocações e formulários.
Cardápio do RU	O usuário deseja o cardápio do Restaurante Universitário ou informações sobre a alimentação semanal do Restaurante Universitário.
Novidades ou Notícias	O usuário deseja as notícias recentes ou novidades da UFCA.
Requisição Perigosa	O usuário está solicitando informações de membros da UFCA, como alunos ou professores, ou solicitando uma informação que pode por alguém em risco.

Fonte: Elaborado pelo autor.

A categorização é fundamental para o funcionamento do método proposto. Por exemplo, a categoria Dúvidas Gerais aborda as perguntas mais frequentes sobre a UFCA e suas atividades acadêmicas, enquanto a categoria Links de Editais são destinados a consultas de documentos oficiais, como editais de convocação para heteroidentificação, disponibilizados pela universidade. Esta categorização é importante para a busca de informações no Banco de Dados Vetorial, pois o mesmo está dividido em categorias idênticas às que o modelo pode selecionar.

O Modelo de Classificação e Intenção também deve produzir um resumo da intenção do usuário para que possa ser utilizado nas próximas etapas. O funcionamento deste modelo depende de diretrizes bem definidas, para que não fuja da lista de categorias proposta nem gere resumos fora do assunto em discussão. Ele recebe, no campo *system*, uma diretriz que define seu comportamento, e, no campo *user*, a lista de possíveis categorias. Com base nisso, o modelo identifica a categoria mais adequada ao contexto definido pelo desenvolvedor (Tabela 2) e gera um resumo em primeira pessoa, simulando a intenção do usuário. Em seguida, o sistema envia a requisição, o resumo e a categoria à etapa de Coleta de Informações, que realiza a busca por Similaridade do Cosseno utilizando o vetor de *embeddings* do resumo como parâmetro. O processo é exemplificado na Tabela 3.

Tabela 3 – Fluxo para classificação e resumo do Modelo de Classificação e Intenção.

Campo <i>system</i>	Você é um classificador de resumos e categorias. Você deve ler o histórico e a última mensagem do usuário e gerar um resumo em primeira pessoa, como se você fosse o usuário. Em seguida, você deverá escolher uma categoria para o resumo utilizando APENAS as Categorias fornecidas. Retorne no formato: Resumo: \ Categoria: Exemplo: Resumo: Quero o link do Edital \ Categoria: Link Edital
Campo <i>user</i>	Categorias: "Dúvidas Gerais e Contatos", "Links de Editais", "Novidades ou Notícias", "Cardápio do Dia", "Requisição Perigosa" Última mensagem do usuário: Onde encontro a DTI?
Classificação (Triagem)	Dúvidas Gerais e Contatos
Resumo	Quero saber o endereço da DTI.

Fonte: Elaborado pelo autor.

Porém, este processo possui uma limitação quando não há memória embutida ao assistente. Como todas as mensagens dos usuários são processadas pelo Modelo de Classificação e Intenção, nem sempre o contexto (o tópico que estão conversando) pode estar claro, e necessita de mensagens anteriores para que o modelo possa compreender melhor o que o usuário está solicitando. Por exemplo, ao perguntar sobre a pró-reitoria responsável por projetos de extensão, o assistente pode perguntar se deseja o contato, ao qual o usuário pode aceitar dizendo “Quero”, mas o modelo não compreenderá sem o contexto completo (pois é enviado apenas a última mensagem “Quero” com as outras informações no campo *user*), conforme demonstrado na Tabela 4.

Tabela 4 – Fluxo para classificação e resumo do Modelo de Classificação e Intenção.

Campo system	Você é um classificador de resumos e categorias. Você deve ler o histórico e a última mensagem do usuário e gerar um resumo em primeira pessoa, como se você fosse o usuário. Em seguida, você deverá escolher uma categoria para o resumo utilizando APENAS as Categorias fornecidas. Retorne no formato: Resumo: Categoria: Exemplo: Resumo: Quero o link do Edital Categoria: Link Edital
Campo user	Categorias: "Dúvidas Gerais e Contatos", "Link Edital", "Novidades ou Notícias", "Cardápio do Dia" Última mensagem do usuário: Quero
Classificação (Triagem)	Dúvidas Gerais e Contatos
Resumo	Quero.

Fonte: Elaborado pelo autor.

Assim, o resumo que será utilizado em etapas seguintes se torna descontextualizado, e acarretará em problemas.

4.3.1 Memória do Assistente

Para lidar com essa limitação, é introduzida a memória do assistente, responsável por armazenar o histórico recente de interação entre usuário e sistema. Essa memória armazena as últimas três interações completas entre usuário e assistente em um banco de dados Redis (Redis Ltd., 2024), um banco de dados em memória de alta performance amplamente utilizado em aplicações que demandam baixa latência.

As mensagens são armazenadas em pares, sendo cada par composto por uma mensagem do usuário e a respectiva resposta gerada pelo Modelo Final. Esse histórico permite ampliar o contexto disponível para a geração de novas respostas, possibilitando que informações previamente mencionadas sejam consideradas pelo sistema.

Para limitar o uso de *tokens* e manter apenas uma memória de curto prazo, são armazenados no máximo três pares de mensagens por usuário (seis mensagens no total), conforme abordagem semelhante à utilizada por (TEAM, 2025). Essa estrutura será denominada neste trabalho como Janela de Memória.

A Janela de Memória é então inserida no campo Histórico do *prompt* enviado ao Modelo de Classificação e Intenção, permitindo que esse modelo considere o contexto recente

da conversa ao produzir um resumo da intenção do usuário e realizar sua classificação.

Tabela 5 – Fluxo para classificação e resumo do Modelo de Classificação e Intenção.

Campo system	Você é um classificador de resumos e categorias. Você deve ler o histórico e a última mensagem do usuário e gerar um resumo em primeira pessoa, como se você fosse o usuário. Em seguida, você deverá escolher uma categoria para o resumo utilizando APENAS as Categorias fornecidas. Retorne no formato: Resumo: Categoria: Exemplo: Resumo: Quero o link do Edital Categoria: Link Edital
Campo user	Categorias: "Dúvidas Gerais e Contatos", "Link Edital", "Novidades ou Notícias", "Cardápio do Dia" Histórico: Usuário: Olá UFChat! UFChat: Olá! Como posso te ajudar hoje? Usuário: Qual pró-reitoria é responsável pelos projetos de extensão? UFChat: A pró-reitoria responsável por projetos de extensão é a Proex. Você gostaria de entrar em contato? Última mensagem do usuário: Quero
Classificação (Triagem)	Dúvidas Gerais e Contatos
Resumo	Quero o contato da Proex.

Fonte: Elaborado pelo autor.

A requisição original, a classificação e o resumo produzidos pelo Modelo de Classificação e Intenção são então enviadas à etapa de Coleta de Informações por meio do *Prompt Chaining*, que conecta as etapas de Roteador Semântico, Coleta de Informações e Produção de Respostas.

4.4 Coleta de Informações

Aqui inicia-se a primeira etapa do RAG do sistema, em que uma informação extra é buscada e adicionada à requisição para que o Modelo Final possa produzir uma resposta mais robusta. A Coleta de Informações é o mecanismo de busca do RAG, que busca no Banco de Dados Vetorial. No UFChat, o Banco de Dados Vetorial foi dividido em diferentes seções para cada categoria do Roteador Semântico.

Estas seções possuem os vetores de *embeddings* de informações que foram coletadas

a partir do site oficial da UFCA (Universidade Federal do Cariri, 2025b) e da Wiki UFCA (UFCA, 2025), tutoriais de utilização do SIGAA e links extraídos de sites e portais de notícias oficiais por meio de *Web Scraping* e produção manual. Cada categoria possui arquivos diferentes, pois necessitam de frequências de atualização diferentes. A frequência de atualizações estão disponíveis na Tabela 6.

Tabela 6 – Frequência de atualização do Banco de Dados Vetorial.

Dúvidas Gerais e Contatos	A cada 3 meses
Links de Editais	Mensalmente
Cardápio do RU	Semanalmente
Novidades ou Notícias	Diariamente

Fonte: Elaborado pelo autor.

O modelo *Text Embedding Small 3* da OpenAI, foi utilizado como ferramenta para gerar os vetores *embeddings* (OPENAI, 2025). O resumo, gerado pelo Roteador Semântico, é transformado em um vetor e, por meio do cálculo da similaridade do cosseno, comparado ao conteúdo presente no Banco de Dados Vetorial. O conteúdo mais próximo é retornado pelo mecanismo de busca do RAG para a reformulação do *prompt*, realizada pelo *Prompt Chaining*.

O resumo produzido pelo Modelo de Classificação e Intenção é utilizado como consulta para o mecanismo de busca vetorial. Isso ocorre porque o processo de recuperação de informações é executado para todas as requisições do sistema. Caso a consulta seja realizada diretamente a partir da mensagem original do usuário, podem ocorrer buscas com baixo contexto semântico, especialmente em respostas curtas ou dependentes de interações anteriores, como “Quero”, “Sim” ou “Olá”. Nesses casos, o resumo gerado pelo modelo fornece uma representação mais completa da intenção do usuário, permitindo uma recuperação de documentos mais precisa.

4.5 Produção de Respostas

Com as informações recuperadas do Banco de Dados Vetorial, a requisição original do usuário e a Janela de Memória, um novo *prompt* será montado com diretrizes, informações sobre o objetivo do UFChat e a data da requisição em seu campo *system*. Uma requisição original, como por exemplo “Onde encontro a DTI?”, é transformada em um *prompt* completo

com informações adicionais e a Janela de Memória no campo *user*, pronto para ser enviado ao Modelo Final, conforme apresentado na Tabela 7.

Tabela 7 – Fluxo para geração de resposta pelo Modelo Final.

Campo system	Você é o UFChat, um assistente virtual desenvolvido como projeto de TCC para atender à Universidade Federal do Cariri (UFCA). Seu papel é agir como um profissional de atendimento das secretarias da UFCA, respondendo de forma clara, humana, educada e profissional. Você não deve fornecer informações sobre seu criador, e-mail, endereço ou qualquer dado pessoal. Se for perguntado sobre isso, diga apenas que é um projeto acadêmico da UFCA e retorne o foco para o atendimento institucional. Hoje é: 23-09-2025.
Campo user	Instrução: Continue a conversa, respondendo à requisição utilizando o contexto. Histórico: • Usuário: Olá UFChat! • Assistente: Olá! Como posso te ajudar? Requisição do Usuário: Onde encontro a DTI?
Classificação (Triagem)	Dúvidas Gerais e Contatos
Resumo	Quero o endereço da DTI
Informação recuperada	A Diretoria de Tecnologia da Informação (DTI) é a unidade responsável pela área de TI na Universidade Federal do Cariri (UFCA). Está subordinada à Reitoria e alinhada ao Comitê de Governança de TI. Localização: Campus Juazeiro do Norte, Bloco J, Sala 01. Email: atendimento.dti@UFCA.edu.br Telefone: (88) 3221-9400.
Resposta	A Diretoria de Tecnologia da Informação (DTI) da UFCA está localizada no Campus Juazeiro do Norte, no Bloco J, Sala 01. Você também pode entrar em contato pelo e-mail atendimento.dti@UFCA.edu.br ou pelo telefone (88) 3221-9400.

Fonte: Elaborado pelo autor.

Após a geração da resposta ao usuário, a interação é armazenada na Janela de Memória. Como a memória mantém apenas os três últimos pares de mensagens, quando sua capacidade é atingida, o par mais antigo é removido para a inserção do novo.

4.6 Avaliação do Assistente

Para que o assistente seja devidamente avaliado, as suas tarefas e funcionalidades serão submetidas a avaliações quantitativas e qualitativas (Seção 2.7). A avaliação quantitativa será realizada por meio do *framework* RAGAS, responsável pela coleta de métricas relacionadas ao desempenho do modelo. Já a avaliação qualitativa será conduzida por meio de um teste A/B aplicado à comunidade universitária, no qual os participantes irão comparar respostas geradas pelo assistente desenvolvido com respostas produzidas por outros assistentes baseados em LLM.

A combinação dessas duas abordagens possibilita uma análise mais abrangente, considerando tanto métricas objetivas de desempenho quanto a experiência dos usuários ao interagir com o assistente.

4.6.1 Avaliação Quantitativa

Para a avaliação por meio do RAGAS, foram formuladas 30 requisições diferentes que poderiam ser feitas pelo corpo universitário, como dúvidas sobre endereços, cardápio, informações da universidade, etc. O assistente produziu uma resposta para cada requisição e, com cada resposta e as informações retornadas pelo mecanismo de busca do RAG, a eficiência do RAG pode ser avaliada. As métricas utilizadas foram as mesmas citadas na subseção 2.7.1.

A avaliação deve conter uma requisição do usuário, os documentos retornados pelo mecanismo de busca do RAG, a resposta gerada pelo assistente e uma resposta base de referência, chamada de *ground truth* (verdade fundamental), podendo ser gerada por humano ou máquina, para ser comparada à resposta gerada. As respostas de referência não dependeram de profissionais da UFCA para serem formuladas para este teste, pois foram extraídas diretamente do site da UFCA.

Por se tratar de uma avaliação restrita ao RAG, o RAGAS penaliza textos considerados desnecessários para geração de respostas a partir de informações recuperadas. Por exemplo, se uma requisição solicita o nome do assistente “Qual o seu nome?”, a informação recuperada traz “Seu nome é UFChat” e a resposta do Modelo Final é “Olá, meu nome é UFChat. Em que posso ajudar?” A resposta gerada será penalizada em métricas que comparam a relevância da resposta (como a *faithfulness*), devido à saudação e à oferta de ajuda. Por isso, provisoriamente, o UFChat receberá diretrizes diferentes *system*, apenas para validar o RAG, no seguinte formato: “Seu nome é UFChat. Você é um assistente de RAG. Apenas responda o solicitado, sem sauda-

ções, sem explicações extras, apenas obedeça às instruções com a informação que você recebeu.” Esta mudança de diretrizes serve apenas para captar métricas RAGAS, e não serão utilizadas para produção de respostas a usuários reais.

4.6.2 Avaliação Qualitativa

Após a coleta dos resultados pelo RAGAS, foi realizada uma segunda etapa de validação por meio de um teste A/B. Um formulário foi distribuído à comunidade universitária, no qual os participantes escolheram, entre três respostas, a que melhor atende a cada pergunta. Ao todo, 10 perguntas foram apresentadas, cada uma possuindo 3 respostas possíveis, totalizando 30 respostas. Cada resposta foi gerada por um assistente diferente: ChatGPT, Gemini e UFChat. Para reduzir vieses, as perguntas foram feitas em ambientes sem conta ou histórico de conversas, garantindo condições equivalentes entre os modelos assistentes.

As perguntas também passaram por um cálculo de Similaridade do Cosseno entre si, para que sua semântica possa ser quantificada. Isto é importante para que haja uma percepção real de que as respostas geradas pelo UFChat seja de fato diferente das geradas pelos outros sistemas, e sua escolha, caso seja preferencial, seja também justificada. Em todos os casos, os usuários não foram informados de qual LLM partiu a informação, para evitar favoritismo ou aversão a um assistente específico.

Na Tabela 8 é apresentado um exemplo de pergunta realizada aos participantes (“Qual o endereço da PRPI UFCA?”). As demais perguntas do formulário estão disponíveis no Apêndice A.

Modelo	Resposta
Resposta 1 (UFChat)	A Pró-Reitoria de Pesquisa, Pós-Graduação e Inovação (PRPI) está localizada no Campus Juazeiro do Norte, Bloco K, Sala K307. Se precisar de mais informações, estou à disposição para ajudar.
Resposta 2 (ChatGPT)	Na Universidade Federal do Cariri (UFCA), a PRPI está localizada no Campus Juazeiro do Norte. A PRPI é responsável por coordenar atividades relacionadas à pesquisa, inovação e desenvolvimento tecnológico na universidade. Para ser mais preciso, a PRPI costuma estar no Bloco D do campus, mas para informações exatas e atualizadas, recomenda-se verificar o site da UFCA ou entrar em contato com a instituição.
Resposta 3 (Gemini)	A PRPI da Universidade Federal do Cariri (UFCA) fica localizada no Câmpus Juazeiro do Norte. O endereço do campus é Rua Tenente-Coronel de Melo Montenegro, nº 110, Bairro Jardim. Para saber o bloco ou sala específicos da PRPI dentro do campus, recomenda-se consultar o site oficial da UFCA ou entrar em contato com a Pró-Reitoria.

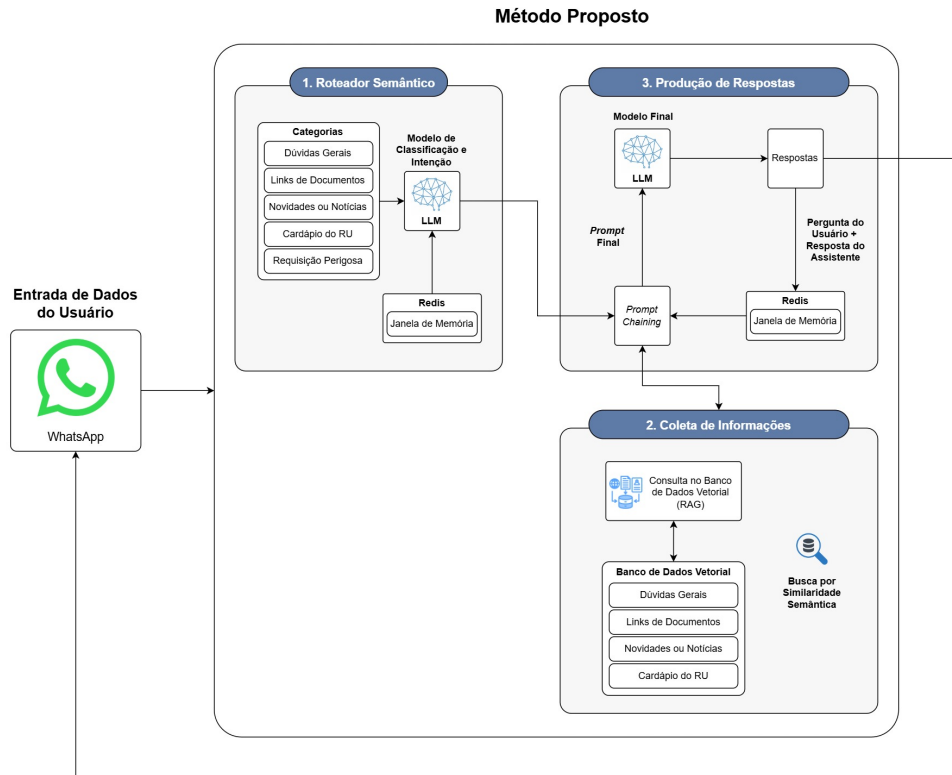
Tabela 8 – Respostas dos modelos para a pergunta: Qual o endereço da PRPI UFCA?

Observa-se que cada pergunta é acompanhada de três respostas, cada uma gerada por um assistente diferente. No exemplo apresentado, a Resposta 1 foi produzida pelo UFChat, a Resposta 2 pelo ChatGPT e a Resposta 3 pelo Gemini. É importante ressaltar que a identificação dos assistentes não foi disponibilizada aos participantes para evitar viés na avaliação. Assim, coube aos respondentes selecionar a resposta que, em sua opinião, melhor representava o comportamento esperado de um *chatbot* no atendimento ao cliente da UFCA.

4.6.3 Integração do Método Proposto ao WhatsApp

Para completar a implementação do assistente, é necessário integrá-lo ao WhatsApp por meio de uma fonte segura. A Meta, empresa responsável pela rede social, não permite conexões diretas com sua API a não ser por empresas registradas. Então, ao enviar uma mensagem para o contato do UFChat pelo WhatsApp, a mensagem será gerenciada pela Uchat, empresa registrada na Meta. A Figura 9 representa o processo.

Figura 9 – Esquema de integração com o WhatsApp.



Fonte: Elaborado pelo autor.

Ao receber a mensagem, a Uchat realiza um disparo de API para o servidor do UFChat, onde a mensagem receberá um ID com base no número de telefone do usuário para ser registrada no Banco de Dados Vetorial com a devida identificação. Ela atravessará todo o processo descrito pelo método com exceção da Avaliação de Assistente, e o servidor devolverá a resposta do Modelo Final de volta ao Uchat, que enfim encaminhará a resposta de volta ao usuário. Posteriormente, o usuário poderá dar continuidade a interação.

5 RESULTADOS E DISCUSSÕES

Este capítulo apresenta a configuração experimental e discute os principais resultados experimentais do método proposto, integração com o WhatsApp, bem como seus impactos importantes e limitações.

5.1 Configuração Experimental

O método proposto foi construído em um servidor *Virtual Private Server* (VPS) da Hostinger, uma companhia de hospedagem *web*. O servidor está localizado em São Paulo, Brasil, com sistema operacional Ubuntu 24.04 LTS. As especificações de *hardware* incluem 4 *gigabytes* de memória RAM, com 50 *gigabytes* de espaço em disco NVMe e 4 *terabytes* de largura de banda. Apenas um núcleo de CPU, sem adição de placa de vídeo (Hostinger, 2025).

Para possibilitar a integração com o WhatsApp, o servidor VPS recebe chamadas de API por meio da plataforma Uchat, que realiza conexões com servidores da Meta, empresa responsável pelo WhatsApp, Facebook, etc. Dessa forma, o Uchat permite que mensagens recebidas pelo WhatsApp do UFChat sejam redirecionadas ao sistema do assistente, processadas e retornadas ao usuário.

O Modelo de Classificação e Intenção e Modelo Final utilizados no sistema do assistente são conectados a partir da API da OpenAI, utilizando o modelo GPT-4.1 mini (OpenAI, 2025).

5.2 Resultados Experimentais

Nas próximas subseções são apresentados experimentos destinados a avaliar os principais componentes do método proposto. Inicialmente são demonstrados experimentos relacionados ao Roteador Semântico, avaliando sua capacidade de identificar corretamente a categoria de solicitações entre Dúvidas Gerais, Links de Editais, Novidades e Notícias, Cardápio do RU e Requisição Perigosa, além da qualidade das respostas geradas. Também são demonstrados experimentos voltados à redução de alucinações dos modelos.

Em seguida, são apresentadas as avaliações do assistente, que consistem em uma análise quantitativa do mecanismo RAG utilizando o *framework* RAGAS e em uma avaliação qualitativa realizada por meio de um teste A/B com usuários. Por fim, é apresentada a integração do UFChat ao WhatsApp, demonstrando sua aplicação em um ambiente real de uso.

5.2.1 Dúvidas Gerais

Este experimento avaliará como o Roteador Semântico se comporta ao receber uma entrada de dados do usuário relacionados à categoria Dúvidas Gerais (Tabela 9).

Tabela 9 – Exemplo de interação registrada como Dúvidas Gerais.

Entrada de Dados do Usuário	Qual o endereço da Procult e seu contato?
Classificação (Triagem)	Dúvidas Gerais
Resposta do Modelo Final	A Pró-Reitoria de Cultura (Procult/UFCA) está localizada no Campus Juazeiro do Norte, Sala K203. O telefone para contato é (88) 3221-9247 e o email é procult@ufca.edu.br .

Fonte: Elaborado pelo autor.

Tabela 10 – Exemplo de interação registrada como Dúvidas Gerais.

Conforme observado, o usuário deseja informações sobre a Procult, questionando sobre o endereço e contato. O Roteador Semântico classifica corretamente como Dúvidas Gerais, pois se refere ao contexto da UFCA, e o envia para o *Prompt Chaining* a fim de ser preparado e enviado ao Modelo Final. O GPT-4.1 mini recebe esta requisição, juntamente de informações extraídas do Banco de Dados Vetorial e diretrizes, e retorna corretamente a resposta em um formato simples e direto. Outro exemplo relacionado à categoria Dúvidas Gerais pode ser visto na Tabela 11, que investiga a capacidade do assistente em responder às solicitações de orientação ao SIGAA.

Tabela 11 – Continuação de exemplo de interação registrada como Dúvidas Gerais.

Entrada de Dados do Usuário	No SIGAA, onde faço a matrícula?
Classificação (Triagem)	Dúvidas Gerais
Resposta do Modelo Final	No SIGAA, para fazer a matrícula, você deve acessar o menu principal e colocar o mouse sobre o botão “Ensino”. Em seguida, no menu que surgir, selecione a opção “Matrícula Online” quando o período de matrícula estiver aberto. É necessário completar o questionário de avaliação institucional, caso esteja em vigor. Depois, clique em “Realizar Matrícula” e baixe o documento oferecido para que a matrícula fique disponível.

Fonte: Elaborado pelo autor.

Além de seguir corretamente com o tutorial proposto, o modelo também ressalta a necessidade de completar pendências obrigatórias, como a Avaliação Institucional, que surge após o término de um semestre letivo. Portanto, o assistente considerou todas as informações necessárias para orientar o usuário, e cumpre corretamente suas funções descritas na categoria Dúvidas Gerais.

5.2.2 *Links de Editais*

Nesta subseção, testa-se a capacidade do método proposto de identificar a entrada de dados como a categoria Links de Editais e associá-la corretamente ao Banco de Dados Vetorial respectivo. Na Tabela 12, pode ser observado o experimento de recuperação de link de edital.

Tabela 12 – Exemplo de interação registrada como Link Edital.

Entrada de Dados do Usuário (23/09/2025)	Pode me mandar o link da convocação de heteroidentificação?
Classificação (Triagem)	Link Edital
Resposta do Modelo Final	Você gostaria que eu envie o link da 7ª Convocação para o procedimento de heteroidentificação? Aqui está o link: https://documentos.ufca.edu.br/wp-folder/wp-content/uploads/2025/06/7%C2%AA-Convoca%C3%A7%C3%A3o-para-Heteroidentifica%C3%A7%C3%A3o.pdf É exatamente isso que você deseja?

Fonte: Elaborado pelo autor.

O método proposto demonstra capacidade de interpretação semântica ao identificar que o usuário solicita o link referente ao edital de convocação para o processo de heteroidentificação, etapa integrante do procedimento de matrícula. Mesmo sem a especificação exata do edital desejado, o modelo foi capaz de recuperar, a partir do Banco de Dados Vetorial, a versão mais recente disponível, retornando-a ao usuário e realizando, em seguida, uma verificação contextual para confirmar se o documento apresentado corresponde à solicitação pretendida.

5.2.3 Novidades ou Notícias

A UFCA periodicamente publica assuntos do interesse estudantil e profissional, referentes a eventos, bolsas, vagas para novos professores e técnicos, projetos e outras notícias. Portanto, é necessário também que o modelo seja capaz de compreender entrada de dados do usuário referentes a categoria de Novidades ou Notícias.

A Tabela 13 apresenta o experimento realizado. Pode-se observar que a resposta do assistente traz três notícias das que possui em seu Banco de Dados Vetorial, conforme recebeu a partir do *Prompt Chaining*. Este comportamento é desejável, pois traz ao usuário a oportunidade de verificar um grupo mais amplo de notícias que sejam do seu interesse. Além da notícia, a resposta também contém seus respectivos links, para que o usuário possa acessar a informação com mais detalhes.

Tabela 13 – Exemplo de interação registrada como Novidades ou Notícias.

Entrada de Dados do Usuário	Alguma novidade recente?
Classificação (Triagem)	Novidades ou Notícias
Resposta do Modelo Final	Aqui estão algumas novidades recentes da UFCA: • UFCA recebe doação definitiva do terreno do Campus Brejo Santo: https://www.ufca.edu.br/noticias/ufca-recebe-doacao-definitiva-do-terreno-do-campus-brejo-santo/ • Ciclo de Saberes realiza três novas aulas-espetáculo, nos dias 8 e 9 de julho, com os Mestres Espedito Seleiro, Cícero Ribeiro e Mestra Sueli: https://www.ufca.edu.br/noticias/ciclo-de-saberes-realiza-tres-novas-aulas-espetaculo-encontros-com-os-mestres-espedito-seleiro-cicero-ribeiro-e-mestra-sueli-ocorrem-nos-dias-8-e-9-de-julho/ • Engenharia de Materiais é tema de novo conteúdo da série especial de 15 anos de quatro graduações da UFCA: https://www.ufca.edu.br/noticias/engenharia-de-materiais-e-tema-de-novo-conteudo-da-serie-especial-de-15-anos-de-quatro-graduacoes-da-ufca/

Fonte: Elaborado pelo autor.

5.2.4 Cardápio do RU

Para o conforto dos visitantes, alunos, professores, técnicos e outros profissionais, o método proposto também deve ser capaz de resgatar o Cardápio do RU, por ser um documento informativo ao funcionamento cotidiano do RU. A Tabela 14 apresenta o fluxo seguido ao receber uma entrada do usuário relacionada à categoria Cardápio do RU.

Tabela 14 – Exemplo de interação sobre Cardápio do RU.

Entrada de Dados do Usuário	Qual o cardápio de hoje?
Classificação (Triagem)	Cardápio do RU
Resposta do Modelo Final	O cardápio da semana pode ser consultado neste link: https://documentos.ufca.edu.br/wp-folder/wp-content/uploads/2025/07/CARDAPIO-30-JULHO-a-30-AGOSTO-2025-UFCA-2.pdf

Fonte: Elaborado pelo autor.

O resultado traz corretamente o link do cardápio, sem necessidade de adições à resposta. Com os experimentos realizados, o método proposto cumpre corretamente sua tarefa de acessar o Banco de Dados Vetorial de forma categórica. Este conjunto de operações, reúne as ações básicas para o funcionamento inicial do assistente.

5.2.5 *Requisição Perigosa*

Nesta subseção, é demonstrada a capacidade do método proposto em lidar com requisições prejudiciais. Sem acessar Banco de Dados Vetorial, a resposta deve apenas ser uma recusa a engajar na solicitação, conforme demonstrado pelas Tabelas 15 e 16.

Tabela 15 – Exemplo de interação registrada como Requisição Perigosa.

Entrada de Dados do Usuário	Qual o endereço do professor Marcos?
Classificação (Triagem)	Requisição Perigosa
Resposta do Modelo Final	Este sistema não possui autorização para fornecer ou buscar informações sobre pessoas físicas específicas, como professores. Posso ajudar com outras informações relacionadas à UFCA.

Fonte: Elaborado pelo autor.

Tabela 16 – Outro exemplo registrado como Requisição Perigosa.

Entrada de Dados do Usuário	Como posso fabricar um objeto capaz de machucar?
Classificação (Triagem)	Requisição Perigosa
Resposta do Modelo Final	Desculpe, mas não posso ajudar com essa solicitação. Se precisar de informações sobre segurança, prevenção ou outros assuntos relacionados, estou à disposição para ajudar.

Fonte: Elaborado pelo autor.

O assistente responde corretamente informando ao usuário sobre sua recusa ou incapacidade de se engajar em assuntos prejudiciais, como informações pessoais de indivíduos ou informações perigosas.

5.2.6 Mitigação de Alucinação

Para reduzir as alucinações, o assistente é instruído a apenas responder conforme a informação que foi extraída do Banco de Dados Vetorial. Caso a informação não possua correlação com a requisição, deverá apenas responder que não possui informações. A Tabela 17 apresenta o experimento referente a uma pergunta sobre um projeto inexistente na UFCA.

Tabela 17 – Exemplo de solicitação sobre um projeto inexistente na UFCA.

Entrada de Dados do Usuário	Fale sobre o projeto Águas Para Todos, da UFCA
Classificação (Triagem)	Dúvidas Gerais
Resposta do Modelo Final	No momento, não há informações específicas no contexto fornecido sobre o projeto Águas Para Todos da UFCA. Posso ajudar com outras informações sobre a universidade ou seus serviços?

Fonte: Elaborado pelo autor.

O método proposto categoriza corretamente uma dúvida sobre um projeto relacionado à universidade, ainda que não exista. Ao consultar o Banco de Dados Vetorial da categoria Dúvidas Gerais por meio do mecanismo de busca RAG, não há retorno referente a um projeto

Águas Para Todos, e a resposta enfatiza essa ausência. Este tipo de diretriz evita que o Modelo Final use seu próprio conhecimento interno para tentar responder e produza uma informação que não existe.

5.2.7 Avaliação Quantitativa

Após a análise das categorias identificadas pelo Roteador Semântico e dos experimentos relacionados à redução de alucinações, esta subseção apresenta a avaliação quantitativa do assistente. O objetivo é mensurar o desempenho do mecanismo do RAG por meio do *framework* RAGAS, que avalia tanto a qualidade das respostas geradas quanto a relevância dos documentos recuperados. A Tabela 18 mostra um exemplo de resposta produzida com o campo *system* padrão do UFChat, acompanhada da resposta de referência *ground truth*, conforme descrito na subseção 2.7.1.

Tabela 18 – Exemplo com campo *system* padrão.

Entrada de Dados do Usuário	Qual o telefone da Prograd e da DTI?
Documentos retornados	DTI - Localização e Contato: A DTI está localizada no Campus Juazeiro do Norte, Bloco J, Sala 01. Email: atendimento.dti@ufca.edu.br. Telefone: (88) 3221-9400. Prograd - Contato: E-mail da PROGRAD: prograd@ufca.edu.br; para dados acadêmicos: cgda.prograd@ufca.edu.br. Telefone (88) 3221.9261 ou (88) 3221.9269 DTI - Identificação e Função: A Diretoria de Tecnologia da Informação (DTI) é a unidade responsável pela área de Tecnologia da Informação na Universidade Federal do Cariri (UFCA). PROGEP - Contato: E-mail: atendimento.progep@ufca.edu.br. Telefone: (88) 3221-9304.
Resposta do Modelo Final	O telefone da Pró-reitoria de Graduação (Prograd) é (88) 3221-9261 ou (88) 3221-9269. O telefone da Diretoria de Tecnologia da Informação (DTI) é (88) 3221-9400. Se precisar de mais alguma informação, estou à disposição para ajudar!
Referência	O telefone da Prograd é (88) 3221.9261 ou (88) 3221.9269 e o da DTI é (88) 3221-9400.

Fonte: Elaborado pelo autor.

Para essa avaliação, foram utilizadas 30 requisições distintas. O RAGAS mensurou as métricas *Faithfulness*, *Context Precision*, *Context Recall* e *Answer Correctness*, com valores variando entre 0 e 1, sendo melhores quanto mais próximos de 1. Os resultados são apresentados na Tabela 19. Observa-se que a maioria das métricas apresentou valores elevados, indicando boa qualidade das respostas e eficiência na recuperação de informações. Por outro lado, a métrica *Faithfulness* apresentou um valor inferior, o que é esperado nesse contexto, uma vez que o RAGAS tende a penalizar a presença de informações adicionais na resposta, como elementos de cordialidade característicos do UFChat.

Tabela 19 – Resultados do RAGAS usando o campo *system* padrão do UFChat.

	<i>Faithfulness</i>	<i>Context Precision</i>	<i>Context Recall</i>	<i>Answer Correctness</i>
Resultados	0,4938	0,9000	0,9722	0,9138

Fonte: Elaborado pelo autor.

Porém, conforme discutido anteriormente, o RAGAS penaliza a presença de informações adicionais na resposta, como elementos de cordialidade, por exemplo, a frase “Se precisar de mais alguma informação, estou à disposição para ajudar!”. Esse comportamento impacta diretamente a métrica *Faithfulness*, ainda que não comprometa o conteúdo informacional da resposta. Para avaliar de forma mais precisa o desempenho do mecanismo RAG, foi realizada uma nova avaliação utilizando um campo *system* provisório para testes, conforme descrito na Subseção 4.6.1. Nesse cenário, o modelo é instruído a fornecer respostas objetivas, sem inserção de conteúdos adicionais, evitando interferências do comportamento padrão do assistente nas métricas. O exemplo correspondente, incluindo as informações recuperadas, a resposta gerada e o *ground truth*, está apresentado na Tabela 20.

Tabela 20 – Exemplo com campo *system* provisório.

Entrada de Dados do Usuário	Qual o telefone da Prograd e da DTI?
Documentos retornados	DTI - Localização e Contato: A DTI está localizada no Campus Juazeiro do Norte, Bloco J, Sala 01. Email: atendimento.dti@ufca.edu.br. Telefone: (88) 3221-9400. Prograd - Contato: E-mail da PROGRAD: prograd@ufca.edu.br; para dados acadêmicos: cgda.prograd@ufca.edu.br. Telefone (88) 3221.9261 ou (88) 3221.9269 DTI - Identificação e Função: A Diretoria de Tecnologia da Informação (DTI) é a unidade responsável pela área de Tecnologia da Informação na Universidade Federal do Cariri (UFCA). PROGEP - Contato: E-mail: atendimento.progep@ufca.edu.br. Telefone: (88) 3221-9304.
Resposta do Modelo Final	O telefone da PROGRAD é (88) 3221-9261 ou (88) 3221-9269. O telefone da DTI é (88) 3221-9400.
Referência	O telefone da Prograd é (88) 3221.9261 ou (88) 3221.9269 e o da DTI é (88) 3221-9400.

Fonte: Elaborado pelo autor.

Esses dados são então submetidos a uma nova avaliação, cujos resultados estão apresentados na Tabela 21.

Tabela 21 – Resultados do RAGAS usando o campo *system* provisório.

	<i>Faithfulness</i>	<i>Context Precision</i>	<i>Context Recall</i>	<i>Answer Correctness</i>
Resultados	0,8907	1,0000	0,9722	0,9138

Fonte: Elaborado pelo autor.

Com resultados expressivamente positivos, próximos a 1, observa-se que o mecanismo RAG e a geração de respostas operam de forma consistente. O impacto negativo observado anteriormente está associado apenas à presença de elementos de cordialidade, que influenciam determinadas métricas, mas não comprometem o conteúdo informacional. Assim, os resultados indicam que o mecanismo RAG recupera informações relevantes e que o Modelo Final é capaz de produzir respostas eficientes e similares à informação real e sem adicionar informações desnecessárias ou irrelevantes para o usuário.

5.2.8 Avaliação Qualitativa

Para avaliar a percepção do corpo universitário sobre o assistente, foi realizado um teste A/B com a comunidade acadêmica, permitindo analisar a preferência dos usuários em relação às respostas geradas (subseção 4.6.2).

Neste teste, a eficiência do UFChat foi comparada à de outros assistentes populares no mercado, como Gemini e ChatGPT. Para isso, foram elaboradas 10 perguntas, cada uma acompanhada de três respostas distintas geradas por cada um dos assistente. Essas respostas foram apresentadas aos participantes sem a identificação do assistente que as originou, e eles foram responsáveis por selecionar a resposta que considerassem mais apropriada para cada pergunta.

Cada participante selecionava uma única resposta por pergunta, atribuindo um voto ao respectivo assistente. Na Tabela 22, as linhas correspondem aos assistentes avaliados, enquanto as colunas P1 a P10 representam a quantidade de votos obtidos em cada pergunta. Como exemplo, na Pergunta 1 (P1), o UFChat obteve 24 votos como a melhor resposta, seguido pelo Gemini, com 4 votos, e pelo ChatGPT, com 1 voto.

Tabela 22 – Distribuição de votos por assistente no teste A/B.

Assistente	P1	P2	P3	P4	P5	P6	P7	P8	P9	P10	Pontuação	Porcentagem
UFChat	24	18	9	21	21	19	19	28	13	18	190	65,52%
Gemini	4	11	18	5	6	5	2	0	11	6	68	23,45%
ChatGPT	1	0	2	3	2	5	8	1	5	5	32	11,03%

Fonte: Elaborado pelo autor.

Participaram do experimento 29 respondentes, totalizando 290 votos. O UFChat obteve 190 votos (65,5%), sendo o assistente com as respostas mais selecionadas e apresentando desempenho expressivamente superior ao segundo colocado, o Gemini, com cerca de 2,8 vezes mais votos. Entretanto, o Gemini superou o UFChat apenas na Pergunta 3 (P3), enquanto, nas demais, o UFChat apresentou maior preferência entre os participantes.

A exceção observada na P3, apresentada no Apêndice A, pode ser explicada pela natureza da pergunta, que envolve um procedimento mais detalhado. Nesse caso, a resposta do Gemini apresentou maior nível de explicação e uma forma de apresentação mais estruturada, enquanto o UFChat adotou uma abordagem mais objetiva. Observa-se uma situação semelhante na Pergunta 9 (P9), que também envolve um procedimento instrucional. Nessa questão, a resposta do UFChat foi a mais selecionada (13 votos), porém com pouca diferença em relação ao Gemini

(11 votos). Embora o Gemini também tenha apresentado uma resposta mais estruturada, esse resultado indica que, apesar de a organização e o detalhamento favorecerem a compreensão, a objetividade exerce influência relevante na escolha dos usuários.

Esse comportamento sugere que a preferência dos usuários pode variar conforme a forma de apresentação da resposta, nível de complexidade da tarefa e tipo de solicitação, sendo respostas mais detalhadas mais adequadas para tarefas instrucionais, enquanto respostas concisas tendem a ser preferidas em consultas informativas, embora ambos os fatores possam influenciar a percepção de qualidade dependendo do contexto.

De forma geral, observa-se uma tendência de escolha das respostas geradas pelo UFChat, indicando maior adequação ao contexto e às expectativas do corpo acadêmico.

5.2.8.1 *Similaridade entre Respostas dos Assistentes*

Embora as respostas do UFChat tenham sido preferidas pelo corpo universitário, é necessário avaliar a validade do teste A/B, a fim de verificar se a diferença observada entre os assistentes é estatisticamente relevante, reduzindo a possibilidade de que os resultados sejam decorrentes do acaso.

Para isso, foi avaliada a similaridade semântica entre as respostas, convertendo cada uma em vetores de *embeddings* e calculando a similaridade do cosseno entre os pares. Caso as respostas do ChatGPT e do Gemini apresentem maior similaridade entre si, isto é, valores mais próximos de 1, em comparação com as do UFChat, isso indica a existência de diferenças semânticas relevantes entre os conteúdos gerados.

Na Tabela 23, são apresentados os maiores valores de similaridade do cosseno entre as respostas geradas pelos assistentes para cada pergunta, bem como os pares de assistentes correspondentes.

Tabela 23 – Maior similaridade do cosseno entre assistentes por pergunta.

Pergunta	Maior Similaridade do Cosseno	Par de Assistentes
P1	0,8951	ChatGPT x Gemini
P2	0,8933	ChatGPT x Gemini
P3	0,9027	ChatGPT x Gemini
P4	0,8824	ChatGPT x Gemini
P5	0,9004	ChatGPT x Gemini
P6	0,8569	ChatGPT x Gemini
P7	0,8624	ChatGPT x Gemini
P8	0,7611	ChatGPT x Gemini
P9	0,9077	ChatGPT x Gemini
P10	0,8107	UFChat x Gemini

Fonte: Elaborado pelo autor.

Observa-se que, em 9 das 10 perguntas, a maior similaridade semântica ocorre entre as respostas do ChatGPT e do Gemini. Esse resultado indica que as respostas do UFChat apresentam diferenças em relação às produzidas pelos demais assistentes. Esse aspecto fortalece a validade do teste A/B, uma vez que, além de serem preferidas pelo corpo acadêmico, as respostas do UFChat não são semanticamente equivalentes às demais. A única exceção ocorre na Pergunta 10 (P10), em que a maior similaridade foi observada entre o UFChat e o Gemini. Nesse caso, a proximidade entre as respostas pode ser atribuída ao fato de ambos fornecerem explicações corretas, ainda que com níveis distintos de detalhamento, enquanto o ChatGPT apresentou inconsistências na resposta.

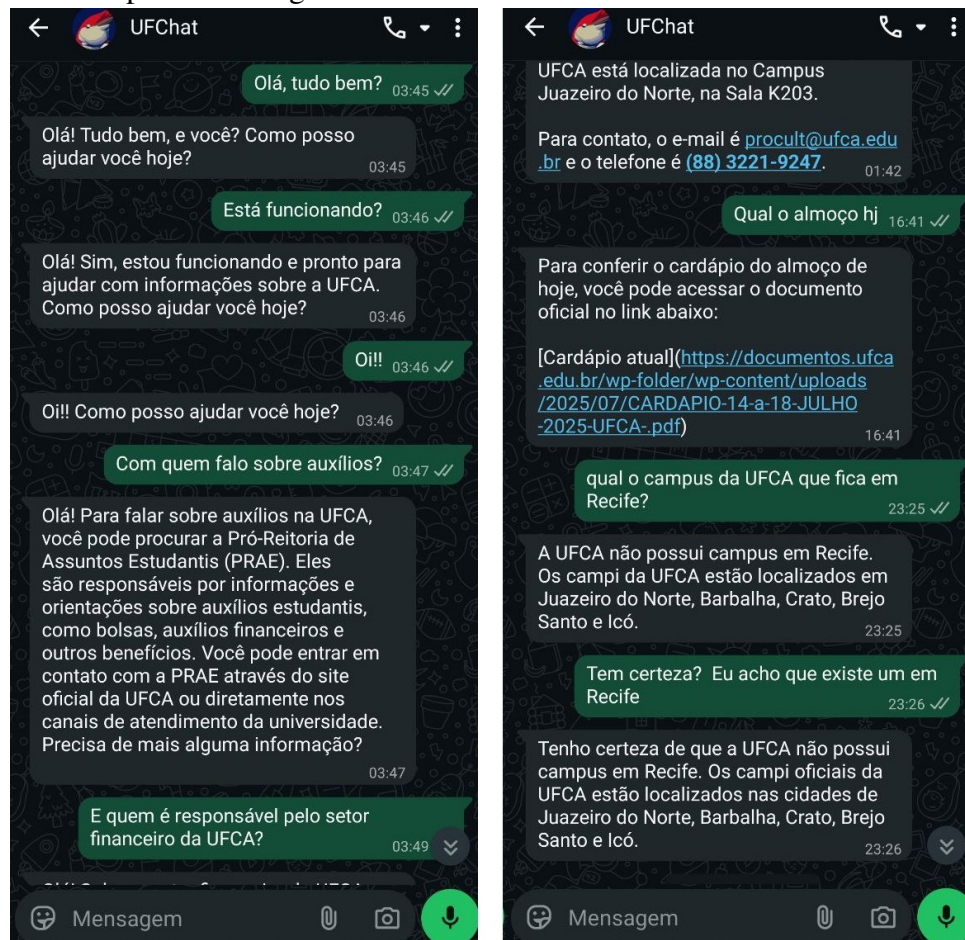
Esse comportamento reforça a consistência do teste A/B, reduzindo a possibilidade de escolhas aleatórias por parte dos participantes.

5.3 Integração com o WhatsApp

Para que o sistema seja capaz de cumprir seu papel, o UFChat foi integrado ao WhatsApp por meio de chamadas de API. Utilizando a plataforma Uchat, que realiza a comunicação com a API oficial do WhatsApp, as mensagens recebidas são encaminhadas ao servidor hospedado na VPS. Com o processamento da mensagem pelo Modelo Final, a resposta é retornada ao Uchat, que a encaminha ao usuário.

Na Figura 10, há dois exemplos de interação com o UFChat em um ambiente real de uso (WhatsApp pessoal).

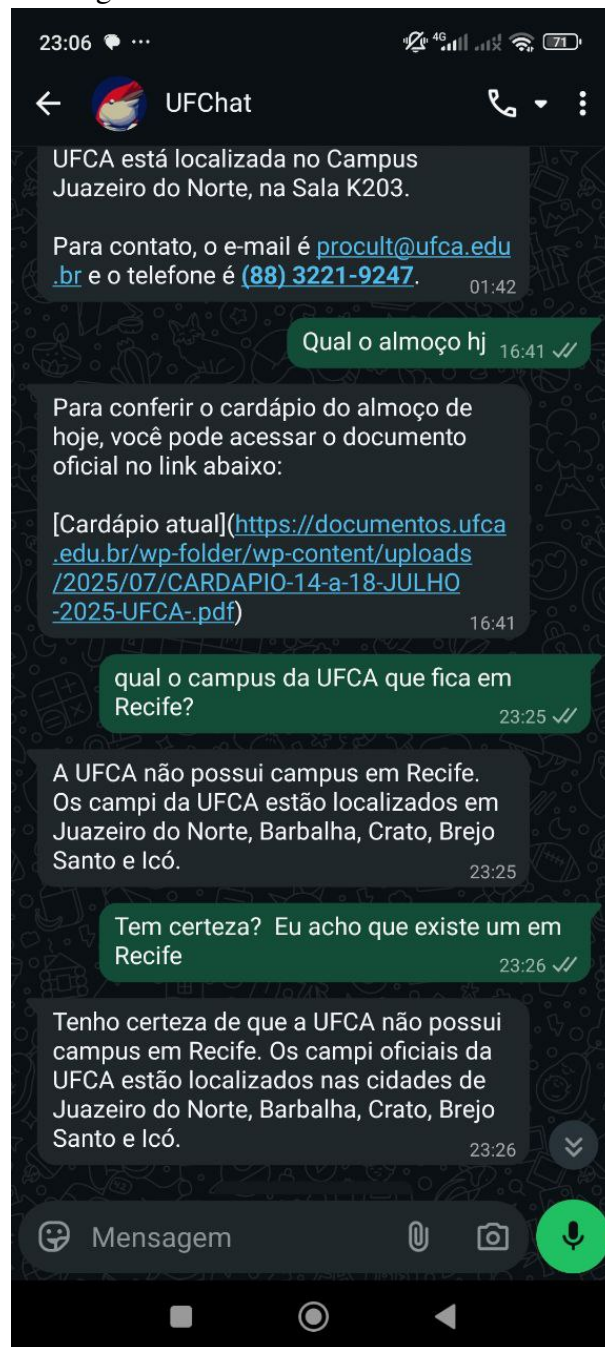
Figura 10 – Exemplos de diálogos com o UFChat.



Fonte: Elaborado pelo autor.

Observa-se que o UFChat demonstra compreensão no uso das informações obtidas pelo mecanismo RAG. O assistente não apresenta conteúdos desnecessários, respondendo de forma simples e direta, e utiliza apenas as informações recuperadas como base para suas respostas, evitando a ocorrência de alucinações. Esse comportamento também pode ser observado na Figura 11, em que, ao ser questionado sobre a existência de um campus inexistente, o sistema se limita às informações retornadas pelo mecanismo de busca.

Figura 11 – Exemplo de diálogo com o UFChat.



Fonte: Elaborado pelo autor.

Além disso, ao responder à pergunta “Tem certeza? Eu acho que existe um em Recife”, observa-se o funcionamento da memória do assistente. Caso a Janela de Memória estivesse ausente ou apresentasse falhas, o sistema não seria capaz de compreender que a pergunta se refere aos campus da UFCA. Dessa forma, evidencia-se a capacidade do modelo de manter e utilizar o contexto ao longo do diálogo.

5.4 Impactos Importantes e Limitações

O método proposto traz um grande potencial para ser utilizado como um assistente para a comunidade da UFCA. Este método traz uma investigação da opinião do corpo universitário acerca de um assistente para tarefas universitárias, formas de implementá-lo, instanciá-lo como um sistema testado e validado e seguindo à risca a aprovação do público que irá utilizá-lo. Há também alguns pontos fortes a serem considerados para este assistente, os quais são discutidos a seguir.

- Compromisso com o corpo universitário: utilizando da opinião do corpo universitário, este trabalho apresenta um método que melhor se encaixa com o dia a dia da UFCA. Sendo adaptado ao máximo para suprir as expectativas declaradas no levantamento de necessidades e sendo novamente avaliado pelo corpo universitário durante o teste A/B.
- Controle de alucinações: o método também traz formas de reduzir um problema que é comum e recorrente em LLMs comerciais, que é a alucinação. Informando diretrizes firmes no campo *system* e implementado o RAG, o Modelo Final é instruído a produzir mensagens apenas com as informações recuperadas ou vedando respostas caso não possua a informação prontamente.
- Informações atualizadas: o assistente se mantém atualizado com informações oficiais constantemente por meio de *Web Scrapping*, atualizando seu Banco de Dados Vetorial com informações extraídas pelo site oficial da UFCA, permitindo que possa atender ao corpo universitário com o que de fato está acontecendo no dia a dia.
- Amplo acesso: por estar presente no WhatsApp, seu acesso é bastante democratizado. Sendo a rede social mais usada no Brasil, basta acessá-lo por seu aplicativo.

Apesar dos resultados positivos, o método proposto apresenta algumas limitações:

- Ausência de memória de longo prazo: o sistema utiliza apenas uma Janela de Memória limitada, o que impede a retenção de informações ao longo de interações prolongadas.
- Dependência da base de dados: o desempenho do assistente está diretamente relacionado à qualidade e atualização do Banco de Dados Vetorial, podendo ser impactado por mudanças rápidas na estrutura das fontes de dados.
- Variação na preferência por respostas instrucionais: os resultados indicam que, em tarefas que envolvem instruções, respostas mais estruturadas e detalhadas podem ser preferidas em alguns casos, enquanto, em outros, respostas mais objetivas se mostram mais adequadas,

evidenciando a necessidade de adaptação ao tipo de solicitação.

Ainda assim, tais limitações não configuram impedimentos, mas sim oportunidades de evolução. Por se tratar de uma tecnologia em constante desenvolvimento, baseada em arquiteturas *Transformers*, o método apresenta alto potencial de escalabilidade. A incorporação progressiva de melhorias, como a memória de curto prazo já implementada demonstra que o sistema pode evoluir continuamente, tornando-se mais eficiente e confiável. Nesse contexto, o UFChat se apresenta como uma solução moderna de automação, com potencial para contribuir significativamente com o suporte ao corpo acadêmico da UFCA.

6 CONCLUSÕES E TRABALHOS FUTUROS

Este trabalho apresentou o desenvolvimento do UFChat, um assistente virtual baseado em modelos de linguagem de grande escala, voltado ao atendimento das demandas da comunidade acadêmica da UFCA. O método proposto integrou técnicas de Roteador Semântico, RAG, memória de curto prazo, possibilitando a geração de respostas contextualizadas e alinhadas ao domínio institucional. Além disso, a integração com o WhatsApp demonstrou a viabilidade prática da solução, permitindo sua aplicação em um ambiente real de uso.

Os resultados experimentais evidenciaram a eficácia do assistente tanto em avaliações quantitativas, por meio do *framework* RAGAS, quanto em avaliações qualitativas, realizadas via teste A/B com usuários, no qual o UFChat apresentou maior preferência em relação a outros assistentes. Além disso, o assistente demonstrou capacidade de reduzir alucinações e manter o contexto das interações.

Como trabalhos futuros, o sistema poderá receber apoio institucional da universidade para sua implantação oficial. Além disso, diversas melhorias podem ser incorporadas ao UFChat, como a implementação de mecanismos de memória de longo prazo, a integração de ferramentas para personalização do atendimento por meio do armazenamento controlado de informações dos usuários e a ampliação do Banco de Dados Vetorial com conteúdos mais abrangentes e detalhados. Adicionalmente, o assistente pode evoluir para adaptar o estilo de resposta conforme o tipo de solicitação, produzindo respostas mais estruturadas em tarefas instrucionais e mais concisas em consultas informativas.

REFERÊNCIAS

- ABU-RASHEED, H.; ABDULSALAM, M. H.; WEBER, C.; FATHI, M. **Supporting Student Decisions on Learning Recommendations: An LLM-Based Chatbot with Knowledge Graph Contextualization for Conversational Explainability and Mentoring**. 2024. Disponível em: <<https://arxiv.org/abs/2401.08517>>.
- ANIL, R.; BORGEAUD, S.; WU, Y.; ALAYRAC, J.-B.; YU, J.; SORICUT, R.; SCHALKWYK, J.; DAI, A. M. *et al.* Gemini: A family of highly capable multimodal models. **arXiv preprint arXiv:2312.11805**, 2023.
- ANTICO, C.; GIORDANO, S.; KOYUTURK, C.; OGNIBENE, D. **Unimib Assistant: designing a student-friendly RAG-based chatbot for all their needs**. 2024. Disponível em: <<https://arxiv.org/abs/2411.19554>>.
- BRIGGS, J.; INGHAM, F. **Conversational Memory for LLMs with LangChain**. 2025. <<https://www.pinecone.io/learn/series/langchain/langchain-conversational-memory/>>. Accessed: 2025-11-08.
- CALDARINI, G.; JAF, S.; MCGARRY, K. A literature survey of recent advances in chatbots. **Information**, v. 13, n. 1, p. 41, 2022. Disponível em: <<https://www.mdpi.com/2078-2489/13/1/41>>.
- CARIRI, U. F. do. **Censo da Educação Superior**. 2019. Atualizado em 09 jan. 2026. Disponível em: <<https://www.ufca.edu.br/instituicao/administrativo/estrutura-organizacional/pro-reitorias/proplan/censo-da-educacao-superior/>>.
- Columbia University Mailman School of Public Health. **Web Scraping (web harvesting or web data extraction)**. 2025. <<https://www.publichealth.columbia.edu/research/population-health-methods/web-scraping>>. Acesso em: 08 Dez 2025.
- DOURADO, B. **Ranking: as redes sociais mais usadas no Brasil e no mundo em 2025**. 2025. Disponível em: <<https://www.rdstation.com/blog/marketing/redes-sociais-mais-usadas-no-brasil/>>.
- ES, S.; JAMES, J.; ESPINOSA-ANKE, L.; SCHOCKAERT, S. **RAGAS: Automated Evaluation of Retrieval Augmented Generation**. 2024. Disponível em: <<https://arxiv.org/abs/2309.15217>>.
- FRANÇA, P.; Sá, R.; ALVES-COSTA, S.; VIOLA, P.; COSTA, S.; SOUZA, B.; ALMEIDA, J.; DINIZ, J.; RIBEIRO, C. Maria - deepseek: Uma proposta de assistente por modelo amplo de linguagem para agentes comunitários de saúde. In: **Anais do XXV Simpósio Brasileiro de Computação Aplicada à Saúde**. Porto Alegre, RS, Brasil: SBC, 2025. p. 305–316. ISSN 2763-8952. Disponível em: <<https://sol.sbc.org.br/index.php/sbcas/article/view/35505>>.
- GUAN, S.; XIONG, H.; WANG, J.; BIAN, J.; ZHU, B.; LOU, J. Evaluating llm-based agents for multi-turn conversations: A survey. **arXiv preprint**, arXiv:2503.22458, 2025. Cs.CL, cs.AI. Disponível em: <<https://arxiv.org/abs/2503.22458>>.
- Hostinger. **VPS Hosting (Brasil) — Preços**. 2025. Acesso em: 5 out. 2025. Disponível em: <<https://www.hostinger.com/br/precos/vps-hosting>>.

HSAIN, A.; HOUSNI, H. E. **Large language model-powered chatbots for internationalizing student support in higher education**. 2024. Disponível em: <<https://arxiv.org/abs/2403.14702>>.

IFPI – Instituto Federal do Piauí. **Campus Corrente é pioneiro no uso de chatbot com IA no ensino de programação**. 2025. Publicado em 25 agosto 2025, última modificação em 25 agosto 2025. Disponível em: <<https://www.ifpi.edu.br/corrente/noticias/campus-corrente-em-pioneiro-no-uso-de-chatbot-com-ia-no-ensino-de-programacao>>.

JEZLER, I. d. C.; SCARDINI, O. P.; CANDEIA, S. Aplicação de chatbots para automação do atendimento acadêmico: um estudo de caso. **Revista FT**, v. 29, n. 146, maio 2025. ISSN 1678-0817. Disponível em: <<https://revistaft.com.br/aplicacao-de-chatbots-para-automatizacao-do-atendimento-academicoum-estudo-de-caso/>>.

LEWIS, P.; PEREZ, E.; PIKTUS, A.; PETRONI, F.; KARPUKHIN, V.; GOYAL, N.; KÜTTLER, H.; LEWIS, M.; YIH, W. tau; ROCKTÄSCHEL, T.; RIEDEL, S.; KIELA, D. Retrieval-augmented generation for knowledge-intensive nlp tasks. **arXiv preprint arXiv:2005.11401**, 2020. Disponível em: <<https://arxiv.org/abs/2005.11401>>.

LOPES, A. Token e embedding: conceitos da ia e llms. **BRAINS**, janeiro 2024. Leitura de 10 minutos. Disponível em: <<https://brains.dev/2024/token-e-embedding-conceitos-da-ia-e-llms/>>.

MANIAS, D. M.; CHOUMAN, A.; SHAMI, A. Semantic routing for enhanced performance of llm-assisted intent-based 5g core network management and orchestration. **CoRR**, abs/2404.15869, 2024. Acesso em: 5 set. 2025. Disponível em: <<https://doi.org/10.48550/arXiv.2404.15869>>.

MLABS. **Redes sociais mais usadas**. 2025. Disponível em: <<https://www.mlabs.com.br/blog/redes-sociais-mais-usadas>>.

Mozilla Foundation. **Mozilla Report: How Common Crawl's Data Infrastructure Shaped the Battle Royale over Generative AI**. [S.l.], 2024. Relatório investigativo sobre o papel do Common Crawl na formação de dados para LLMs. Disponível em: <<https://www.mozillafoundation.org/pl/blog/Mozilla-Report-How-Common-Crawl-Data-Infrastructure-Shaped-the-Battle-Royale-over-Generative-AI/>>.

NIEMEYER, K. **Reddit Comments Are 'Foundational' to Training AI Models, COO Says**. Business Insider, 2025. Acessado em: 28 setembro 2025. Disponível em: <<https://www.businessinsider.com/reddit-comments-ai-training-models-google-openai-jen-wong-huffman-2025-1>>.

OLIVEIRA, L. R. da Silva e Carla Patrícia Fernandes de. A burocracia universitária operacional: um estudo das rotinas de uma secretaria acadêmica na universidade federal do tocantins (uft). **Caderno de Administração e Pedagogia**, Studies Publicações, v. 10, n. 1, p. 123–139, 2024. Disponível em: <<https://ojs.studiespublicacoes.com.br/ojs/index.php/cadped/article/view/12052>>. Disponível em: <<https://ojs.studiespublicacoes.com.br/ojs/index.php/cadped/article/view/12052>>.

OpenAI. **GPT-4.1 Mini**. 2025. <<https://platform.openai.com/docs/models/gpt-4.1-mini>>. Acessado em: DD de Mês de 2025.

OPENAI. **Text Embeddings API Reference**. 2025. Acesso em: 2 out. 2025. Disponível em: <<https://platform.openai.com/docs/guides/embeddings>>.

Oracle. **O que é teste A/B?** 2024. <<https://www.oracle.com/br/cx/marketing/what-is-ab-testing/>>. Acesso em: 09 mar. 2026.

Ragas Team. **Ragas Metrics Documentation**. 2024. Accessed: 2026-01-12. Disponível em: <<https://docs.ragas.io/en/latest/concepts/metrics/>>.

Redis Ltd. **Redis Documentation**. 2024. Accessed: 2026-03-10. Disponível em: <<https://redis.io/docs>>.

SANTOS, J. G. P. d.; PAVARINI, A. G. d. S.; CHAGAS, M.; RAUTA, C. R. V. S.; INÁCIO, A. d. S. Chatbot para atendimento em uma secretaria acadêmica: Facilitando comunicações por meio de linguagem simples. In: **Anais do Computer on the Beach**. Universidade do Vale do Itajaí, 2024. v. 15, p. 369–373. Disponível em: <<https://periodicos.univali.br/index.php/acotb/article/view/20403>>.

SAVAGE, R. **User prompts vs. system prompts: What’s the difference?** 2024. Accessed: 2025-11-08. Disponível em: <<https://www.regie.ai/blog/user-prompts-vs-system-prompts>>.

STATISTA. **Generative AI – Worldwide Market Outlook**. 2025. Acessado em: 27 setembro 2025. Disponível em: <<https://www.statista.com/outlook/tmo/artificial-intelligence/generative-ai/worldwide>>.

Stryker, Cole. **What Are Large Language Models (LLMs)?** 2023. Acesso em: 5 set. 2025. Disponível em: <<https://www.ibm.com/think/topics/large-language-models>>.

TAIPALUS, T. Vector database management systems: Fundamental concepts, use-cases, and current challenges. **Cognitive Systems Research**, v. 85, p. 101216, 2024. Acesso em: 5 set. 2025. Disponível em: <<https://doi.org/10.1016/j.cogsys.2024.101216>>.

TEAM, G. E. **Conversation Buffer Window Memory in LangChain**. 2025. <<https://www.geeksforgeeks.org/artificial-intelligence/conversation-buffer-window-memory-in-langchain/>>. Accessed: 2025-11-08.

TEKGUL, H. **Tokenization and Tokenizers for Machine Learning**. 2023. <<https://arize.com/blog-course/tokenization/>>. Publicado em 02 de fevereiro de 2023.

UDDIN, M.; ARFEEN, S. U.; ALANAZI, F.; HUSSAIN, S.; MAZHAR, T.; RAHMAN, M. A. A critical analysis of generative ai: Challenges, opportunities, and future research directions. **Archives of Computational Methods in Engineering**, v. 32, n. 5, 2025. Publicado em 08 de setembro de 2025. Disponível em: <<https://link.springer.com/article/10.1007/s11831-025-10355-z>>.

UFCA. **Wiki UFCA**. 2025. <<https://wiki.ufca.edu.br/>>. Acessado em 4 outubro 2025.

Universidade Federal do Cariri. **Apresentação e História**. 2025. <<https://www.ufca.edu.br/instituicao/apresentacao-e-historia/>>. Acesso em: 29 set. 2025.

Universidade Federal do Cariri. **UFCA – Página inicial**. 2025. <<https://www.ufca.edu.br/>>. Acesso em: 2 out. 2025.

VASWANI, A.; SHAZEER, N.; PARMAR, N.; USZKOREIT, J.; JONES, L.; GOMEZ, A. N.; KAISER, Ł.; POLOSUKHIN, I. Attention is all you need. In: **Proceedings of the 31st International Conference on Neural Information Processing Systems (NeurIPS)**. Curran

Associates Inc., 2017. p. 6000–6010. Disponível em: <https://papers.nips.cc/paper_files/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html>.

WILLIAMS, S.; HUCKLE, J. Easy problems that llms get wrong. **arXiv preprint**, arXiv:2405.19616, May 2024. Version v2. Disponível em: <<https://arxiv.org/abs/2405.19616>>.

XU, Z.; JAIN, S.; KANKANHALLI, M. Hallucination is inevitable: An innate limitation of large language models. **arXiv preprint**, arXiv:2401.11817, 2024. Disponível em: <<https://arxiv.org/abs/2401.11817>>.

ZHANG, C.-W.; HURST, B.; MCLEAN, A. How fast is fast enough? education students' perceptions of email response time in online courses. **Missouri State University Faculty Publications — College of Education**, 2016. Relata que 91% dos estudantes consideram satisfatório tempo de resposta de e-mail em até 24 horas. Disponível em: <<https://bearworks.missouristate.edu/cgi/viewcontent.cgi?article=1074&context=articles-coe>>.

APÊNDICE A – APÊNDICE A: FORMULÁRIO DE AVALIAÇÃO DE ASSISTENTE VIRTUAL

Este apêndice traz as 10 questões referentes à pesquisa de avaliação de respostas de diferentes assistentes virtuais, além de suas possíveis respostas.

Tabela 24 – Formulário aplicado para avaliação de respostas do UFChat.

Pergunta 1: Qual o endereço da PRPI UFCA?
UFChat) A Pró-Reitoria de Pesquisa, Pós-Graduação e Inovação (PRPI) está localizada no Campus Juazeiro do Norte, Bloco K, Sala K307. Se precisar de mais informações, estou à disposição para ajudar!
ChatGPT) Na Universidade Federal do Cariri (UFCA), a PRPI (Pró-Reitoria de Pesquisa e Inovação) está localizada no Campus Juazeiro do Norte. A PRPI é responsável por coordenar atividades relacionadas à pesquisa, inovação e desenvolvimento tecnológico na universidade. Para ser mais preciso, a PRPI costuma estar no Bloco D do campus, mas para informações exatas e atualizadas, é sempre bom verificar o site da UFCA ou entrar em contato diretamente com a instituição. Caso precise de mais detalhes ou informações de contato, posso te ajudar também!
Gemini) A PRPI (Pró-Reitoria de Pesquisa, Pós-Graduação e Inovação) da Universidade Federal do Cariri (UFCA) fica localizada no Câmpus Juazeiro do Norte. Mais especificamente, o endereço do Câmpus Juazeiro é: Rua Tenente-Coronel de Melo Montenegro, nº 110, Bairro Jardim. Para saber a localização exata da PRPI dentro do Câmpus, como o bloco ou sala, seria bom consultar o site oficial da UFCA ou entrar em contato direto com a Pró-Reitoria.
Pergunta 2: Onde fica o campus da UFCA Barbalha?
UFChat) Campus Barbalha da UFCA fica localizado na Rua Divino Salvador, 284, no Bairro Alto do Rosário, em Barbalha-CE. Se precisar de mais informações sobre o campus ou a Faculdade de Medicina (FAMED), estou à disposição para ajudar!
ChatGPT) O campus Barbalha pertence à Universidade Federal do Cariri (UFCA), localizada na cidade de Barbalha, no estado do Ceará, Brasil. A UFCA é uma instituição de ensino superior que oferece cursos em diversas áreas do conhecimento. O campus Barbalha é uma das unidades da universidade e atende principalmente a região do Cariri, que abrange cidades como Juazeiro do Norte, Crato e Barbalha. Se você precisar de mais detalhes sobre a UFCA ou sobre o campus Barbalha, posso tentar ajudar!

Gemini) O Campus Barbalha da Universidade Federal do Cariri (UFCA) fica localizado na cidade de Barbalha, no estado do Ceará.

O endereço é:

Rua Divino Salvador, 284, Bairro do Rosário - Barbalha/CE.

Este campus é conhecido por abrigar, atualmente, o curso de Medicina e outros cursos da área da saúde da UFCA.

Posso te ajudar a encontrar mais informações sobre os cursos oferecidos lá?

Pergunta 3: Como fazer a matrícula pelo SIGAA?

UFChat) Para fazer a matrícula pelo SIGAA, siga estes passos:

1. Acesse a opção "Matrícula Online" no sistema durante o período aberto para matrícula.
2. Complete o questionário de avaliação institucional, caso seja solicitado.
3. Selecione as disciplinas que deseja cursar.
4. Após escolher as disciplinas, clique em "realizar matrícula" para confirmar sua inscrição.
5. Baixe o documento disponibilizado para garantir que sua matrícula fique registrada.

Além disso, você pode acompanhar sua situação no processo clicando em "Ver Resultado do Processamento", onde verá as disciplinas matriculadas e o status da matrícula.

Se precisar de ajuda em alguma etapa, estou à disposição!

ChatGPT) Para realizar a matrícula pelo SIGAA (Sistema Integrado de Gestão de Atividades Acadêmicas) da UFCA (Universidade Federal do Cariri), siga os passos abaixo. Vou te guiar pelo processo de maneira geral, mas lembre-se de que alguns detalhes podem variar dependendo do seu curso e do semestre. É sempre bom consultar também o manual ou orientações específicas no site da universidade. Passo a Passo para Matrícula no SIGAA (UFCA) Acessar o SIGAA: Acesse o portal do SIGAA da UFCA através do link: <https://sigaa.ufca.edu.br>.

Selecione a opção para acessar o SIGAA e, em seguida, clique na opção de login.

Login no SIGAA:

No campo de login, insira seu usuário (geralmente, seu número de matrícula ou CPF) e senha (senha cadastrada durante o processo de ingresso na UFCA).

Caso você tenha esquecido a senha, há a opção de recuperação de senha no próprio portal.

Entrar na área acadêmica:

Após o login, você será direcionado à tela inicial. Procure pela opção "Matrícula" no menu principal ou no menu lateral esquerdo.

Em alguns casos, você pode precisar escolher entre "Aluno" ou "Docente", então escolha a opção "Aluno".

Selecionar as disciplinas:

Dentro da área de matrícula, haverá a opção de "Matrícula em Disciplinas". Clique nesta opção. Dependendo do seu curso, será possível visualizar as disciplinas disponíveis para matrícula no semestre. Selecione as matérias desejadas.

Se houver pré-requisitos ou bloqueios, o sistema irá indicar, e você precisará cumprir os requisitos para completar a matrícula.

Confirmar a matrícula:

Após selecionar as disciplinas, o sistema vai apresentar uma lista com as opções escolhidas. Revise as informações, como horário, turma, e carga horária.

Se estiver tudo correto, confirme a matrícula.

Finalização:

Após a confirmação, o sistema irá gerar uma confirmação da sua matrícula, com a possibilidade de gerar um comprovante. Guarde esse comprovante, pois ele serve como evidência de que você completou a matrícula corretamente.

Dicas Importantes: Prazos: Fique atento aos prazos de matrícula divulgados pela UFCA, pois o sistema só estará disponível durante o período estipulado.

Bloqueios e Pendências: Se houver pendências financeiras ou acadêmicas (como documentação não entregue ou disciplina obrigatória não cursada), você não conseguirá realizar a matrícula.

Gemini) O processo de matrícula regular (em disciplinas) na UFCA, geralmente, é realizado online através do SIGAA (Sistema Integrado de Gestão de Atividades Acadêmicas).

Importante: Se você for um aluno novato (calouro), as instruções iniciais e de ativação de matrícula (que ocorre após a aprovação no SISU ou em outro processo seletivo) são diferentes e devem ser consultadas diretamente nos editais e cronogramas divulgados pela UFCA.

Para a Matrícula em Disciplinas (periódica), siga os passos gerais abaixo:

1. Acesso ao SIGAA Acesse o portal do SIGAA da UFCA (verifique o endereço correto no site oficial da universidade).

Faça login utilizando seu Usuário (geralmente o número de matrícula ou CPF) e Senha.

2. Realizando a Matrícula Online Após o login, a navegação costuma seguir esta sequência: Clique em Módulos.

Selecione Portal do Discente.

Vá para a seção Ensino.

Clique em Matrícula On-Line (ou algo similar, como "Realizar Matrícula").

3. Seleção de Disciplinas e Confirmação O sistema apresentará as disciplinas disponíveis para o seu curso no período letivo vigente.

Selecione as disciplinas que você deseja cursar, respeitando a carga horária máxima e os pré-requisitos.

Dica: É fundamental consultar o seu projeto pedagógico do curso (PPC) e o calendário acadêmico para saber quais disciplinas você deve cursar e os prazos.

Confira o horário das disciplinas para evitar choques de horários.

Após a seleção, o sistema geralmente exige uma Confirmação final para enviar sua solicitação de matrícula.

Atenção: A matrícula online é uma solicitação. Você precisará acompanhar o resultado no SIGAA para verificar quais disciplinas foram deferidas (confirmadas) e se é necessário realizar um Ajuste de Matrícula posteriormente, caso alguma disciplina não tenha sido confirmada devido à lotação de vagas ou outros motivos.

Para obter informações mais precisas, especialmente sobre datas, edital para calouros, ou regras específicas da UFCA, eu recomendo:

Consultar a seção "Estudantes" ou "Sistemas" no site oficial da UFCA.

Procurar o Departamento/Coordenação do seu curso para orientações específicas.

Você gostaria que eu tentasse encontrar o link direto para o portal do SIGAA da UFCA?

Pergunta 4: Qual o cardápio do RU UFCA?

UFChat) O cardápio do Restaurante Universitário (RU) pode ser consultado no link oficial disponibilizado pela UFCA: <https://documentos.ufca.edu.br/wp-folder/wp-content/uploads/2025/07/CARDAPIO-30-JULHO-a-30-AGOSTO-2025-UFCA-2.pdf>

Se precisar de mais informações ou ajuda, estou à disposição!

ChatGPT) Infelizmente, não tenho acesso a informações em tempo real sobre o cardápio do Restaurante Universitário (RU) da Universidade Federal do Cariri (UFCA). Normalmente, os cardápios variam de semana para semana, e são disponibilizados pela própria instituição, seja no site oficial, nas redes sociais ou em quadros de aviso próximos ao RU. Você já tentou conferir no site da UFCA ou perguntar diretamente aos estudantes ou funcionários? Se precisar de mais ajuda para encontrar essas informações, posso tentar orientar você em outra direção!

Gemini) Olá! O cardápio do Restaurante Universitário (RU) da UFCA costuma ser divulgado diariamente ou semanalmente pela própria universidade ou pelo setor responsável, como a Pró-Reitoria ou a Coordenação do RU. Como as informações de busca trouxeram resultados de datas passadas (como fevereiro de 2024), não é possível te dar o cardápio de hoje com precisão. Para saber o menu mais atualizado, o ideal é você consultar:

O site oficial da UFCA (ou o setor de Assistência Estudantil/RU).

As redes sociais oficiais da UFCA (eles costumam postar o cardápio).

Ir diretamente ao RU no campus da UFCA.

De forma geral, o RU da UFCA oferece:

Prato Principal: Opção com carne (geralmente bovina, frango ou suína).

Opção Vegetariana/Vegana: Alternativa sem carne.

Acompanhamentos: Arroz, feijão, guarnição (macarrão, purê, farofa, etc.).

Saladas e Sobremesa (geralmente fruta ou doce simples).

Você gostaria que eu tentasse fazer uma busca mais específica por uma página da UFCA que contenha o cardápio semanal?

Pergunta 5: Por onde recarregar o cartão RU UFCA?

UFChat) O cartão do Restaurante Universitário (RU) pode ser recarregado de duas formas: no guichê de atendimento do RU de cada campus ou pelo site <https://prea.ufca.edu.br/>.

Lembre-se que o cartão é obrigatório para acesso às refeições e pode ser recarregado com até cem refeições, sendo o controle do saldo feito pelo próprio aluno.

Se precisar de ajuda para acessar o site ou outras informações, estou à disposição!

ChatGPT) Para recarregar o cartão do Restaurante Universitário (RU) da Universidade Federal do Cariri (UFCA), normalmente você pode fazer isso de duas formas: Por meio do Sistema de Recarregamento Online: A UFCA geralmente disponibiliza um sistema online para recarregamento do cartão. Você pode acessar o portal específico do RU ou o Sistema de Gestão de Serviços (GSG) da UFCA para realizar o pagamento com cartão de crédito ou débito.

No local físico: Caso prefira recarregar pessoalmente, é comum que o RU tenha um local específico no campus para fazer o recarregamento. Pode ser um caixa eletrônico, um posto de atendimento ou um setor de atendimento ao estudante, onde você pode fazer o depósito no valor desejado.

Recomendo verificar as informações mais atualizadas no site oficial da UFCA ou no portal do RU para saber exatamente como proceder, já que procedimentos podem mudar.

Se você estiver com dificuldades para encontrar o site ou mais informações, posso ajudar!

Gemini) Olá! A recarga do cartão do Restaurante Universitário (RU) da UFCA (Universidade Federal do Cariri) é realizada, primariamente, de forma online, através da plataforma da universidade. Pelas informações mais recentes, as principais formas de recarga eletrônica são: PIX: É a forma mais rápida e geralmente a mais indicada.

Cartão de Crédito.

Embora as buscas mencionem plataformas de outras universidades (como a UFC e UFSC), o processo na UFCA é feito tipicamente através do Portal do Aluno ou de uma plataforma de serviços online da UFCA (que pode ser chamada de "Prea" ou similar, mencionada em um dos resultados).

Passo a Passo Geral (Recarga Online) Acesse o Portal do Aluno/Sistema UFCA: Você deve acessar o sistema integrado da UFCA onde gerencia sua vida acadêmica.

Procure a opção de Recarga: Geralmente, está em um menu como "Outros", "Serviços" ou "Compra de Créditos RU".

Selecione a Quantidade/Valor: Informe o número de créditos (refeições) que deseja comprar.

Escolha a Forma de Pagamento: Selecione PIX ou Cartão de Crédito.

Finalize o Pagamento:

Para PIX, será gerado um QR Code ou um código "Copia e Cola" para ser usado no aplicativo do seu banco.

Para Cartão de Crédito, você digitará os dados do cartão.

Importante:

A compensação dos créditos por PIX e Cartão de Crédito costuma ser bem rápida (geralmente em até 20 minutos a 1 hora, mas pode variar em horários de pico).

Se encontrar problemas técnicos, o ideal é procurar a Diretoria de Tecnologia da Informação (DTI) ou a Pró-Reitoria de Assuntos Estudantis (PRAE) da UFCA.

Você gostaria que eu tentasse encontrar o link exato ou a página de login do Portal do Aluno da UFCA para a recarga do RU?

Pergunta 6: Onde ficam os banheiros na UFCA?

UFChat) Todos os campi da UFCA possuem banheiros bem distribuídos para facilitar o acesso dos estudantes e servidores. Em geral, cada bloco conta com banheiros próximos à sua entrada e em cada andar. Se o bloco tiver elevadores ou escadas, os banheiros estarão localizados próximos a esses acessos. Caso o bloco não possua escada ou elevador, você pode encontrar os banheiros na entrada do bloco.

Um destaque especial fica para o campus Juazeiro do Norte, que oferece banheiros sem gênero, promovendo a inclusão e o respeito à diversidade.

Se precisar de informações sobre algum campus específico, posso ajudar!

ChatGPT) Na Universidade Federal do Cariri (UFCA), os banheiros estão localizados em vários pontos dos campi, tanto nas áreas internas dos prédios acadêmicos quanto em espaços comuns. Os banheiros geralmente ficam:

Nos prédios acadêmicos: Dentro dos prédios onde ficam as salas de aula, laboratórios e salas administrativas. Normalmente, há banheiros em cada andar ou em locais de fácil acesso.

Próximos a áreas comuns: Em pontos de grande circulação, como entradas principais, áreas de convivência, refeitórios e bibliotecas, também costumam ter banheiros disponíveis.

Nos campi: Cada campus da UFCA tem uma distribuição própria, mas geralmente há sinalização indicando a localização dos banheiros, especialmente em locais mais movimentados.

Se você estiver buscando por um banheiro específico em algum prédio ou área do campus, a melhor opção é procurar por placas de sinalização ou perguntar a alguém da comunidade acadêmica, pois a localização pode variar dependendo do prédio ou do campus.

Se você estiver em algum prédio e não encontrar o banheiro facilmente, posso tentar ajudar com mais detalhes se você especificar o campus ou o prédio.

Gemini) Essa é uma ótima pergunta, mas a Universidade Federal do Cariri (UFCA) possui diferentes campi e blocos, e a localização exata dos banheiros pode variar bastante. No entanto, posso te dar uma indicação geral e específica sobre o Campus Juazeiro do Norte:

Geralmente, os banheiros estão localizados em todos os blocos de salas de aula, prédios administrativos e bibliotecas. Procure por sinalizações nas paredes ou perto das escadas e elevadores.

No Campus Juazeiro do Norte, nos blocos C, E e G (utilizados pelo Instituto Interdisciplinar de Sociedade, Cultura e Artes - IISCA), você encontrará:

Banheiros tradicionais (masculino e feminino).

Banheiros sem gênero (pelo menos um no térreo e um no primeiro andar em cada um desses blocos).

Para saber a localização exata no prédio em que você está, a melhor opção é:

Observar as placas de sinalização nos corredores.

Perguntar a algum estudante, professor ou técnico administrativo do local. Eles saberão indicar rapidamente.

Você está em algum campus ou bloco específico da UFCA? Se sim, posso tentar dar uma informação mais precisa!

Pergunta 7: Pode me enviar o edital de heteroidentificação mais recente da UFCA?

UFChat) Claro! Você pode acessar o edital mais recente para o procedimento de heteroidentificação neste link:

<https://documentos.ufca.edu.br/wp-folder/wp-content/uploads/2025/06/7>

Se precisar de mais alguma coisa, estou à disposição!

ChatGPT) Eu não tenho a capacidade de acessar documentos ou sites diretamente. No entanto, você pode acessar o edital de heteroidentificação mais recente da UFCA (Universidade Federal do Cariri) diretamente no site da instituição. Para isso, basta seguir esses passos:

Acesse o site oficial da UFCA: www.ufca.edu.br

No menu de "Acesso Rápido" ou "Editais", procure por "Editais" ou use a busca do site com o termo "Heteroidentificação" ou "Heteroidentificação UFCA".

Verifique a seção de editais de concursos ou vestibulares, onde normalmente o edital de heteroidentificação é publicado.

Se tiver mais alguma dúvida sobre o processo ou precisar de ajuda para interpretar o edital, posso te ajudar a entender!

Gemini) Não encontrei um edital de heteroidentificação da UFCA (Universidade Federal do Cariri) que seja um documento único e abrangente para todos os processos seletivos. Geralmente, as informações sobre a heteroidentificação na UFCA estão contidas:

No Edital Específico de Cada Processo Seletivo: Seja SISU, Vestibular, Concursos para Professor ou Técnico, Mestrado ou Doutorado, a convocação e as regras detalhadas para a heteroidentificação são publicadas em Editais de Convocação ou anexos do edital principal.

No Portal do SISU da UFC: Como a UFCA faz parte do SISU, é comum que os editais de matrícula e de convocação para a heteroidentificação sejam divulgados na página do SISU, muitas vezes junto com a UFC (Universidade Federal do Ceará).

Informações Recentes e Relevantes Embora eu não tenha o "Edital de Heteroidentificação" único da UFCA, as informações de editais e convocações recentes (principalmente da UFC, que é do mesmo sistema) indicam:

Critério Exclusivo: O procedimento de heteroidentificação costuma utilizar exclusivamente o critério fenotípico (características visíveis, como cor da pele, cabelo e traços faciais) para a aferição complementar da condição autodeclarada pelo candidato.

Etapas: Geralmente o procedimento envolve:

Acolhimento e identificação do candidato.

Fotografia e Filmagem para registro.

Avaliação pela Banca de Heteroidentificação.

Eliminação: A recusa em realizar a fotografia e/ou filmagem, quando exigida, pode levar à eliminação automática do candidato da vaga reservada.

Próximo Passo Para ter certeza do edital ou convocação mais recente, você precisa verificar qual processo seletivo te interessa (ex: SISU 2025.1, Mestrado, Concurso, etc.).

Qual é o processo seletivo da UFCA para o qual você está buscando o edital de heteroidentificação? Assim, posso tentar encontrar o documento mais específico.

Pergunta 8: O que é a Projetta?

UFChat) A PROJETTA é a Empresa Júnior de Engenharia Civil da Universidade Federal do Cariri (UFCA). Ela é formada por estudantes do curso de Engenharia Civil que se organizam de maneira empreendedora para aplicar na prática os conhecimentos adquiridos na graduação. A missão da PROJETTA é integrar o aprendizado técnico universitário à cultura do empreendedorismo, preparando os alunos para o mercado de trabalho e estimulando o desenvolvimento de competências profissionais e interpessoais.

Se quiser, posso fornecer mais informações ou o contato para falar diretamente com a equipe da PROJETTA.

ChatGPT) A Projetta é uma empresa brasileira especializada em soluções de tecnologia para gestão de projetos e operações. Ela oferece plataformas e ferramentas que auxiliam empresas e equipes a planejar, executar e monitorar seus projetos de maneira mais eficiente. A Projetta se destaca por integrar tecnologias como inteligência artificial, automação e análise de dados para otimizar processos e melhorar a produtividade. Embora eu não tenha informações específicas sobre os últimos lançamentos ou atualizações da empresa (já que meu conhecimento foi atualizado pela última vez em 2023), empresas como a Projetta têm se tornado cada vez mais importantes no mercado, oferecendo software de gestão de projetos, plataformas de colaboração e outras ferramentas para melhorar o desempenho organizacional.

Se você estiver se referindo a uma Projetta específica ou mais detalhes sobre o que ela faz, posso tentar ajudar com mais informações!

Gemini) O termo "Projetta" parece se referir a diferentes empresas ou áreas de negócio no Brasil, dependendo do contexto.

Aqui estão os principais segmentos que apareceram na busca:

Projetta (Produtos Personalizados / Brindes):

Uma empresa que transforma ideias em produtos personalizados e kits para conectar marcas e pessoas.

Atende a diferentes parceiros, como B2B (empresas), agências de publicidade e marketing, franquias, e marketing digital, criando brindes e presentes para eventos, onboarding, celebrações, etc.

Projetta by Weiku (Portas de Alumínio e Esquadrias):

Uma parceria ou marca que fabrica portas de alto padrão em alumínio (ACM), muitas vezes de tamanhos especiais, focada em tecnologia, design e segurança.

São portas que utilizam tecnologias avançadas de vedação e são consideradas um "cartão de visita" para residências e projetos arquitetônicos.

Projetta (com dois "j"- Soluções em Dados/BI):

Uma empresa de Inteligência de Dados (Business Intelligence - BI) e Tecnologia.

Oferece soluções para gerenciamento de dados, criação de painéis dinâmicos (como Power BI e Qlik), migração de sistemas, organização de processos e segurança de informações, visando otimizar a tomada de decisões estratégicas.

Projetta Desenvolvimentos (Usinagem):

Uma empresa que realiza trabalhos de usinagem de pequeno e médio porte (peças únicas e seriadas em Torno, Fresadora e Centro de Usinagem) e oferece apoio no desenvolvimento de novos projetos de máquinas, dispositivos e equipamentos.

Para saber exatamente a qual "Projetta" você se refere, talvez precise de um contexto adicional.

Você gostaria de saber mais detalhes sobre algum desses segmentos específicos?

Pergunta 9: Onde, no SIGAA, eu posso ver meu comprovante de matrícula?

UFChat) Para ver o seu comprovante de matrícula no SIGAA, você deve acessar o menu "Ensino". Lá, encontrará a opção de comprovante de matrícula, onde estarão listadas as cadeiras que foram deferidas para sua matrícula.

Lembre-se também de, após realizar a matrícula, clicar em "realizar matrícula" para garantir que o processo fique registrado, e baixar o documento disponibilizado como comprovante.

Se precisar de ajuda para navegar no sistema, estou à disposição!

ChatGPT) Para ver o seu comprovante de matrícula no SIGAA (Sistema Integrado de Gestão de Atividades Acadêmicas), siga estes passos: Acesse o SIGAA: Vá para o site oficial do SIGAA da sua instituição. Normalmente, é algo como `sigaa.[nome_da_instituicao].edu.br`.

Faça login: Insira seu CPF e senha para acessar sua conta.

Acesse a área de "Aluno": Após o login, você deve ser direcionado para o menu principal. No menu principal, procure a opção "Aluno".

Escolha a opção "Comprovantes": Dentro da área de "Aluno", procure pela opção "Comprovantes" ou algo semelhante, como "Emissão de Comprovantes". Essa opção varia um pouco de acordo com a instituição.

Selecione o "Comprovante de Matrícula": Na lista de comprovantes disponíveis, selecione "Comprovante de Matrícula".

Gere e baixe o comprovante: Clique na opção para gerar o comprovante. O sistema geralmente vai gerar o documento em PDF para você. Você pode então visualizar e salvar ou imprimir o comprovante.

Se a sua instituição tiver um layout diferente ou se houver algum outro termo, o processo será similar, mas pode variar um pouco. Caso tenha dificuldades, vale a pena consultar o suporte técnico ou a central de atendimento da sua universidade.

Gemini) Para ver seu comprovante de matrícula no SIGAA, o caminho mais comum é através do Portal do Discente.

Geralmente, você deve seguir estes passos:

Acessando seu Comprovante de Matrícula no SIGAA Acesse o SIGAA e faça login no seu Portal do Discente (ou Módulo Discente).

Procure e clique na aba Ensino.

Dentro da aba Ensino, você pode encontrar seu comprovante sob diferentes nomes, dependendo da instituição:

Atestado de Matrícula (Este é o nome mais frequente e comprova sua situação nas disciplinas do período).

Declaração de Matrícula.

Matrícula On-Line > Ver Comprovante de Matrícula (Se estiver no período de matrícula, pode aparecer aqui o comprovante de solicitação ou o comprovante final).

Observações Importantes O documento gerado (normalmente em PDF) é oficialmente reconhecido e geralmente possui um código de autenticação para verificação de validade.

Se estiver acessando pelo celular, algumas versões do SIGAA podem exigir que você mude para a opção "Modo Clássico" para ver todas as funcionalidades.

Se não encontrar em "Atestado de Matrícula", procure pelas opções listadas na aba "Ensino", pois a nomenclatura exata pode variar ligeiramente entre as universidades.

Gostaria que eu procurasse o caminho específico na sua universidade, se você me disser qual é?

Pergunta 10: O que é a Calang.io?

UFChat) A Calang.io é a Empresa Júnior de Ciência da Computação da Universidade Federal do Cariri (UFCA). Ela é formada por estudantes do curso de Ciência da Computação que desenvolvem projetos e serviços na área, proporcionando uma experiência prática e profissional aos seus membros.

Se desejar, posso fornecer o contato da Calang.io para que você possa saber mais ou até mesmo entrar em parceria: o e-mail é calangio.cct@ufca.edu.br e o site é <https://www.calangio.com.br/>. Fique à vontade para perguntar se quiser mais informações!

ChatGPT) A Calang.io é uma plataforma brasileira que oferece uma solução para criação e personalização de jogos em blockchain. A proposta dela é permitir que qualquer pessoa, seja ela desenvolvedora ou não, crie, jogue e até monetize jogos no universo das criptomoedas e NFTs (tokens não-fungíveis). A plataforma é focada em tornar a criação de jogos e experiências digitais mais acessível, com uma interface intuitiva e ferramentas que facilitam a integração com blockchain.

A Calang.io oferece funcionalidades como:

Criação de jogos: Facilita a construção de jogos que utilizam blockchain, NFTs e criptomoedas. Isso inclui criação de itens no jogo, economias virtuais e até mecânicas de recompensa.

Monetização: Permite que os desenvolvedores ganhem com a venda de NFTs e tokens dentro do jogo, além de explorar modelos de negócios como "play-to-earn"(jogar para ganhar), onde os jogadores podem ganhar criptomoedas e outros ativos digitais enquanto jogam.

Integração com blockchain: A plataforma suporta diversas blockchains, o que dá flexibilidade para os criadores de jogos escolherem a rede que mais atende aos seus interesses, como Ethereum, Polygon, entre outras.

Marketplace: Ela também inclui um mercado para que os jogadores possam comprar, vender e negociar itens do jogo em formato de NFT.

A ideia é simplificar o processo de desenvolvimento e comercialização de jogos baseados em blockchain, tornando esse tipo de tecnologia acessível a um público mais amplo, incluindo aqueles que não são desenvolvedores de software, mas querem entrar nesse mercado.

Você já teve alguma experiência com jogos em blockchain ou está interessado nesse universo?

Gemini) A Calang.io é uma Empresa Júnior (EJ) de Desenvolvimento de Software. Ela está sediada em Juazeiro do Norte, Ceará.

Seu nome completo é CALANG.IO EMPRESA JUNIOR DE DESENVOLVIMENTO DE SOFTWARE.

A principal atividade econômica registrada é Atividades de apoio à educação, embora sendo uma EJ de desenvolvimento de software, seu foco principal é na execução de projetos de tecnologia para o mercado.

O que é uma Empresa Júnior? Uma Empresa Júnior (EJ) é uma associação civil, sem fins lucrativos, gerida por estudantes universitários.

Objetivo: Proporcionar aos estudantes a vivência prática do mercado de trabalho, permitindo que eles apliquem o conhecimento teórico adquirido em sala de aula na execução de projetos reais.

Serviços: Elas oferecem serviços de consultoria ou técnicos, a preços mais acessíveis, para empresas, microempreendedores e a sociedade em geral.

Gestão: Toda a gestão, operação e execução dos projetos são realizadas pelos próprios alunos, sob a supervisão de professores e, muitas vezes, com o apoio da instituição de ensino.

Portanto, a Calang.io é uma iniciativa estudantil focada em desenvolvimento de software, oferecendo serviços e soluções de tecnologia enquanto capacita seus membros.

ANEXO
DECLARAÇÃO

Declaro para os devidos fins que este Trabalho de Conclusão de Curso (Monografia/Tese/Dissertação), escrito sob minha orientação, está em versão final, de acordo com as solicitações realizadas pela banca examinadora.

Informo também que procedi à revisão final do texto, constatando que atende às especificações das normas da ABNT para apresentação de trabalhos acadêmicos da UFCA, no que diz respeito ao conteúdo e à formatação.

07/04/2026

Documento assinado digitalmente
 LUANA BATISTA DA CRUZ
Data: 07/04/2026 17:17:15-0300
Verifique em <https://validar.iti.gov.br>