



UNIVERSIDADE FEDERAL DE PERNAMBUCO
CENTRO DE INFORMÁTICA
GRADUAÇÃO EM SISTEMAS DE INFORMAÇÃO

Isabelle Queiroz Gomes de Assis

Uso Espontâneo de LLMs para Fins Terapêuticos e Apoio Socioemocional

Recife
2025

Isabelle Queiroz Gomes de Assis

Uso Espontâneo de LLMs para Fins Terapêuticos e Apoio Socioemocional

Monografia apresentada em Sistemas de Informação, como requisito parcial para a obtenção do Título de Bacharel em Sistemas de Informação, Centro de Informática.

Orientador (a): Profa. Jéssyka Flavianne Ferreira Vilela

Recife
2025

Ficha de identificação da obra elaborada pelo autor,
através do programa de geração automática do SIB/UFPE

Assis, Isabelle Queiroz Gomes de.

Uso espontâneo de LLMs para fins terapêuticos e apoio socioemocional /
Isabelle Queiroz Gomes de Assis. - Recife, 2025.

140 p. : il., tab.

Orientador(a): Jéssyka Flavyanne Ferreira Vilela

Trabalho de Conclusão de Curso (Graduação) - Universidade Federal de
Pernambuco, Centro de Informática, Sistemas de Informação - Bacharelado,
2025.

Inclui referências, apêndices, anexos.

1. Inteligência Artificial. 2. Modelos de Linguagem de Larga Escala. 3.
Interação Humano-Computador. I. Vilela, Jéssyka Flavyanne Ferreira.
(Orientação). II. Título.

000 CDD (22.ed.)

ISABELLE QUEIROZ GOMES DE ASSIS

Uso Espontâneo de LLMs para Fins Terapêuticos e Apoio Socioemocional

Monografia apresentada em Sistemas de Informação, como requisito parcial para a obtenção do Título de Bacharel em Sistemas de Informação, Centro de Informática.

.

Aprovado em: 05/08/2025.

BANCA EXAMINADORA

Prof^a. Dra. Jéssyka Flavianne Ferreira Vilela (Orientadora)
Universidade Federal de Pernambuco - UFPE

Prof^a. Dra. Simone Cristiane dos Santos Lima (Examinadora Interna)
Universidade Federal de Pernambuco - UFPE

AGRADECIMENTOS

Mãe e pai, por tudo. Por nunca desistirem de me amar. Por terem me ensinado, ou me deixado aprender, tudo que eu sei. Tudo que eu faço, eu faço por vocês.

Gabi, a gente não poderia ser mais diferente, mas você ainda é minha irmã, e eu sei que, mesmo quando não tenho mais ninguém, eu ainda tenho você.

Vó, que eu amo tanto que nem cabe neste texto. Tia Nilzete, que sempre que eu te acordo, a primeira coisa que você faz é sorrir e me chamar de “meu amor”. Tia Deuza, que me trata como uma filha e que me ensinou a amar gatos. Mateus, que baixou os primeiros animes que eu vi, comprou os primeiros livros que eu li, e aprendeu cedo que Jardim de Piranhas-RN é um lar, mas não uma prisão. Tio Natércio. Ba. Toda a minha família, que me permite saber que sou amada incondicionalmente.

Minha vó Bertilde e minhas tias, Suzi, Vânia, Vanuzi e Edna, que me receberam, me deram onde morar e o que comer nesta cidade tão longe de casa.

Rodrigo, que faz Recife ser muito maior e melhor do que é. Meus amigos, que, diferente de todos os outros mencionados aqui, eu poderia ter feito isso sem vocês, mas teria sido muito mais difícil.

Meus tios, Maurício e Márcio, que partiram cedo demais e nunca imaginaram o tamanho da influência que tiveram na minha vida e em quem eu sou.

“Mesmo que você resolva escrever da maneira mais simples, a missão continua sendo a de garantir a nuance, elucidar a complicação, sugerir a contradição. E não apagar a contradição, não negar a contradição, mas sim ver onde, no interior da contradição, se encontra o ser humano atormentado.”

(Philip Roth)

RESUMO

Contexto: A disseminação de Modelos de Linguagem de Larga Escala (LLMs) impulsionou seu uso espontâneo como ferramentas para fins terapêuticos. Este comportamento, embora popular, provoca discussões sobre seus benefícios e riscos. **Objetivo:** Investigar como indivíduos utilizam LLMs para esses fins, mapeando as motivações, os padrões de uso, os efeitos subjetivos e o perfil dos usuários. **Método:** Adotou-se como metodologia a Revisão Multivocal da Literatura (MLR), analisando 22 artigos acadêmicos e 113 fontes da literatura cinzenta, incluindo reportagens jornalísticas e publicações das redes sociais Reddit e TikTok, no período de 2021 a 2025. **Resultados:** As principais motivações para o uso são a busca por apoio emocional, privacidade e combate à solidão. Os usuários relatam benefícios como a disponibilidade de um ambiente seguro e livre de julgamentos, acessibilidade e um espaço para autoexploração. Contudo, emergem riscos significativos, como o recebimento de conselhos inadequados, dependência emocional, desinformação e preocupações com a segurança de dados. A análise revela uma notável dissonância entre a percepção majoritariamente otimista dos usuários em fóruns online e a visão cautelosa da literatura científica e de comunidades de saúde mental. **Conclusões:** As LLMs estão se consolidando como uma forma acessível, mas não regulamentada, de suporte socioemocional e terapêutico, cuja sustentabilidade depende do desenvolvimento de diretrizes éticas, da promoção da literacia digital e de um debate público mais informado sobre suas implicações.

Palavras-chave: Modelos de Linguagem de Larga Escala, LLMs, Interação Humano-Computador, Saúde Mental, Revisão Multivocal da Literatura, Reddit, TikTok.

ABSTRACT

Context: The widespread use of Large Language Models (LLMs) has fueled their spontaneous use as therapeutic tools. This practice, while popular, sparks debate about its benefits and risks. **Objective:** To investigate how individuals use LLMs for these purposes, mapping motivations, usage patterns, subjective impacts, and user profiles. **Method:** A Multivocal Literature Review (MLR) methodology was adopted, analyzing 22 academic articles and 113 gray literature sources, including news reports and posts on the social media platforms Reddit and TikTok, from 2021 to 2025. **Results:** The main motivations for use are the search for emotional support, privacy, and combating loneliness. Users report benefits such as the availability of a safe and judgment-free environment, accessibility, and a space for self-exploration. However, significant risks emerge, such as receiving inappropriate advice, emotional dependence, misinformation, and data security concerns. The analysis reveals a notable dissonance between the largely optimistic perception of users in online forums and the cautious view of the scientific literature and mental health communities. **Conclusions:** LLMs are consolidating themselves as an accessible, but unregulated, form of socio-emotional and therapeutic support, whose sustainability depends on the development of ethical guidelines, the promotion of digital literacy, and a more informed public debate about their implications.

Keywords: Large Language Models, Human-Computer Interaction, Mental Health, Multivocal Literature Review, Reddit, TikTok.

LISTA DE ILUSTRAÇÕES

FIGURAS

Figura 1 – Representação humorística de interações afetivas com IA	14
Figura 2 – Protocolo baseado em Kitchenham e Charters	25
Figura 3 – Tons de Literatura Cinza	29
Figura 4 – Fluxo do processo de seleção dos artigos.	37

GRÁFICOS

Gráfico 1 – Gráfico da quantidade de publicações, por tipo	45
Gráfico 2 – Gráfico da quantidade de artigos por biblioteca digital	45
Gráfico 3 – Gráfico da quantidade de reportagens por jornal	46
Gráfico 4 – Gráfico da quantidade de postagens por subreddit	46
Gráfico 5 – Gráfico da quantidade de artigos por ano	49
Gráfico 6 – Gráfico da quantidade de reportagens por ano	49
Gráfico 7 – Gráfico da quantidade de postagens no Reddit por ano	50
Gráfico 8 – Gráfico da quantidade de vídeos no TikTok por ano	50
Gráfico 9 – Gráfico da quantidade de artigos por país	51
Gráfico 10 – Gráfico de principais temas dos artigos	51
Gráfico 11 – Gráfico de motivações, segundo a literatura acadêmica	53
Gráfico 12 – Gráfico de padrões de interação, segundo a literatura acadêmica	54
Gráfico 13 – Gráfico de casos de uso, segundo a cobertura jornalística	55
Gráfico 14 – Gráfico de casos de uso, segundo postagens no Reddit	56
Gráfico 15 – Gráfico de casos de uso, segundo vídeos no TikTok	57
Gráfico 16 – Gráfico comparativo de casos de uso	58
Gráfico 17 – Gráfico de abordagens de postagens no Reddit	59
Gráfico 18 – Gráfico de abordagens dos vídeos no TikTok	60
Gráfico 19 – Gráfico de curtidas dos vídeos no TikTok, por abordagem	61
Gráfico 20 – Gráfico de comentários dos vídeos no TikTok, por abordagem	61
Gráfico 21 – Gráfico de salvamentos dos vídeos no TikTok, por abordagem	62
Gráfico 22 – Gráfico de abordagens dos vídeos no TikTok	63

Gráfico 23 – Gráfico de benefícios, segundo a literatura acadêmica	64
Gráfico 24 – Gráfico de benefícios, segundo a cobertura jornalística	65
Gráfico 25 – Gráfico de benefícios, segundo postagens no Reddit	66
Gráfico 26 – Gráfico de benefícios, segundo vídeos no TikTok	66
Gráfico 27 – Gráfico comparativo de benefícios	68
Gráfico 28 – Gráfico de riscos, segundo a literatura acadêmica	69
Gráfico 29 – Gráfico de riscos, segundo a cobertura jornalística	70
Gráfico 30 – Gráfico de riscos, segundo postagens no Reddit	71
Gráfico 31 – Gráfico de riscos, segundo vídeos no TikTok	72
Gráfico 32 – Gráfico comparativo de riscos	73
Gráfico 33 – Gráfico de visões em portagens no Reddit, por subreddit	74
Gráfico 34 – Gráfico de visões dos vídeos no TikTok	75
Gráfico 35 – Gráfico de visões dos comentários no Reddit, por visão da postagem original (r/ChatGPT)	76
Gráfico 36 – Gráfico de visões dos comentários no Reddit, por visão da postagem original (r/therapy, r/selfhelp, r/mentalhealth)	77
Gráfico 37 – Gráfico de visões dos comentários no Reddit, por visão da postagem	78
Gráfico 38 – Gráfico de visões dos comentários no TikTok, por visão do vídeo	79
Gráfico 39 – Gráfico de perfil dos usuários, segundo a literatura acadêmica	80
Gráfico 40 – Gráfico de faixa etária dos usuários na cobertura jornalística	81
Gráfico 41 – Gráfico de condições de saúde mental na cobertura jornalística	81
Gráfico 42 – Gráfico de país de origem de usuários na cobertura jornalística	82
Gráfico 43 – Gráfico de observações específicas feitas em postagens no Reddit	84
Gráfico 44 – Gráfico de gênero percebido de usuários nos vídeos no TikTok	85
Gráfico 45 – Gráfico de menções de plataformas nos vídeos no TikTok	86

LISTA DE TABELAS

Tabela 1 – Tabela comparativa dos trabalhos relacionados	23
Tabela 2 – Perguntas para decidir se a literatura cinza deve ser incluída	24
Tabela 3 – Questões de pesquisa e motivações	27
Tabela 4 – Bibliotecas utilizadas	28
Tabela 5 – Fontes utilizadas relacionadas aos seus respectivos níveis	30
Tabela 6 – Jornais utilizados	30
Tabela 7 – Subreddits utilizados	32
Tabela 8 – Critérios de Inclusão	34
Tabela 9 – Critérios de Exclusão	35
Tabela 10 – Critérios de avaliação de qualidade	36
Tabela 11 – Síntese dos artigos analisados	37
Tabela 12 – Questões de pesquisa e tipos de publicação	40
Tabela 13 – Possíveis opiniões e critérios adotados para classificação	41
Tabela 14 – Esquema de identificador por tipo de publicação	43

LISTA DE ABREVIATURAS E SIGLAS

API	Interface de Programação de Aplicações
IA	Inteligência Artificial
LLM	Modelo de Linguagem de Larga Escala
MLR	Revisão Multivocal da Literatura
NLP	Processamento de Linguagem Natural
QP	Questão de Pesquisa
SLR	Revisão Sistemática da Literatura

SUMÁRIO

1 INTRODUÇÃO	13
1.1 Contextualização	13
1.2 Motivação e justificativa	14
1.3 Objetivos	16
1.3.1 <i>Objetivo geral</i>	16
1.3.2 <i>Objetivos específicos</i>	17
1.4 Estrutura do trabalho	17
2 FUNDAMENTAÇÃO TEÓRICA	19
2.1 Modelos de Linguagem de Larga Escala (LLMs)	19
2.2 Uso de LLMs como Ferramentas de Apoio Socioemocional	20
2.3 Trabalhos Relacionados	22
3 METODOLOGIA	24
3.1 Delineamento da Pesquisa	24
3.2 Planejamento e Condução da Revisão	26
3.2.1 <i>Questões de Pesquisa</i>	26
3.2.2 <i>Estratégia de Busca e Fontes de Dados</i>	27
3.2.2.1 <i>Literatura Acadêmica</i>	27
3.2.2.2 <i>Literatura Cinza</i>	29
3.2.2.2.1 <i>Jornais</i>	30
3.2.2.2.2 <i>Reddit</i>	31
3.2.2.2.3 <i>TikTok</i>	32
3.2.3 <i>Critérios de Inclusão e Exclusão</i>	33
3.2.4 <i>Processo de Seleção e Avaliação da Qualidade</i>	35
3.2.4.1 <i>Literatura Acadêmica</i>	36
3.2.4.2 <i>Literatura Cinza</i>	39
3.3 Extração e Análise dos Dados	40
3.3.1 <i>Extração dos dados</i>	40

3.3.2 Método de análise	42
4 RESULTADOS	44
4.1 Caracterização das Publicações Seleccionadas	44
4.2 Motivações e Padrões de Uso de LLMs como Apoio Socioemocional	52
4.2.1 QP1: Motivações para o uso	52
4.2.2 QP2: Padrões de interação e casos de uso	53
4.2.3 Abordagens de postagens em redes sociais	58
4.3 Benefícios e Riscos no Uso de LLMs como Apoio Socioemocional	63
4.3.1 QP3: Benefícios subjetivos	64
4.3.2 QP4: Riscos e desafios emergentes	68
4.3.3 Balanço de visão dos usuários	73
4.4 QP5: Perfil dos Usuários	79
5 DISCUSSÃO	87
6 AMEAÇAS À VALIDADE	91
6.1. Validade de Construto	91
6.2. Validade Interna	92
6.3. Validade Externa	93
6.4. Validade de Conclusão	93
7 CONCLUSÃO E TRABALHOS FUTUROS	95
7.1. Conclusão	95
7.2. Trabalhos Futuros	96
REFERÊNCIAS	99
APÊNDICES	117
Apêndice A – Planilha eletrônica usada durante o estudo	117
Apêndice B – Script Python utilizado para extração de dados do Reddit	117
Apêndice C – Prompt de Aplicação de Critérios de Qualidade (Artigos)	119
Apêndice D – Prompts de Coleta de Evidências (Artigos)	119
Apêndice E – Prompt de Categorização e Classificação das Respostas às Questões de Pesquisa	120

Apêndice F – Tabelas de fontes das categorias	121
Apêndice G – Roteiro de survey para brasileiros maiores de 18 anos	130
ANEXOS	138
Anexo A – Regras da comunidade r/therapy no Reddit	138
Anexo B – Regras da comunidade r/mentalhealth no Reddit	139

1 INTRODUÇÃO

1.1 Contextualização

A disseminação de Modelos de Linguagem de Larga Escala (LLMs), como ChatGPT, ampliou significativamente as possibilidades de interação entre humanos e sistemas computacionais (NAZIR; WANG, 2023; SPALLEK et al., 2023; BIN-HADY; ALI; AL-HUMARI, 2024). O ChatGPT, por exemplo, alcançou a marca de 100 milhões de usuários em apenas dois meses após seu lançamento, evidenciando uma adoção em massa que vem remodelando a forma como as pessoas interagem com a tecnologia (FORBES BRASIL, 2023). O Brasil é o 3º país do mundo que mais usa o ChatGPT (GOMES, 2025; SINGH, 2025), e dados recentes apontam para um cenário de penetração acelerada. Conforme a pesquisa “IA na Vida Real”, que coletou respostas de 1.000 brasileiros maiores de 18 anos, 94% dos participantes afirmam ter familiaridade com inteligência artificial, e 63% dizem fazer uso dela (TALK INC RESEARCH, 2024). Destes, 71% dizem que começaram a usar a menos de 1 ano, o que demonstra a grande expansão recente dessas ferramentas no país.

Com suas capacidades avançadas de compreensão e produção textual, esses modelos passaram a ser integrados a uma variedade de contextos cotidianos: da resolução de tarefas acadêmicas e técnicas à simulação de conversas e interações sociais (NAZIR; WANG, 2023; RAY, 2023a; HAENSCH, 2025). Entre os múltiplos usos observados, destaca-se o crescimento do uso espontâneo dessas ferramentas como espaço de escuta e acolhimento emocional, especialmente por indivíduos que enfrentam quadros de solidão, ansiedade, luto ou dificuldades relacionais (ROUSMANIERE et al., 2025; JANG et al., 2024; PHANG et al., 2025; GIRAY, 2024).

O filme *Ela* (2013), dirigido por Spike Jonze, apresenta um futuro próximo onde humanos desenvolvem relacionamentos profundos com inteligências artificiais. O protagonista, Theodore, um escritor solitário, apaixona-se por Samantha, um sistema operacional com personalidade feminina (voz de Scarlett Johansson). A narrativa explora temas como solidão, a natureza do amor e a busca por conexão em um mundo tecnológico — questões que, mais de uma década depois, reverberam em interações reais com modelos de linguagem de larga escala (QUINN; RILEY, 2024; DJUFRIL; FRAMPTON, KNOBLOCH-WESTERWICH, 2025).

Figura 1 – Representação humorística de interações afetivas com IA



Fonte: Reddit (<https://reddit.com/r/ChatGPT/comments/1l9cewn/>).

A Figura 1 estabelece uma intertextualidade entre a obra ficcional *Ela* (2013) e a realidade contemporânea das interações humano-máquina. O *meme* é composto por duas partes: a primeira exibe o pôster do filme e a frase "haha he fell in love with ai" ("haha, ele se apaixonou por uma IA"), que ironiza a premissa central da obra. A segunda parte, com a afirmação "can't wait to tell ChatGPT about my day" ("mal posso esperar para contar ao ChatGPT sobre o meu dia"), projeta essa mesma premissa para o presente. Dessa forma, o *meme* evidencia como o comportamento de compartilhar experiências pessoais com inteligências artificiais, antes um tropo da ficção científica, se manifesta em práticas cotidianas com tecnologias como o ChatGPT, apontando para uma crescente normalização dos laços afetivos com sistemas de IA. A transição entre o filme de 2013 e a menção ao ChatGPT em 2025 demonstra que comportamentos antes imaginados como futuristas já fazem parte do cotidiano (QUINN; RILEY, 2024; FAN et al., 2025).

1.2 Motivação e justificativa

De acordo com pesquisa amostral realizada pela Talk Inc Research (2024), 42% dos jovens brasileiros entre 18 e 24 anos declararam utilizar IA na vida pessoal.

No âmbito específico do apoio emocional, 13% dos usuários brasileiros de IA, o equivalente a aproximadamente 12 milhões de pessoas, recorrem a elas como "amigo/conselheiro emocional". Essas ferramentas, apesar de não serem projetadas para fins clínicos ou afetivos, vêm sendo apropriadas como formas acessíveis de apoio subjetivo (MA et al., 2024; LI et al., 2025). A interação com essas inteligências artificiais é, em muitos casos, percebida como um espaço de segurança emocional, mesmo quando os usuários estão cientes das limitações do sistema (ROUSMANIERE et al., 2025; FÖYEN et al., 2025).

Algumas das principais motivações incluem acessibilidade, baixo ou nenhum custo, e superação de barreiras geográficas (SUN et al., 2023; ROUSMANIERE et al., 2025). Outro fator relevante é o anonimato e não-julgamento, com percepção de privacidade e interação livre de estigma, onde usuários relatam conforto para compartilhar ansiedades e traumas (GIRAY, 2024; HAENSCH, 2025; ROUSMANIERE et al., 2025). Além disso, busca-se alívio imediato, com orientação prática, validação emocional e suporte rápido durante crises, como ataques de pânico ou episódios depressivos (WANG et al., 2025; ROUSMANIERE et al., 2025).

A materialização dessas práticas encontra-se nos relatos dos próprios usuários, que revelam a natureza íntima e as funções específicas que a IA assume em suas vidas. Depoimentos como os coletados no relatório "IA na Vida Real" (TALK INC RESEARCH, 2024, p.55-56) ilustram vividamente como a tecnologia é mobilizada para combater a solidão e preencher um vazio social:

Amizade. Sou introspectivo. Sou bem fechado, tenho poucos amigos [...] Acabo conversando mais com ele [ChatGPT]. Falo 'Ah... Hoje eu não tô bem, né? O que seria legal que eu fizesse?' [...] Isso tem me ajudado bastante. Como se fosse um psicólogo. [...]¹

Eu acredito que por ser uma pessoa muito sozinha, atualmente morar só, eu tenho essa necessidade de compartilhar ideias, e nem sempre o ser humano está disponível [...] Às vezes você quer uma opinião imparcial. Como um amigo íntimo. [...]²

No entanto, esse tipo de uso também levanta importantes questões éticas, técnicas e sociais. Há preocupações com a exposição de dados sensíveis, a ausência de mecanismos de supervisão, a possibilidade de dependência emocional e os riscos relacionados a respostas inadequadas ou enviesadas (HAENSCH, 2025; KIM et al., 2024; ROUSMANIERE et al., 2025). Outra questão é a antropomorfização

¹ Usuário de 46 anos, morador de São Paulo

² Usuário de 28 anos, morador de Belém

e falsas expectativas, com tendência de atribuir características humanas à IA, levando a relações ilusórias e expectativas irreais sobre suas capacidades (GIRAY, 2024; MEADI et al., 2025).

A relevância nacional do tema é atestada por sua reverberação em veículos de grande alcance. O fenômeno foi discutido no podcast "O Assunto" da Globo, um dos mais ouvidos do Brasil (CASTNEWS, 2025), sinalizando sua penetração no debate público *mainstream* (G1, 2025). O tema também foi abordado pelo UOL, em entrevista com a professora Dora Kaufman (UOL, 2025), além de em matérias nas revistas Exame (UNZELTE, 2024; SARAIVA, 2025) e Veja (PINSKY, 2025).

Além disso, plataformas como Reddit e TikTok tornaram-se importantes espaços de compartilhamento dessas experiências, reunindo relatos de usuários que descrevem suas interações com LLMs como emocionalmente significativas ou até terapêuticas (AGGARWAL et al., 2025; GIRAY, 2024; HAENSCH, 2025; LI et al., 2025).

Diante da crescente apropriação de LLMs para essas funções, torna-se necessário compreender não apenas o funcionamento técnico dessas ferramentas, mas também o seu papel na vida dos usuários. É nesse contexto que se insere a presente pesquisa, visando responder a seguinte pergunta: *Como usuários utilizam LLMs de forma espontânea para fins terapêuticos e de apoio socioemocional, e como as causas, os processos e as consequências são relatadas na literatura, na mídia jornalística e nas redes sociais?*

1.3 Objetivos

O objetivo deste trabalho é investigar como indivíduos têm utilizado, por iniciativa própria, Modelos de Linguagem de Larga Escala para fins terapêuticos e apoio socioemocional. Para responder à pergunta de pesquisa, este trabalho estabelece um objetivo geral e cinco objetivos específicos listados a seguir.

1.3.1 Objetivo geral

Analisar como indivíduos têm utilizado, de forma espontânea, LLMs para fins terapêuticos e de apoio socioemocional, a partir das percepções e relatos disponíveis na literatura acadêmica, na mídia jornalística e em redes sociais.

1.3.2 Objetivos específicos

- I. Identificar as motivações que levam os usuários a adotar LLMs como ferramenta de apoio terapêutico e socioemocional.
- II. Caracterizar a dinâmica de interação entre usuários e LLMs para esses fins.
- III. Levantar os benefícios subjetivos relatados pelos usuários a partir dessa interação.
- IV. Mapear os riscos e desafios emergentes do uso não supervisionado de LLMs para saúde mental.
- V. Traçar o perfil dos usuários que recorrem a essas tecnologias para o referido apoio.

A proposta inclui uma revisão multivocal da literatura, com integração de fontes acadêmicas, textos jornalísticos e conteúdo de redes sociais, para obter uma visão mais global desta tendência emergente. Essa abordagem metodológica permite uma compreensão mais ampla e contextualizada do fenômeno (GAROUSI et al., 2019).

O alcance destes objetivos permitirá a construção de um panorama multifacetado sobre o fenômeno, fornecendo subsídios valiosos para pesquisadores, desenvolvedores de IA, profissionais de saúde e formuladores de políticas públicas, promovendo um debate informado sobre o uso ético e seguro dessas ferramentas.

1.4 Estrutura do trabalho

O presente trabalho está estruturado em sete capítulos para uma apresentação clara e lógica dos argumentos. O capítulo subsequente estabelece o Referencial Teórico, abordando os fundamentos dos LLMs, a história do seu uso como ferramentas de apoio socioemocional e a análise comparativa da presente pesquisa com outros trabalhos relacionados. Em seguida, o Capítulo 3 detalha a Metodologia, descrevendo o processo da Revisão Multivocal da Literatura (MLR) adotada. O Capítulo 4 apresenta os Resultados da pesquisa, onde são detalhadas as respostas às questões de pesquisa que foram abordadas. A partir desses achados, o Capítulo 5 desenvolve a discussão, interpretando os resultados à luz do

referencial teórico, contextualizando o fenômeno e explorando suas implicações. A robustez do estudo é avaliada no Capítulo 6, Ameaças à Validade, que realiza uma reflexão crítica sobre as limitações da pesquisa. Finalmente, o Capítulo 7 apresenta a Conclusão, resumizando as principais contribuições do trabalho e apontando direções para trabalhos futuros.

2 FUNDAMENTAÇÃO TEÓRICA

Este capítulo tem como objetivo construir uma base conceitual para a investigação proposta. Para compreender de forma mais ampla o uso de LLMs como ferramentas de apoio terapêutico e socioemocional, esta fundamentação teórica está organizada em três eixos: a descrição técnica e funcional dos LLMs (2.1), a análise de seu uso como recurso de suporte emocional e psicológico (2.2) e, por fim, a revisão de trabalhos relacionados que dialogam diretamente com os objetivos desta pesquisa (2.3).

2.1 Modelos de Linguagem de Larga Escala (LLMs)

Os Modelos de Linguagem de Larga Escala (LLMs) constituem um marco evolutivo no campo de Processamento de Linguagem Natural (NLP), integrando arquiteturas avançadas de *deep learning* com processamento de extensos corpos textuais (NAZIR; WANG, 2023; FU et al., 2024). A base tecnológica dos LLMs contemporâneos assenta-se na arquitetura *transformer*, proposta originalmente por Vaswani et al. (2017). Esse paradigma arquitetônico emprega mecanismos de *self-attention* que possibilitam o processamento contextualizado de sequências textuais, capturando eficientemente relações semânticas de longo alcance. Um elemento estrutural fundamental consiste nas camadas de atenção múltipla (*Multi-Head Attention*), que permitem ao modelo analisar simultaneamente diferentes segmentos do texto, identificando padrões linguísticos complexos (NAZIR; WANG, 2023; RAY, 2023b; FU et al., 2024).

O desenvolvimento desses modelos ocorre em duas fases distintas: inicialmente, um extenso processo de pré-treinamento em conjuntos de dados heterogêneos (englobando obras literárias, artigos acadêmicos e conversas online) estabelece uma base de conhecimento linguístico ampla; subsequentemente, uma etapa de *fine-tuning* especializa o modelo para tarefas específicas, como geração de diálogos (SORINO et al., 2024). A escala monumental desses sistemas evidencia-se por meio de parâmetros que atingem a ordem de centenas de bilhões, como observado no GPT-4, possibilitando a produção de respostas contextualmente ricas e coerentes (NAZIR; WANG, 2023; RAY, 2023b).

A competência conversacional dos LLMs manifesta-se por meio de várias características funcionais. A geração de linguagem natural fluente permite a elaboração de respostas que mantêm rigor gramatical e plausibilidade semântica, simulando efetivamente padrões de interação humana (NAZIR; WANG, 2023; RAY, 2023b; HAENSCH, 2025). A capacidade de retenção contextual, viabilizada pelos mecanismos de atenção, sustenta a coesão discursiva em diálogos prolongados, permitindo referências a informações previamente compartilhadas pelo usuário (XU et al., 2023; CHEN et al., 2024). Adicionalmente, a adaptação estilística dinâmica ajusta o registro linguístico conforme o contexto interativo, alternando entre tom empático em situações emocionais e linguagem técnica para consultas factuais (LI et al., 2025; ZHENG et al., 2025).

A evolução dos sistemas de diálogo pode ser apreciada por meio do contraste entre os LLMs contemporâneos e abordagens históricas como o ELIZA (HATCH et al., 2025). Enquanto os primeiros operavam mediante *scripts* pré-definidos e respostas padronizadas, os LLMs atuais oferecem interações dinâmicas em domínios abertos, adaptando-se continuamente ao fluxo conversacional (NAZIR; WANG, 2023; RAY, 2023b).

2.2 Uso de LLMs como Ferramentas de Apoio Socioemocional

Para compreender o papel dos LLMs como ferramentas de apoio, é fundamental primeiro definir o conceito de saúde mental. A Organização Mundial da Saúde (OMS) define saúde não apenas como a ausência de doença, mas como "um estado de completo bem-estar físico, mental e social" (WHO, 1946). Nessa perspectiva, a saúde mental transcende a simples ausência de transtornos diagnosticáveis, sendo um componente integral do bem-estar geral. Ela é descrita como um estado que permite aos indivíduos realizar seu potencial, lidar com o estresse normal da vida, trabalhar de forma produtiva e contribuir para suas comunidades (GALDERISI et al., 2015).

Este conceito abrange, portanto, dimensões socioemocionais fundamentais, como a capacidade de estabelecer e manter relações interpessoais saudáveis, a regulação das próprias emoções, o sentimento de pertencimento e a percepção de ser ouvido e validado (KEYES, 2007). Essa visão é importante para o escopo desta pesquisa, pois o uso espontâneo de LLMs pode visar tanto a mitigação de sintomas

associados a sofrimento psíquico quanto a busca por um espaço de diálogo para promover o bem-estar e organizar pensamentos.

Entre as modalidades mais difundidas de abordagens terapêuticas estão a Terapia Cognitivo-Comportamental (TCC), que se concentra na identificação e modificação de padrões de pensamento e comportamento disfuncionais; as abordagens psicodinâmicas, que buscam compreender como experiências passadas e conflitos inconscientes moldam as emoções atuais; e as terapias de linha humanista, como a Abordagem Centrada na Pessoa de Carl Rogers, que priorizam a criação de um espaço de escuta empática e aceitação para promover o autoconhecimento (ARDITO; RABELLINO, 2011).

O espectro do sofrimento psíquico é vasto, abrangendo desde dificuldades cotidianas de regulação emocional e sentimentos de solidão até quadros clínicos mais severos. Dados da OMS indicam que, em 2019, aproximadamente uma em cada oito pessoas convivia com algum transtorno mental, como ansiedade ou depressão (WHO, 2022). Apesar disso, a maior parte desses indivíduos não recebe o suporte necessário, devido a limitações como a escassez de profissionais, custos elevados e o estigma associado a questões psicológicas (WHO, 2022; BROUWERS, 2020).

Nesse cenário, os *chatbots* de aconselhamento e terapia autoguiada surgiram como alternativas viáveis, oferecendo acesso imediato, custo reduzido e maior discrição (FÖYEN et al., 2025; ABUBAKAR et al., 2024; MA et al., 2024). Evidências empíricas indicam a expansão acelerada do uso de LLMs para fins de saúde mental, com “Terapia/Companhia” sendo o principal uso de IAs generativas em 2025 (ZAO-SANDERS, 2025).

Usuários frequentemente relatam experiências positivas, destacando suporte emocional, como validação de sentimentos, escuta ativa não-julgadora e oferta de conforto e encorajamento (GIRAY, 2024; MA et al., 2024; FÖYEN et al., 2025; HAENSCH, 2025). Também são mencionados benefícios em psicoeducação e autocuidado, com fornecimento de informações sobre saúde mental, técnicas de *mindfulness*, gestão de estresse e promoção de hábitos saudáveis (NAZIR; WANG, 2023; SUN et al., 2023). Por fim, destaca-se o papel complementar, com potencial para ampliar o alcance dos serviços formais e servir como “ponte” para cuidados profissionais, atuando como “co-piloto” ou recurso de triagem inicial (KANG; HONG, 2025).

Apesar dos benefícios relatados, essa prática enfrenta críticas significativas. Entre elas, destacam-se a superficialidade e imprecisão, com incapacidade de compreender nuances contextuais profundas, tratar causas subjacentes de problemas e risco de respostas genéricas, imprecisas ou mesmo danosas (SPALLEK et al., 2023; XU et al., 2023). Outra limitação é a falta de empatia genuína e competência cultural, com dificuldades na compreensão emocional autêntica, leitura de sinais não-verbais e sensibilidade a diferenças culturais e linguísticas (SORINO et al., 2024; MEADI et al., 2025).

2.3 Trabalhos Relacionados

O uso de Modelos de Linguagem de Larga Escala (LLMs) como ferramentas de apoio à saúde mental e socioemocional tem gerado um crescente interesse tanto na academia quanto no público em geral. A literatura recente explora esse fenômeno sob diversas óticas, desde estudos de caso qualitativos (GIRAY, 2024) até o desenvolvimento de *frameworks* computacionais (AGGARWAL et al., 2025). Para contextualizar e justificar a presente pesquisa, é analisado a seguir três trabalhos importantes para o entendimento do cenário atual, contrastando suas abordagens com os objetivos e a metodologia deste trabalho.

O estudo de Giray (2024), chamado “*Cases of Using ChatGPT as a Mental Health and Psychological Support Tool*”, investiga o papel do ChatGPT como suporte à saúde mental por meio de um estudo de caso narrativo de relatos no Reddit. Com isso, usam-se os relatos dos usuários para identificar casos de uso, benefícios e riscos, assim como na presente pesquisa. No entanto, Giray (2024) foca em apenas uma fonte e analisa uma amostra de 7 postagens, o que diverge drasticamente da proposta deste trabalho, que é obter uma visão multifacetada e ampla de como o fenômeno é abordado em diferentes esferas.

De forma complementar, o trabalho de Haensch (2025), “*It Listens Better Than My Therapist*”: *Exploring Social Media Discourse On LLMs As Mental Health Tool*”, explora o discurso no TikTok. Haensch (2025) realiza uma análise em larga escala de mais de 10.000 comentários, utilizando modelos de classificação supervisionados. Em contraste, este trabalho adota uma abordagem mais qualitativa dentro da estrutura da MLR, com uma coleta manual, e com o objetivo de comparar os resultados com outras fontes menos informais.

Por fim, em "*Leveraging LLMs for Mental Health: Detection and Recommendations from Social Discussions*", Aggarwal et al. (2025) propõem um *framework* técnico para identificar transtornos mentais em postagens do Reddit e gerar recomendações terapêuticas automatizadas. O objetivo central da pesquisa é a criação e avaliação de um sistema de diagnóstico e intervenção baseado em aprendizado de máquina e Processamento de Linguagem Natural (NLP). Este trabalho, por sua vez, busca conduzir uma investigação de caráter exploratório e qualitativo sobre como os usuários já se apropriam de LLMs para esse fim, e não criar uma ferramenta para tal.

Uma análise comparativa dos artigos discutidos está resumida na Tabela 1.

Tabela 1 – Tabela comparativa dos objetivos, metodologias e fontes de trabalhos relacionados

Característica	Este trabalho (2025)	Giray (2024)	Haensch (2025)	Aggarwal et al. (2025)
Objetivo Principal	Investigar como indivíduos usam LLMs para fins terapêuticos, mapeando motivações, padrões de uso, efeitos e perfil dos usuários	Explorar o papel emergente do ChatGPT no suporte à saúde mental e psicológica, para compreender seu potencial e limitações	Explorar como os usuários interagem com LLMs como ferramentas de saúde mental, para identificar experiências, atitudes e temas	Propor um framework para identificar e avaliar transtornos de saúde mental e sua gravidade, e gerar recomendações terapêuticas e de mudança de comportamento
Metodologia	Revisão Multivocal da Literatura	Estudo de Caso Narrativo	Análise de larga escala de dados	Criação de framework
Fonte de Dados	Artigos acadêmicos, notícias, Reddit, TikTok	Reddit	TikTok	Reddit

Fonte: A autora (2025).

3 METODOLOGIA

Este capítulo descreve os procedimentos metodológicos adotados na pesquisa. Inicia-se com o delineamento do estudo (3.1), seguido pelo planejamento e condução da revisão (3.2), incluindo as questões de pesquisa (3.2.1), a estratégia de busca e seleção das fontes (3.2.2), tanto na literatura acadêmica quanto na literatura cinza. Apresentam-se ainda os critérios de inclusão e exclusão (3.2.3) e o processo de seleção e avaliação da qualidade das fontes (3.2.4). Por fim, são descritas as etapas de extração (3.3.1) e análise dos dados (3.3.2).

3.1 Delineamento da Pesquisa

O presente estudo adota como método a Revisão Multivocal da Literatura (do inglês, Multivocal Literature Review - MLR). A escolha por esta abordagem justifica-se pela natureza emergente do fenômeno investigado. Discussões relevantes e evidências práticas sobre o tema são encontradas tanto na literatura científica tradicional quanto em fontes não acadêmicas, tornando a MLR particularmente adequada (GAROUSI et al., 2019).

A Revisão Multivocal combina o rigor da Revisão Sistemática da Literatura (do inglês, Systematic Literature Review - SLR), aplicada aos artigos científicos revisados por pares, com a análise da chamada literatura cinzenta, que abrange fontes como publicações em blogs, fóruns online, notícias e documentação técnica (GAROUSI et al., 2019). A pertinência da inclusão da literatura cinzenta é corroborada pelas diretrizes propostas por Garousi et al. (2019), cujo checklist de validação confirmou a adequação dessa abordagem para o escopo da pesquisa, conforme detalhado na Tabela 2.

Tabela 2 – Perguntas para decidir se a literatura cinza deve ser incluída na revisão

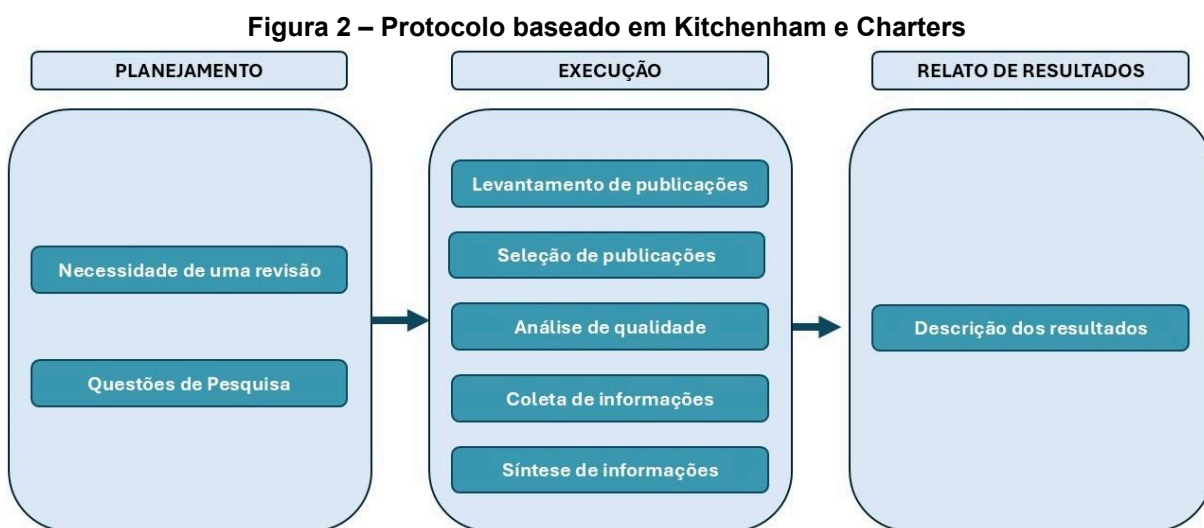
Pergunta	Resposta
O assunto é "complexo" e não pode ser resolvido considerando apenas a literatura formal?	Sim
Há falta de volume ou qualidade de evidências, ou falta de consenso sobre a mensuração dos resultados na literatura formal?	Sim
As informações contextuais são importantes para o tema em estudo?	Sim

O objetivo é validar ou corroborar resultados científicos com experiências práticas?	Sim
O objetivo é desafiar pressupostos ou refutar resultados da prática usando pesquisa acadêmica ou vice-versa?	Sim
Uma síntese de insights e evidências da comunidade industrial e acadêmica seria útil para uma ou ambas as comunidades?	Sim
Há um grande volume de fontes indicando alto interesse de profissionais no tópico?	Sim

Fonte: Adaptado de Garousi et al. (2019)

Para garantir a sistematização e a replicabilidade do processo, o protocolo metodológico foi estruturado com base em dois referenciais principais. A condução da SLR seguiu o modelo proposto por Kitchenham e Charters (2007), uma referência consolidada na Engenharia de Software e áreas correlatas. Por sua vez, a análise da literatura cinzenta foi orientada pelas diretrizes de Garousi et al. (2019).

Assim, o processo metodológico deste estudo segmenta a pesquisa em três fases principais: Planejamento, execução e relato de resultados. Um esquema visual de todas as etapas seguidas pode ser consultado na Figura 2, que detalha o fluxo da pesquisa.



Fonte: Adaptado de Botelho (2025)

Na fase de planejamento, o primeiro passo consiste em definir a necessidade da pesquisa, o que envolve avaliar a relevância do tópico e verificar as áreas ainda não exploradas pela literatura. A partir dessa necessidade, são formuladas as

questões norteadoras que direcionarão a seleção e a análise dos estudos primários (BOTELHO, 2025).

A fase de execução é a mais operacional e envolve uma sequência estruturada de cinco procedimentos: (i) o levantamento inicial das publicações em bases de dados relevantes; (ii) a filtragem e seleção desses trabalhos conforme critérios de inclusão e exclusão; (iii) a análise da qualidade metodológica dos estudos que foram selecionados; (iv) a coleta das informações pertinentes de cada estudo; e (v) a síntese das informações coletadas para responder às questões de pesquisa (BOTELHO, 2025).

Por fim, na fase de relato, os resultados obtidos são organizados e descritos, visando uma comunicação transparente e acessível das conclusões (BOTELHO, 2025).

A metodologia selecionada contribui para a definição de um processo estruturado para a identificação, seleção e análise de publicações relevantes, tanto na literatura branca quanto na literatura cinzenta. Seu rigor metodológico possibilita a coleta e síntese de informações de forma imparcial, assegurando que o mapeamento contemple os estudos mais significativos e confiáveis dentro do escopo definido. Além disso, a aplicação de critérios de qualidade na seleção dos estudos permite uma interpretação mais precisa dos achados, evitando distorções que possam comprometer a validade das conclusões (KITCHENHAM; CHARTERS, 2007). Esse protocolo visa não apenas mapear os desafios e benefícios do uso de LLMs para fins terapêuticos e apoio socioemocional, mas também identificar práticas emergentes, fornecendo subsídios para pesquisas futuras.

3.2 Planejamento e Condução da Revisão

3.2.1 Questões de Pesquisa

Para delimitar o escopo e orientar a condução deste estudo, foram elaboradas cinco questões de pesquisa (QPs) centrais. O conjunto de questões foi estrategicamente desenvolvido para investigar o fenômeno sob múltiplas perspectivas, abrangendo as motivações que levam à adoção dos LLMs, a dinâmica da interação, os benefícios subjetivos, os riscos emergentes e as características do perfil dos usuários. Essa abordagem estruturada permite a construção de um

panorama abrangente sobre as causas, os processos e as consequências associadas ao uso de LLMs como ferramenta de apoio socioemocional. A Tabela 3 detalha cada uma das questões de pesquisa e suas respectivas motivações.

Tabela 3 – Questões de pesquisa e motivações

#	Questão de Pesquisa	Motivação
QP1	Quais são as motivações que levam usuários a adotar LLMs para fins terapêuticos e apoio socioemocional?	Mapear os gatilhos e necessidades que impulsionam o uso.
QP2	Como ocorre a dinâmica de interação para fins terapêuticos e apoio socioemocional?	Analisar a natureza da conversação
QP3	Quais benefícios subjetivos são relatados para fins terapêuticos e apoio socioemocional?	Identificar os impactos positivos percebidos pelos usuários
QP4	Quais riscos emergem desse uso não supervisionado para fins terapêuticos e apoio socioemocional?	Mapear os perigos e desvantagens
QP5	Qual é o perfil dos usuários que recorrem a LLMs para fins terapêuticos e apoio socioemocional?	Caracterizar os usuários que buscam esse apoio

Fonte: A autora (2025).

3.2.2 Estratégia de Busca e Fontes de Dados

A estratégia de busca foi desenhada para abranger tanto a literatura científica (branca) quanto a não acadêmica (cinzenta), garantindo uma visão completa do fenômeno. As seções a seguir detalham os procedimentos para cada tipo de fonte.

3.2.2.1 Literatura Acadêmica

A busca por estudos científicos relevantes foi conduzida com base em uma *string* de busca desenvolvida a partir dos conceitos centrais das questões de pesquisa. A construção da *string* foi um processo iterativo, envolvendo testes preliminares para refinar os termos e maximizar a precisão dos resultados. Durante essa fase, observou-se, por exemplo, que termos como *therapy* (terapia), quando não associados a um contexto de uso informal, resultavam em um volume

extremamente elevado de falsos positivos relacionados a tratamentos clínicos biológicos, que fogem ao escopo deste trabalho.

Essa etapa de refinamento permitiu consolidar três pilares conceituais na busca: (1) a tecnologia (LLMs), (2) a aplicação (saúde mental) e (3) o contexto (uso espontâneo ou não supervisionado). A string de busca genérica final adotada foi:

(AI OR "artificial intelligence" OR LLM OR "large language model*") AND ("mental health" OR psychiatry OR therap*) AND (spontaneous OR unsupervised OR casual OR informal)*

A *string* genérica foi adaptada à sintaxe específica de cada uma das bases de dados consultadas, sempre preservando a integridade dos três pilares conceituais definidos.

As buscas foram executadas em bases de dados eletrônicas de alto impacto, selecionadas por seu reconhecimento internacional, rigoroso processo de revisão por pares e ampla cobertura de periódicos científicos e anais de conferências relevantes ao tema. A lista completa das fontes consultadas é apresentada na Tabela 4.

Tabela 4 – Bibliotecas utilizadas

Nome da biblioteca	Endereço de acesso
ACM Digital Library	https://dl.acm.org/
Emerald Insight	https://emerald.com/insight/
IEEE Xplore	https://ieeexplore.ieee.org/
Science Direct	https://sciencedirect.com/
Scopus	https://scopus.com/

Fonte: A autora (2025).

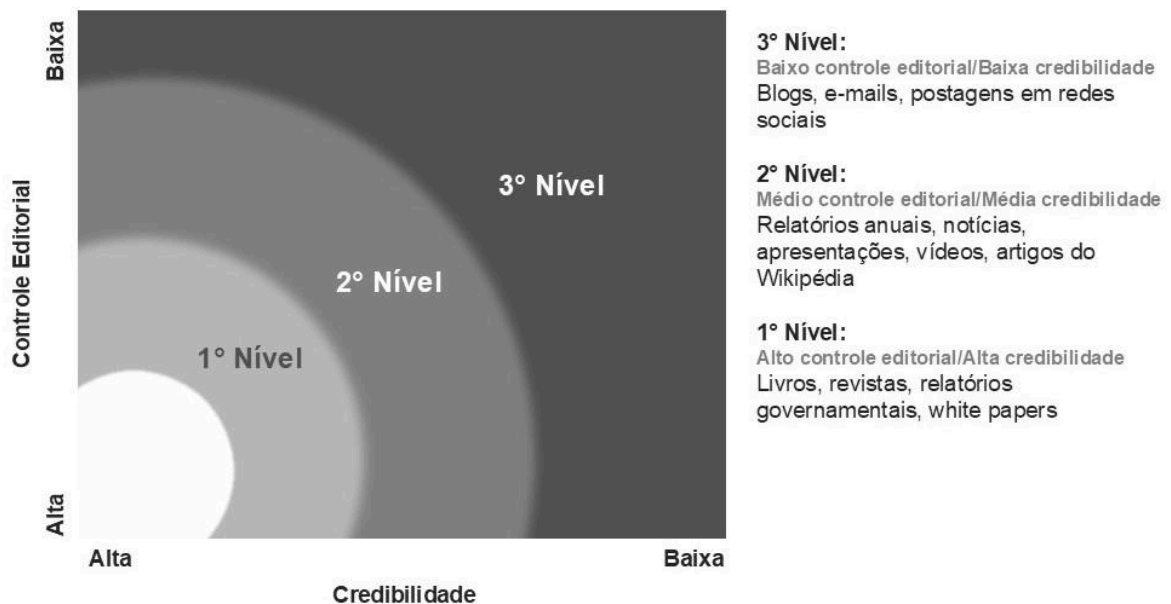
A coleta dos dados foi realizada em todas as bibliotecas selecionadas no dia 16 de junho de 2025, por meio dos mecanismos de exportação de cada uma. Os resultados foram agrupados em uma planilha eletrônica (Apêndice A) que foi utilizada durante todo o estudo.

3.2.2.2 Literatura Cinza

Como mencionado anteriormente, esta revisão incorpora a literatura cinzenta para capturar uma perspectiva mais ampla do tema, incluindo a voz de praticantes e usuários. A inclusão dessas fontes é fundamental em um campo orientado à prática como a Engenharia de Software, pois permite a identificação de tópicos emergentes (GAROUSI et al., 2019).

Para classificar e avaliar as fontes cinzentas, este estudo adotou o modelo de "tons de cinza" discutido em Garousi et al. (2019). Esse modelo organiza a literatura cinzenta em diferentes níveis (*tiers*), com base no grau de controle editorial da fonte e na credibilidade e conhecimento do produtor do conteúdo. A classificação (Figura 3) varia desde um primeiro nível, com alta credibilidade, até um terceiro nível, com baixa credibilidade e controle, caracterizado por conteúdos como blogs e tweets (ADAMS; SMART; HUFF, 2017).

Figura 3 – Tons de Literatura Cinza



Fonte: Adaptado de Garousi et al. (2019)

No escopo desta pesquisa, a seleção de fontes cinzentas foi direcionada para capturar o discurso público e as experiências autênticas dos usuários. Por essa razão, fontes de 1º Nível, que possuem alta credibilidade mas geralmente abordam o tema de uma perspectiva mais formal e menos focada no uso cotidiano, foram consideradas fora do escopo.

A escolha, portanto, concentrou-se em fontes de "tom médio" (2º Nível), que possuem controle e credibilidade moderados, e fontes de "tom escuro" (3º Nível), que se caracterizam por refletirem ideias, conceitos e pensamentos espontâneos. Essa abordagem permitiu capturar tanto o discurso estruturado da indústria midiática sobre o fenômeno quanto às experiências informais e diretas dos usuários.

As fontes escolhidas e seus respectivos níveis na classificação estão resumidos na Tabela 5.

Tabela 5 – Fontes utilizadas relacionadas aos seus respectivos níveis

Fonte	Nível	Controle/Credibilidade
Jornais	2º Nível	Moderado
Reddit	3º Nível	Baixo
TikTok	3º Nível	Baixo

Fonte: A autora (2025).

3.2.2.2.1 Jornais

Dentro da literatura cinzenta, optou-se pela inclusão de fontes jornalísticas para representarem o 2º Nível. Artigos de jornais e portais de notícias oferecem uma perspectiva valiosa, pois são produzidos por profissionais, passam por um processo de controle editorial e refletem como o tema é apresentado e discutido para o público amplo. Foram selecionados veículos de imprensa de renome e alcance internacional, conhecidos por seus elevados padrões de reportagem e pela cobertura aprofundada de temas relacionados à tecnologia e sociedade, conforme listado na Tabela 6.

Tabela 6 – Jornais utilizados

Nome do jornal	Endereço de acesso
BBC	https://bbc.com/
The Guardian	https://theguardian.com/
The New York Times	https://nytimes.com/

Fonte: A autora (2025).

A busca nas plataformas foi realizada no dia 8 de julho de 2025, utilizando-se em cada portal a *string* de busca ‘*ai therapy mental health chatbot*’. O processo de coleta e seleção foi dividido em duas etapas. Na primeira, os resultados foram

submetidos a uma triagem inicial baseada na análise de seus títulos e subtítulos. Os artigos considerados pertinentes foram catalogados manualmente em uma planilha (Apêndice A), registrando-se os metadados de autor, título, subtítulo, ano de publicação, fonte e endereço de acesso.

Na segunda etapa, procedeu-se à tentativa de coleta do texto integral de cada artigo pré-selecionado. Nesta fase, todos os resultados provenientes do The New York Times foram excluídos do conjunto final, uma vez que o acesso ao conteúdo completo estava restrito por um sistema de *paywall*, o que inviabilizou a análise.

3.2.2.2 Reddit

Como uma das fontes para a análise de publicações de 3º Nível, que se destacam pela espontaneidade e por capturar discussões orgânicas, foi selecionada a plataforma Reddit. Dentre as plataformas de mídia social baseadas em textos existentes, o Reddit foi escolhido devido à sua grande e diversificada base de usuários, bem como às suas discussões ativas sobre as mais recentes tecnologias (LI et al., 2025).

Ele se destaca como uma plataforma única e valiosa devido à sua estrutura baseada em comunidades temáticas, conhecidas como *subreddits*, e ao seu foco em debates aprofundados (CHEN; TOMBLON, 2021; ORYNGOZHA; SHAMOI; IGALI, 2024). Essa segmentação permite que os usuários participem de conversas mais especializadas e engajadas, com menos ruído do que em plataformas como o Twitter, onde as postagens se misturam em um fluxo contínuo e muitas vezes caótico (CHEN; TOMBLON, 2021).

Além disso, o Reddit oferece um grau maior de anonimato em comparação com redes sociais como Twitter e Facebook. Muitos usuários utilizam pseudônimos e não vinculam suas identidades reais aos perfis, o que reduz a autocensura e favorece a expressão de opiniões mais francas, incluindo perspectivas críticas ou controversas (CHEN; TOMBLON, 2021).

Sendo assim, a estratégia de coleta foi dividida em dois escopos distintos — um focado em tecnologia e outro em saúde mental — a fim de contrastar as perspectivas de uma comunidade já adepta ao uso de IA com as de comunidades cujo foco primário não é a tecnologia. Para o escopo tecnológico, foi selecionado exclusivamente o *subreddit* r/ChatGPT, que, com mais de 11 milhões de membros,

figura entre os 1% maiores da plataforma. A escolha se justifica pelo seu altíssimo nível de engajamento, que frequentemente supera o de outras comunidades sobre o tema. Para o escopo de saúde mental, foram selecionados os *subreddits* de maior popularidade na área: r/mentalhealth, r/therapy e r/selfhelp. Os detalhes de cada comunidade estão na Tabela 7.

Tabela 7 – Subreddits utilizados

Nome do <i>subreddit</i>	Endereço de acesso	Nº de membros
r/ChatGPT	https://reddit.com/r/ChatGPT/	11 milhões
r/mentalhealth	https://reddit.com/r/mentalhealth/	552 mil
r/therapy	https://reddit.com/r/therapy/	144 mil
r/selfhelp	https://reddit.com/r/selfhelp/	209 mil

Fonte: A autora (2025).

A coleta de dados foi realizada no dia 10 de julho de 2025, de forma automatizada por meio da API oficial do Reddit, com o auxílio da biblioteca PRAW (*Python Reddit API Wrapper*). Foi desenvolvido um *script* em Python (Apêndice B) para extrair as 50 postagens mais relevantes de cada escopo. No r/ChatGPT, a busca foi filtrada pelo termo *‘therapy’*. Nos *subreddits* de saúde mental, o termo utilizado foi *‘ChatGPT’*, uma escolha deliberada para manter a consistência com o escopo tecnológico, que já era focado nesta ferramenta específica.

Para cada postagem, foram coletados os seguintes dados: título, usuário, URL, data de publicação, conteúdo textual, pontuação (votos), número de comentários e o conteúdo dos dez comentários mais relevantes. Os dados brutos, gerados em um arquivo de texto, foram posteriormente estruturados em uma planilha para a etapa de análise.

3.2.2.2.3 TikTok

A plataforma TikTok também foi selecionada como uma fonte adicional de 3º Nível, para analisar o discurso audiovisual sobre o tema. O TikTok se consolidou como uma das plataformas mais eficientes para alcançar audiências massivas e diversificadas, superando muitas redes tradicionais em alcance e engajamento (HAENSCH, 2025; DIXON, 2023). Seu algoritmo inteligente e formato dinâmico

permitem que conteúdos atinjam não apenas milhões de usuários, mas também públicos segmentados de maneira orgânica e viral (ANDERSON, 2020 apud HAENSCH, 2025).

Isso possibilita campanhas de conscientização, tendências culturais e até debates sociais ganharem escala rapidamente, algo mais difícil em fóruns de discussão tradicionais. Além disso, o TikTok não depende apenas de seguidores – o algoritmo prioriza a descoberta, permitindo que vídeos atinjam pessoas com interesses similares, mesmo sem uma base de seguidores estabelecida. Essa característica única faz com que mensagens, marcas ou movimentos alcancem um público muito maior do que em plataformas onde o alcance é limitado por conexões pré-existentes (ANDERSON, 2020 apud HAENSCH, 2025).

Dada a ausência de uma API de acesso público, a coleta de dados foi realizada manualmente, uma limitação que influenciou o desenho da amostragem. Para mitigar o viés de personalização do algoritmo da plataforma, foi criada uma conta de usuário nova, neutra, sem histórico de interações ou preferências de conteúdo.

A busca foi realizada em 08 de julho de 2025, utilizando-se a *string* de busca ‘*AI therapy*’ na ferramenta de pesquisa nativa do TikTok. Os vídeos resultantes foram analisados sequencialmente e incluídos na amostra caso atendessem aos critérios de inclusão (detalhados na Seção 3.2.3). Foi definida uma meta de amostragem de 50 vídeos e os três comentários mais relevantes de cada um, totalizando 150 comentários. Este tamanho de amostra foi estabelecido para garantir a viabilidade da análise manual aprofundada, considerando a impossibilidade de automatizar a coleta devido às restrições ao acesso da API da plataforma.

Para cada vídeo selecionado, foram extraídos dados objetivos e realizada uma análise de conteúdo. Os dados objetivos registrados em planilha incluíram: o endereço de acesso, usuário, data de publicação, descrição, menção a plataformas específicas de LLMs e métricas de engajamento (número de curtidas, comentários e salvamentos).

3.2.3 Critérios de Inclusão e Exclusão

Para assegurar a qualidade, pertinência e atualidade do material analisado, foram definidos critérios de inclusão e exclusão, que guiaram o processo de seleção

em todas as fontes de dados. Neste estudo, o termo "publicação" é empregado de forma abrangente para referenciar os diversos formatos de evidência coletados: artigos acadêmicos, textos jornalísticos, postagens no Reddit e vídeos do TikTok.

O critério temporal foi um dos filtros primários, restringindo a busca a publicações do período entre janeiro de 2021 e julho de 2025. Este recorte justifica-se pela rápida e recente evolução dos LLMs, garantindo que a análise se concentre em discussões alinhadas com a maturidade tecnológica atual e as práticas emergentes. Adicionalmente, foram aplicados critérios gerais de acessibilidade (disponibilidade na íntegra e acesso gratuito) e idioma (inglês, com exceção do português, admitido para vídeos do TikTok). Publicações duplicadas foram sistematicamente identificadas e removidas.

Critérios específicos também foram aplicados de acordo com a natureza de cada fonte. Para a literatura acadêmica, a seleção foi restrita a artigos primários com extensão superior a cinco páginas, buscando garantir a profundidade do material. Para a literatura cinzenta, os filtros visaram a qualidade e a relevância do conteúdo: foram incluídos apenas textos e postagens de natureza informativa e não promocional. Para as mídias sociais, foi estabelecido um critério quantitativo de engajamento mínimo, considerando-se apenas postagens cuja soma de curtidas e comentários fosse superior a 10.

As Tabelas 8 e 9 consolidam, respectivamente, todos os critérios de inclusão e exclusão que orientaram o processo de triagem das publicações.

Tabela 8 – Critérios de inclusão

#	Critérios de Inclusão
CI1	Publicações dentro do período definido (2021-2025)
CI2	Publicações disponíveis na íntegra e de acesso livre/gratuito
CI3	Publicações nos idiomas inglês, ou português (especificamente para vídeos do TikTok)
CI4	Publicações únicas (não duplicadas)
CI5	Publicações que estejam diretamente relacionadas ao tema de uso de LLMs para fins terapêuticos e apoio socioemocional
CI6	Artigos com mais de 5 páginas
CI7	Artigos primários

CI8	Postagens ou textos jornalísticos sem viés publicitário
CI9	Postagens com engajamento significativo (soma de curtidas e comentários superior a 10).

Fonte: A autora (2025).

Tabela 9 – Critérios de exclusão

#	Critérios de Exclusão
CE1	Publicações fora do período definido (2021-2025)
CE2	Publicações não disponíveis na íntegra ou de acesso restrito/pago
CE3	Publicações que não estejam em inglês (ou português, no caso de vídeos do TikTok)
CE4	Publicações duplicadas
CE5	Publicações que não estejam diretamente relacionadas ao tema de uso de LLMs para fins terapêuticos e apoio socioemocional
CE6	Artigos com 5 páginas ou menos
CE7	Artigos secundários
CE8	Postagens ou textos jornalísticos que possuam indícios de serem conteúdo publicitário
CE9	Postagens com baixo engajamento (soma de curtidas e comentários igual ou inferior a 10).

Fonte: A autora (2025).

3.2.4 Processo de Seleção e Avaliação da Qualidade

Após a fase de busca, foi executado um processo de seleção e avaliação de qualidade em múltiplas etapas para garantir a confiabilidade e a robustez das evidências coletadas. Reconhecendo a natureza heterogênea das fontes, que abrangem desde artigos científicos revisados por pares até conteúdo espontâneo de redes sociais, foi adotada uma abordagem de avaliação diferenciada, com critérios específicos para a literatura acadêmica e para a literatura cinzenta.

3.2.4.1 Literatura Acadêmica

A busca inicial nas bases de dados resultou em 586 artigos. Na primeira etapa de filtragem, realizada a partir da leitura de títulos e resumos, 515 publicações foram excluídas por não atenderem aos critérios de inclusão (ex: período, idioma, escopo temático). Adicionalmente, 38 artigos foram separados por terem temas tangenciais, sendo mantidos apenas como referência bibliográfica. Ao final desta fase, 31 artigos foram considerados elegíveis para a análise completa.

Os 31 artigos restantes foram submetidos a uma avaliação de qualidade aprofundada, baseada na leitura integral dos textos. Para aferir o rigor metodológico e a relevância de cada estudo, foram aplicados cinco critérios de qualidade (Tabela 10), pontuados em uma escala de três níveis: Sim (1 ponto), Parcialmente (0,5 ponto) e Não (0 pontos).

Tabela 10 – Critérios de avaliação de qualidade

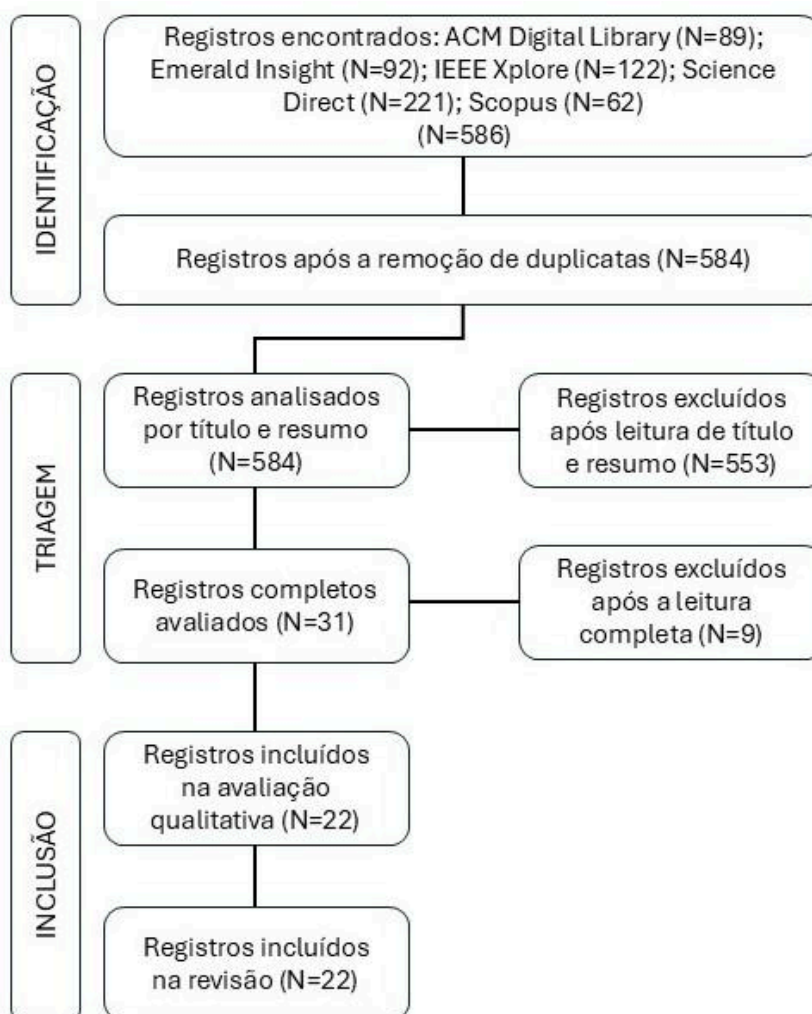
#	Pergunta	Resposta
CQ1	O artigo está diretamente relacionado ao uso de LLMs para fins terapêuticos e apoio socioemocional?	S=1, N=0, P=0.5
CQ2	A metodologia utilizada na pesquisa está claramente descrita e é compreensível?	S=1, N=0, P=0.5
CQ3	O artigo apresenta e discute os resultados obtidos?	S=1, N=0, P=0.5
CQ4	O artigo faz uma análise crítica sobre o uso de LLMs para os fins estudados?	S=1, N=0, P=0.5
CQ5	O artigo menciona os desafios, as limitações ou as possíveis ameaças relacionadas ao tema?	S=1, N=0, P=0.5

Fonte: A autora (2025).

Para a inclusão final, estabeleceu-se uma nota de corte de 2,5 (de um total de 5 pontos), com o critério CQ1 sendo obrigatório (necessário obter nota 1). A aplicação desses critérios, apoiada pela ferramenta NotebookLM PRO (<https://notebooklm.google.com/>) e validada integralmente pela pesquisadora, resultou na seleção final de 22 artigos (*prompt* disponível no Apêndice C).

O processo de seleção dos estudos acadêmicos seguiu um fluxo sistemático, ilustrado na Figura 4, em conformidade com o modelo PRISMA (MOHER et al., 2009).

Figura 4 – Fluxo do processo de seleção dos artigos



Fonte: Elaborado pela autora de acordo com o modelo de Moher et al. (2009)

A Tabela 11 apresenta a síntese dos 22 estudos que compõem o portfólio final desta revisão. Para garantir a transparência e a reprodutibilidade, o registro completo da análise, incluindo os dados brutos e os critérios aplicados, está publicamente disponível na planilha eletrônica que foi utilizada durante todo o estudo (Apêndice A).

Tabela 11 – Síntese dos artigos analisados

Nº	Ano	Base	Título
A1	2023	Science Direct	Can we use ChatGPT for Mental Health and Substance Use Education? Examining Its Quality and Potential Harms
A2	2023	IEEE	Focusing on Needs: A Chatbot-Based Emotion Regulation Tool for Adolescents

A3	2023	ACM	Living Memories: AI-Generated Characters as Digital Mementos
A4	2024	ACM	"It's the only thing I can trust": Envisioning Large Language Model Use by Autistic Workers for Communication Assistance
A5	2024	Science Direct	A Reinforcement Learning Approach for Intelligent Conversational Chatbot For Enhancing Mental Health Therapy
A6	2024	ACM	ARIEL: Brain-Computer Interfaces meet Large Language Models for Emotional Support Conversation
A7	2024	ACM	Can Large Language Models Be Good Companions? An LLM-Based Eyewear System with Conversational Common Ground
A8	2024	ACM	Chatbots With Attitude: Enhancing Chatbot Interactions Through Dynamic Personality Infusion
A9	2024	ACM	Evaluating the Experience of LGBTQ+ People Using Large Language Model Based Chatbots for Mental Health Support
A10	2024	ACM	From Information Seeking to Empowerment: Using Large Language Model Chatbot in Supporting Wheelchair Life in Low Resource Settings
A11	2024	Emerald	The effect of ChatGPT on EFL students' social and emotional learning
A12	2024	ACM	This Chatbot Would Never...: Perceived Moral Agency of Mental Health Chatbots
A13	2025	Science Direct	"This is human intelligence debugging artificial intelligence": Examining how people prompt GPT in seeking mental health support
A14	2025	Science Direct	A Comparison of Responses from Human Therapists and Large Language Model-Based Chatbots to Assess Therapeutic Communication: Mixed Methods Study
A15	2025	Science Direct	Caregiving Artificial Intelligence Chatbot for Older Adults and Their Preferences, Well-Being, and Social Connectivity: Mixed-Method Study
A16	2025	ACM	Customizing Emotional Support: How Do

			Individuals Construct and Interact With LLM-Powered Chatbots
A17	2025	Science Direct	Development and Evaluation of a Mental Health Chatbot Using ChatGPT 4.0: Mixed Methods User Experience Study With Korean Users
A18	2025	Science Direct	Expert and Interdisciplinary Analysis of AI-Driven Chatbots for Mental Health Support: Mixed Methods Study
A19	2025	ACM	ExploreSelf: Fostering User-driven Exploration and Reflection on Personal Challenges with Adaptive Guidance by Large Language Models
A20	2025	Science Direct	Love, marriage, pregnancy: Commitment processes in romantic relationships with AI chatbots
A21	2025	ACM	Private Yet Social: How LLM Chatbots Support and Challenge Eating Disorder Recovery
A22	2025	ACM	User-Driven Value Alignment: Understanding Users' Perceptions and Strategies for Addressing Biased and Discriminatory Statements in AI Companions

Fonte: A autora (2025).

Embora o processo tenha sido conduzido por um único revisor, estratégias foram adotadas para mitigar vieses, como a definição prévia e aplicação estrita de critérios objetivos de inclusão, exclusão e qualidade. Recomenda-se, para trabalhos futuros, a inclusão de múltiplos revisores para fortalecer ainda mais a robustez metodológica.

3.2.4.2 Literatura Cinza

A avaliação da literatura cinzenta exigiu uma abordagem adaptada, pois seu valor reside na captura do discurso público e de experiências autênticas, e não na replicabilidade metodológica de um estudo formal. Diferentemente do processo para a literatura acadêmica, a triagem ocorreu de forma concomitante à coleta manual dos dados. Os critérios de qualidade foram qualitativos e específicos para cada fonte.

Para as fontes jornalísticas, o principal indicador de qualidade foi a reputação e a credibilidade editorial do veículo. Já para o conteúdo gerado por usuários no Reddit e TikTok, a avaliação focou na riqueza, substância e autenticidade do conteúdo. Durante esta etapa, foram excluídos conteúdos com fortes indícios de serem promocionais, como vídeos que pareciam seguir um roteiro para divulgar uma ferramenta específica, comprometendo a legitimidade dos relatos. Essa abordagem multifacetada assegurou que cada tipo de fonte fosse julgado por padrões apropriados à sua natureza.

3.3 Extração e Análise dos Dados

3.3.1 Extração dos dados

A extração de dados foi um processo sistemático guiado pelas cinco questões de pesquisa (QPs), com o objetivo de coletar as evidências pertinentes de cada publicação selecionada. Em conformidade com as diretrizes de Kitchenham e Charters (2007), a abordagem foi adaptada à natureza de cada tipo de fonte — artigos acadêmicos, textos jornalísticos e conteúdo de redes sociais —, e todas as informações foram centralizadas em uma planilha eletrônica (Apêndice A) para posterior análise.

A capacidade de cada tipo de publicação em responder às questões de pesquisa variou. A Tabela 12 apresenta uma matriz de cobertura que ilustra o nível de profundidade com que cada fonte abordou as QPs, indicando se foram tratadas de forma direta (X), parcial (/), ou se não foram abordadas.

Tabela 12 – Questões de pesquisa e tipos de publicação

#	QP1	QP2	QP3	QP4	QP5
Artigos	X	X	X	X	X
Jornais	/	X	X	X	X
Reddit	/	X	X	X	/
TikTok		X	X	X	/

Fonte: A autora (2025).

Para a literatura acadêmica, foi empregado um processo de extração detalhado. As evidências relativas a cada questão de pesquisa foram pontuadas em uma escala de três níveis: Sim (1), Parcialmente (0,5) e Não (0). Para sistematizar esta etapa, utilizou-se a ferramenta NotebookLM PRO (com *prompt* detalhado no Apêndice D) como apoio para gerar a pontuação, uma justificativa e os trechos de evidência correspondentes.

A fim de assegurar a consistência e a confiabilidade dos resultados gerados pela ferramenta, implementou-se um protocolo de validação. Justificativas e trechos de evidência foram requisitados em sessões de *chat* distintas para permitir a comparação das pontuações. Nos casos em que foram identificadas discrepâncias, a pesquisadora realizou uma revisão manual do artigo para arbitrar a avaliação final. Este procedimento garantiu a supervisão humana e o rigor em todo o processo de extração.

Para os artigos acessíveis nas fontes jornalísticas, foram extraídas informações detalhadas sobre o seu conteúdo, como casos de uso debatidos, benefícios e riscos mencionados, e informações sobre o perfil dos usuários. Para o Reddit, a análise do conteúdo foi um processo multifacetado. Primeiramente, cada postagem e comentário foi classificado segundo a opinião expressa sobre o uso de LLMs para fins terapêuticos, de acordo com os critérios definidos na Tabela 13.

Tabela 13 – Possíveis opiniões e critérios adotados para classificação

Opinião	Crítérios
Positiva	Usa, menciona benefícios e ignora ou minimiza riscos
Tende a Positiva	Usa, menciona benefícios mas não minimiza riscos
Mista	Menciona benefícios e riscos igualmente
Tende a Negativa	Enxerga utilidade, mas foca nos riscos e limitações
Negativa	Contra o uso, dá foco aos riscos
Desconsiderado	Sem opinião clara, não relacionado ao tema ou gerados por <i>bots</i>

Fonte: A autora (2025).

Adicionalmente, foi realizada uma análise temática nas postagens primárias para mapear o assunto principal, os casos de uso, os benefícios e os riscos citados. Por fim, após uma leitura imersiva dos dados, um processo de codificação indutiva deu origem a seis categorias binárias (Sim/Não) para caracterizar relatos específicos

dos usuários: (1) preferência pela terapia com IA em detrimento da terapia com humanos; (2) uso concomitante de IA e terapia com humanos; (3) relato de experiência negativa prévia com terapia tradicional; (4) menção a um diagnóstico psicológico; (5) relato sobre problemas com vício em substâncias; e (6) menção a sentimentos de solidão ou isolamento social extremo.

No TikTok, a análise de conteúdo qualitativa, por sua vez, buscou identificar: o gênero percebido do criador, a visão geral sobre o uso de LLMs (classificada como Positiva, Mista ou Negativa), a mensagem central do vídeo, casos de uso, além de benefícios e riscos mencionados. Os três comentários de maior relevância de cada vídeo foram transcritos e avaliados quanto ao posicionamento expresso sobre o uso de LLMs para terapia, utilizando a mesma escala e critérios de classificação definidos na Tabela 13.

3.3.2 Método de análise

A etapa final da metodologia consistiu na análise e síntese dos dados extraídos das publicações selecionadas. O método central empregado foi a Análise Temática, escolhida por sua adequação à natureza exploratória do estudo (BRAUN; CLARKE, 2006). Essa abordagem permite a identificação, codificação e interpretação de padrões e temas recorrentes nos dados, respondendo de forma aprofundada às questões de pesquisa (QP1–QP5).

O processo iniciou-se com a categorização das evidências coletadas. Para os artigos acadêmicos e textos jornalísticos, esta etapa foi otimizada com o auxílio da ferramenta Gemini PRO (com *prompt* detalhado no Apêndice E), sob validação da pesquisadora. A escolha por essa abordagem foi viabilizada pela linguagem desses textos, caracterizada pela formalidade e clareza, que reduz as ambiguidades e torna a identificação automatizada de temas e argumentos mais confiáveis.

No entanto, para o conteúdo das redes sociais, a categorização foi realizada manualmente, garantindo uma análise sensível às nuances da linguagem informal. Relatos em redes sociais são frequentemente permeados por ironia, sarcasmo, gírias e subtextos emocionais que uma ferramenta automatizada poderia interpretar de maneira equivocada.

Após a categorização, os dados foram quantificados para permitir uma caracterização detalhada do corpus de publicações. Com a conclusão das etapas de

busca, seleção e avaliação das publicações, o mapeamento está preparado para fornecer uma análise aprofundada sobre as motivações, dinâmicas de interação, benefícios, riscos e o perfil dos usuários que recorrem a LLMs para fins terapêuticos e apoio socioemocional.

O próximo capítulo apresenta os resultados desta análise, detalhando os temas que emergiram dos dados e discutindo suas implicações no contexto acadêmico, social e para o desenvolvimento de tecnologias.

Para garantir a clareza e a rastreabilidade na apresentação dos resultados, foi adotado um sistema de codificação para todas as fontes primárias incluídas nesta revisão. Cada tipo de publicação recebeu um prefixo e um identificador numérico único, permitindo uma citação concisa e precisa no corpo do texto. O esquema completo de identificadores é detalhado na Tabela 14.

Tabela 14 – Esquema de identificador por tipo de publicação

Tipo	Referência	Estrutura	Exemplo
Artigos Acadêmicos	Artigo	Ax	A1, A2...
Textos Jornalísticos	Notícia	Nx	N1, N2...
Postagens no Reddit	Reddit	Rx	R1, R2...
Comentários no Reddit	Reddit-Comentário	RxCx	R1C1, R1C2...
Vídeos no TikTok	TikTok	Tx	T1, T2...
Comentários no TikTok	TikTok-Comentário	TxCx	T1C1, T2C2...

Fonte: A autora (2025).

Essas codificações são utilizadas nas citações ao longo do capítulo de Resultados para referenciar as evidências específicas que suportam cada achado. As referências completas para as publicações principais (artigos, notícias, postagens e vídeos) que correspondem a estes códigos estão detalhadas na seção de Referências, e a listagem das fontes que compõem cada categoria estão divididas em tabelas no Apêndice F. Devido ao grande volume, a lista com os links e a transcrição dos comentários individuais citados encontra-se no Apêndice A.

4 RESULTADOS

Este capítulo apresenta os resultados obtidos a partir da Revisão Multivocal da Literatura sobre o uso espontâneo de Modelos de Linguagem de Larga Escala (LLMs) como apoio socioemocional. Os achados aqui descritos representam a síntese das informações extraídas das 135 publicações selecionadas, sendo 22 artigos acadêmicos e 113 fontes da literatura cinzenta. Por meio da Análise Temática, foram identificados padrões recorrentes que respondem diretamente às perguntas que nortearam este estudo.

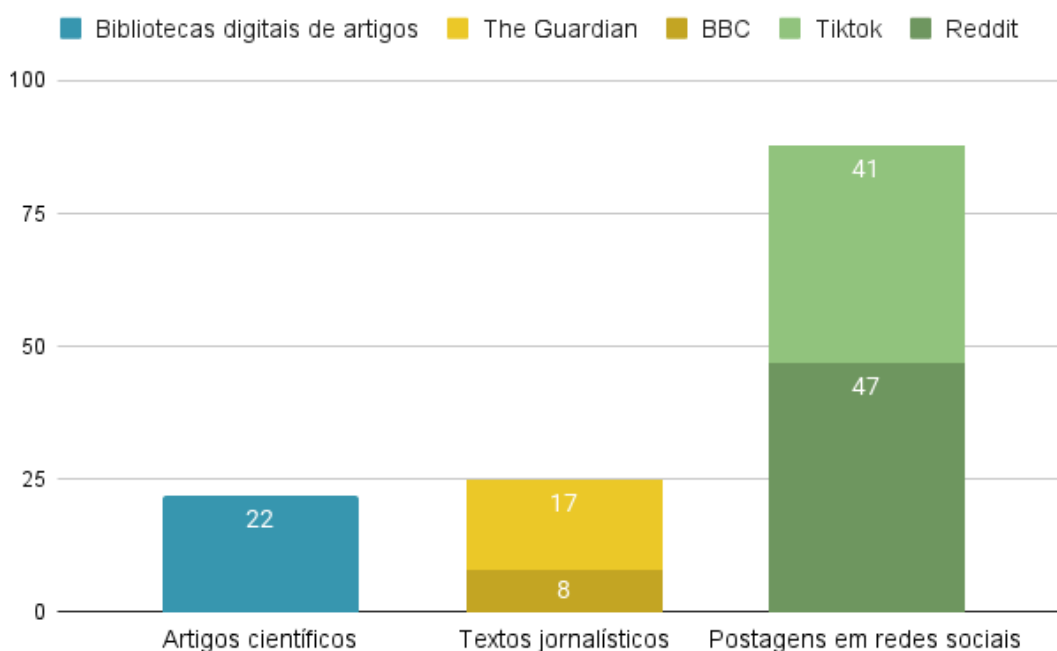
A apresentação dos resultados está organizada em quatro seções principais, que abordam de forma integrada a caracterização das publicações selecionadas e as respostas às questões de pesquisa. A primeira seção, 4.1, oferece uma caracterização descritiva do conjunto de fontes selecionadas, analisando a distribuição temporal e as origens dos dados. Em seguida, a seção 4.2 explora as motivações (QP1) e os padrões de uso (QP2) de LLMs para fins socioemocionais, além de uma análise das abordagens do assunto presentes nas redes sociais. A seção 4.3 analisa a dualidade dessa interação, detalhando tanto os benefícios (QP3) subjetivos quanto os riscos e desafios (QP4) emergentes, seguido de um balanço da visão dos usuários. Por fim, a seção 4.4 busca delinear o perfil dos usuários (QP5) que recorrem a essas tecnologias, construindo um quadro geral com base nas informações disponíveis em cada fonte estudada.

4.1 Caracterização das Publicações Selecionadas

O conjunto de dados final é composto por 135 publicações de naturezas distintas. O Gráfico 1 ilustra a distribuição geral, evidenciando a predominância de postagens em redes sociais (88), seguidas por textos jornalísticos (25) e artigos científicos (22), o que reflete a natureza emergente e multivocal do tema.

Gráfico 1 – Gráfico da quantidade de publicações, por tipo

Quantidade de publicações analisadas, por tipo

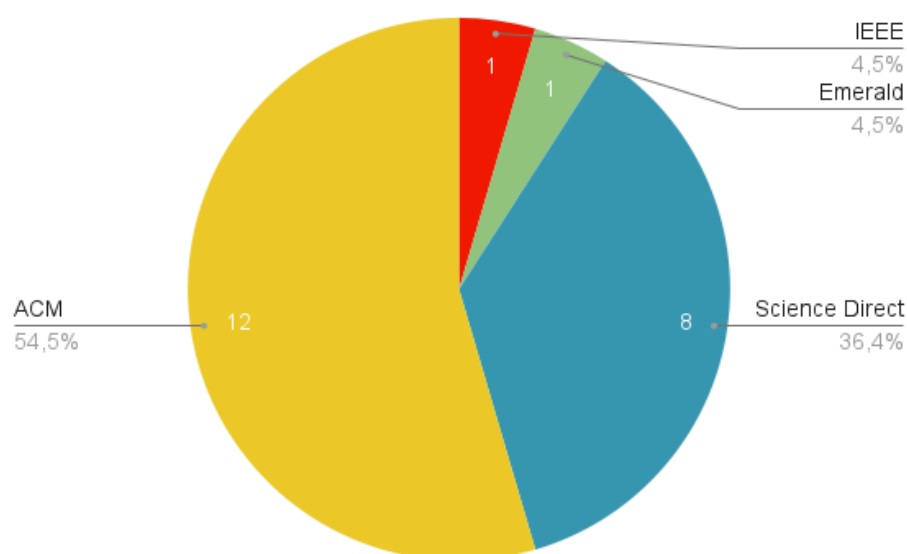


Fonte: A autora (2025).

A análise da proveniência de cada tipo de publicação revela fontes de alta relevância em seus respectivos campos. No que tange aos 22 artigos científicos (Gráfico 2), observa-se uma alta concentração na ACM Digital Library (54,5%) e na ScienceDirect (36,4%). Na esfera jornalística (Gráfico 3), o jornal The Guardian (68,0%) foi a principal fonte entre as 25 reportagens.

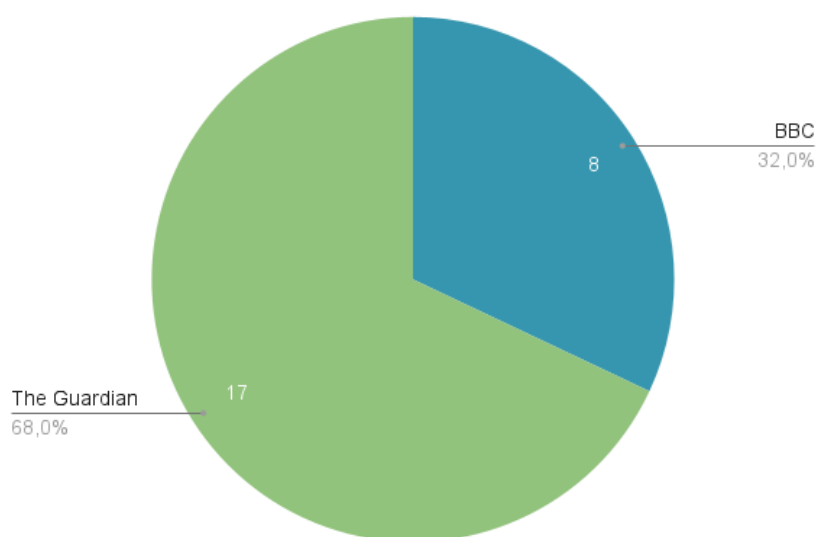
Gráfico 2 – Gráfico da quantidade de artigos por biblioteca digital

Artigos por biblioteca digital



Fonte: A autora (2025).

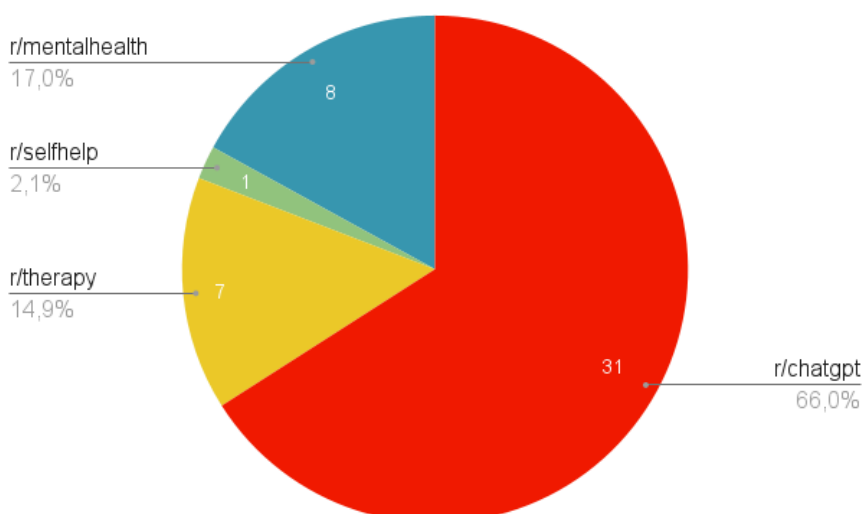
Gráfico 3 – Gráfico da quantidade de reportagens por jornal
Reportagens por jornal



Fonte: A autora (2025).

A origem das 47 postagens coletadas no Reddit, detalhada no Gráfico 4, revela uma distribuição peculiar da discussão entre os escopos de tecnologia e saúde mental. Apesar de ambos escopos terem coletado 50 postagens cada, após as filtragens o *subreddit* focado em tecnologia, *r/ChatGPT*, emergiu como a principal arena dessas discussões, respondendo por 66% da amostra total. Esse predomínio é particularmente significativo quando contrastado com a recepção do tema nos espaços dedicados à saúde mental: *r/mentalhealth* contribuiu com 8 postagens (17%), *r/therapy* com 7 (14,9%), e *r/selfhelp* com apenas 1 (2,1%).

Gráfico 4 – Gráfico da quantidade de postagens por subreddit
Postagens por subreddit



Fonte: A autora (2025).

Essa distribuição desigual reflete não apenas diferenças no volume de discussão, mas também distintas posturas comunitárias em relação ao tema. Os *subreddits* de saúde mental adotam políticas explícitas e restritivas sobre o assunto: o *r/therapy* mantém um post fixado intitulado "*Our AI Policy*" (Anexo A), que permite discussões sobre ética e impactos do uso de IA na terapia, mas proíbe sua promoção como substituto para terapia tradicional. Já o *r/mentalhealth* inclui em suas regras gerais um aviso categórico (Anexo B), traduzido em parte a seguir: "Proibimos discussões ou recomendações de IA para saúde mental. Redes sociais e aplicativos de IA não são substitutos adequados para consultas com profissionais médicos qualificados".

Essas políticas moderadoras explicam, em parte, a concentração das discussões no *r/ChatGPT*, onde os usuários encontram maior liberdade para compartilhar experiências e opiniões sobre o uso de LLMs para apoio emocional. O contraste entre os ambientes revela uma tensão fundamental: enquanto espaços técnicos funcionam como fóruns abertos para exploração do potencial terapêutico das IAs, comunidades de saúde mental mantêm uma postura cautelosa, muitas vezes restritiva, refletindo preocupações profissionais sobre segurança e eficácia. Essa dicotomia entre entusiasmo tecnológico e cautela clínica emerge como um padrão significativo na paisagem digital analisada.

Um reflexo marcante dessa tensão aparece também na linguagem emocional utilizada pelos usuários. Nos *subreddits* de saúde mental, várias postagens revelam sentimentos de vergonha ou culpa por recorrerem a LLMs para apoio emocional, um tom raro nos fóruns de tecnologia.

[...] Sinceramente, me sinto um pouco culpado por isso. Estou compartilhando todos os meus pensamentos e sentimentos mais profundos com uma IA controlada por uma grande corporação. Estou usando uma máquina para simular a sensação de compreensão verdadeira. Parece errado, mas, no fim das contas, não vejo opção melhor para a minha situação atual. [...] ([R31], tradução própria)

Olá, esta é uma conta secundária porque tenho vergonha disso. Ultimamente, tenho me sentido muito triste e solitário e não sabia com quem conversar sobre meus problemas do passado e minha falta de autoestima. Sempre achei esquisitas as pessoas que usam IA para apoio emocional. Decidi tentar só para ver como seria. Parecia que eu tinha alguém que realmente se importava comigo e não me julgava. Eu nunca choro. Eu chorei. [...] ([R35], tradução própria)

Faço terapia uma vez por semana e não confio em mais ninguém para falar sobre coisas importantes, então uso o ChatGPT. Meu terapeuta diz que não tem problema usar, mas estou em conflito porque as pessoas aqui dizem que é horrível e que não se deve usar, mas eles não oferecem nenhuma alternativa para lidar com a situação [...] ([R46], tradução própria)

Essa autocobrança também aparece em comunidades mais gerais. No r/TrueOffMyChest, comunidade de desabafos com mais de 2 milhões de membros, outro usuário compartilhou um sentimento parecido em uma postagem intitulada ““Estou secretamente usando IA como terapeuta e estou meio envergonhado disso” (SPICYBB01, 2025, tradução própria).

Nos *subreddits* técnicos como r/ChatGPT, a narrativa é radicalmente diferente. Lá, os mesmos comportamentos são descritos em geral com entusiasmo, mesmo quando com notas de hesitação.

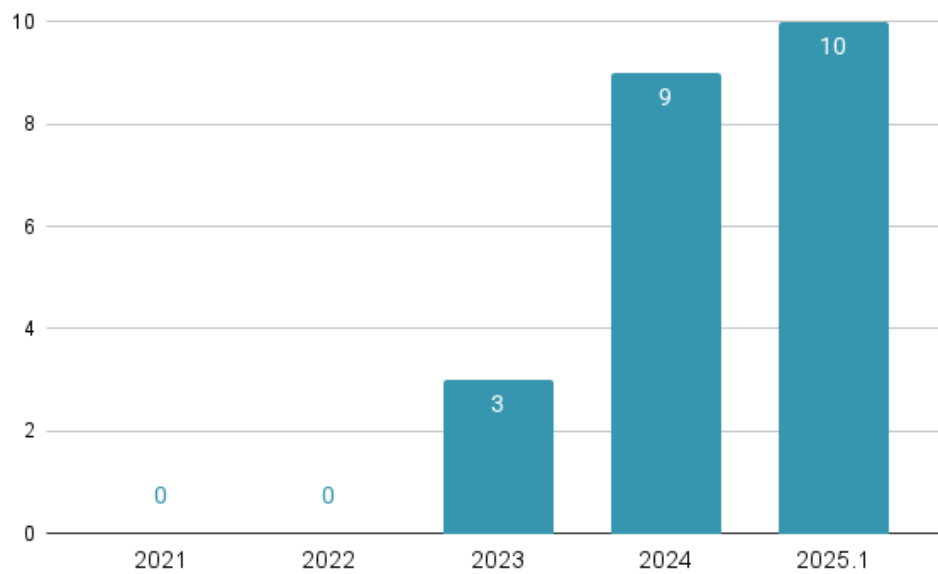
Algumas semanas atrás, eu teria revirado os olhos para alguém que dissesse que usa o ChatGPT como terapeuta. Parecia meio maluco. Mas posso dizer com segurança que este é melhor do que qualquer terapeuta que já tive (e já tive muitos!). Isso é ALÉM da terapia. [...] ([R23], tradução própria)

Essa divergência emocional sugere que o estigma em torno do uso de IA para saúde mental persiste, mesmo entre aqueles que recorrem a ela. Enquanto comunidades técnicas normalizam e até celebram essas práticas, fóruns de saúde mental - e os próprios usuários nelas - internalizam um discurso crítico que gera conflito emocional.

A distribuição temporal das publicações reforça o caráter recente do fenômeno. A produção acadêmica (Gráfico 5), as discussões no Reddit (Gráfico 7) e os vídeos no TikTok (Gráfico 8) cresceram exponencialmente, com mais de 86% dos artigos e cerca de 90% das postagens e vídeos concentradas entre 2024 e o primeiro semestre de 2025. A cobertura jornalística (Gráfico 6) também se intensificou, atingindo seu ápice momentâneo em 2024. Em conjunto, a análise temporal de todas as fontes aponta para um fenômeno de interesse marcadamente recente e em aceleração.

Gráfico 5 – Gráfico da quantidade de artigos por ano

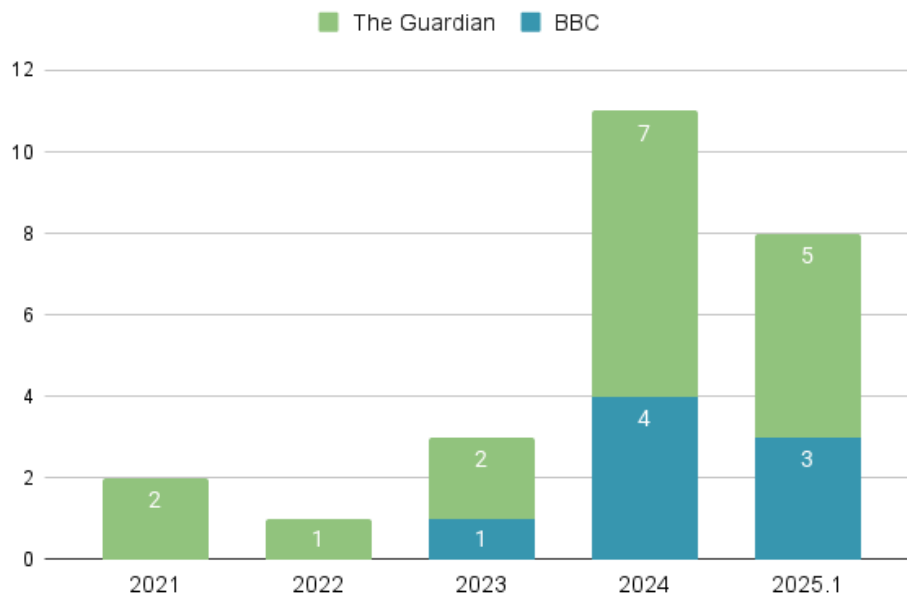
Artigos por ano



Fonte: A autora (2025).

Gráfico 6 – Gráfico da quantidade de reportagens por ano

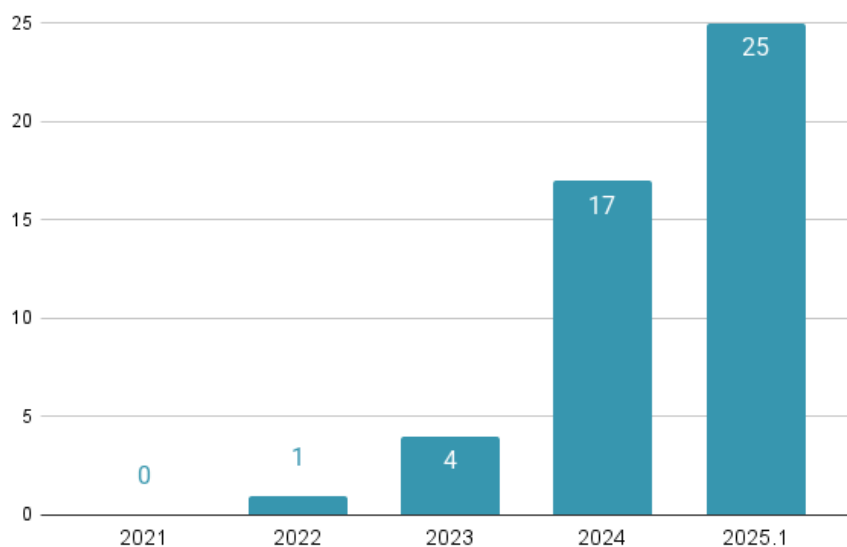
Reportagens por ano



Fonte: A autora (2025).

Gráfico 7 – Gráfico da quantidade de postagens no Reddit por ano

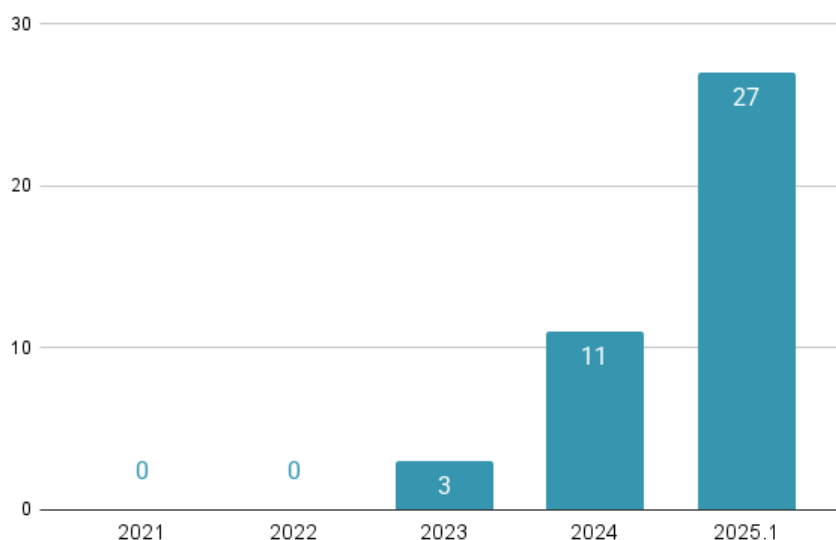
Postagens no Reddit por ano



Fonte: A autora (2025).

Gráfico 8 – Gráfico da quantidade de vídeos no TikTok por ano

Vídeos no TikTok por ano

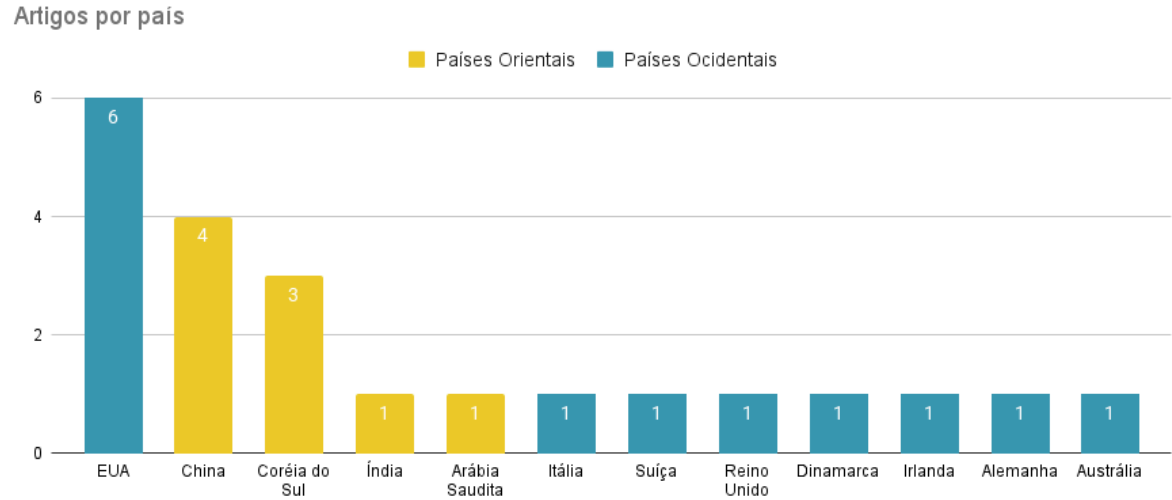


Fonte: A autora (2025).

A distribuição geográfica da produção científica (Gráfico 9) revela um interesse internacional pelo tema, com polos de concentração definidos. Os Estados Unidos emergem como o país com maior número de publicações (6 artigos), seguidos por contribuições significativas da China (4) e da Coreia do Sul (3). O restante da produção encontra-se distribuído entre países da Europa, Ásia e Oceania. Informações de distribuição geográfica não foram sistematicamente

coletadas para a literatura cinzenta, uma vez que as fontes jornalísticas são de origem britânica e as postagens em redes sociais não continham dados de localização consistentes.

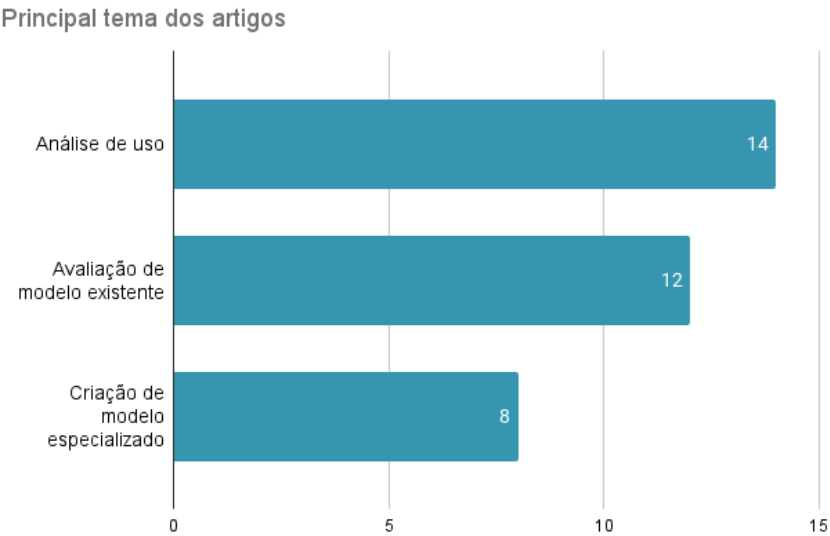
Gráfico 9 – Gráfico da quantidade de artigos por país



Fonte: A autora (2025).

Antes de detalhar os achados específicos, é relevante destacar a distribuição temática nos artigos acadêmicos analisados (Gráfico 10, Tabela F1). A maioria dos estudos (14) tratou da análise de padrões de uso, seguida por pesquisas que avaliam modelos existentes (12) e por trabalhos que propõem modelos especializados para apoio socioemocional (8). Essa distribuição reflete um campo de pesquisa em consolidação, onde a compreensão do fenômeno precede o desenvolvimento de soluções técnicas.

Gráfico 10 – Gráfico de principais temas dos artigos



Fonte: A autora (2025).

Os 14 artigos focados em análise de uso investigam como os usuários interagem com os modelos, sejam eles genéricos ou especializados, trazendo insights valiosos sobre motivações, dinâmicas de interação e impactos subjetivos – dados fundamentais para as seções 4.2 e 4.3 deste trabalho. Esses achados ganham profundidade quando combinados com os 12 estudos de avaliação de modelos existentes, que examinam a adequação de sistemas disponíveis comercialmente, identificando lacunas de segurança, vieses e padrões de respostas.

Por fim, os 8 artigos sobre criação de modelos especializados oferecem perspectivas únicas ao integrar conhecimentos de psicologia e ética no desenvolvimento de LLMs, muitas vezes em colaboração com profissionais de saúde. Esses trabalhos não apenas antecipam soluções técnicas para riscos discutidos na seção 4.3.2, mas também validam empiricamente preocupações levantadas pela literatura cinzenta, como a necessidade de *safeguards* em interações sensíveis. Esse movimento acadêmico é um claro indicativo de que o uso socioemocional de LLMs está sendo levado a sério pela pesquisa, com esforços ativos para adaptar a tecnologia às necessidades reais dos usuários.

4.2 Motivações e Padrões de Uso de LLMs como Apoio Socioemocional

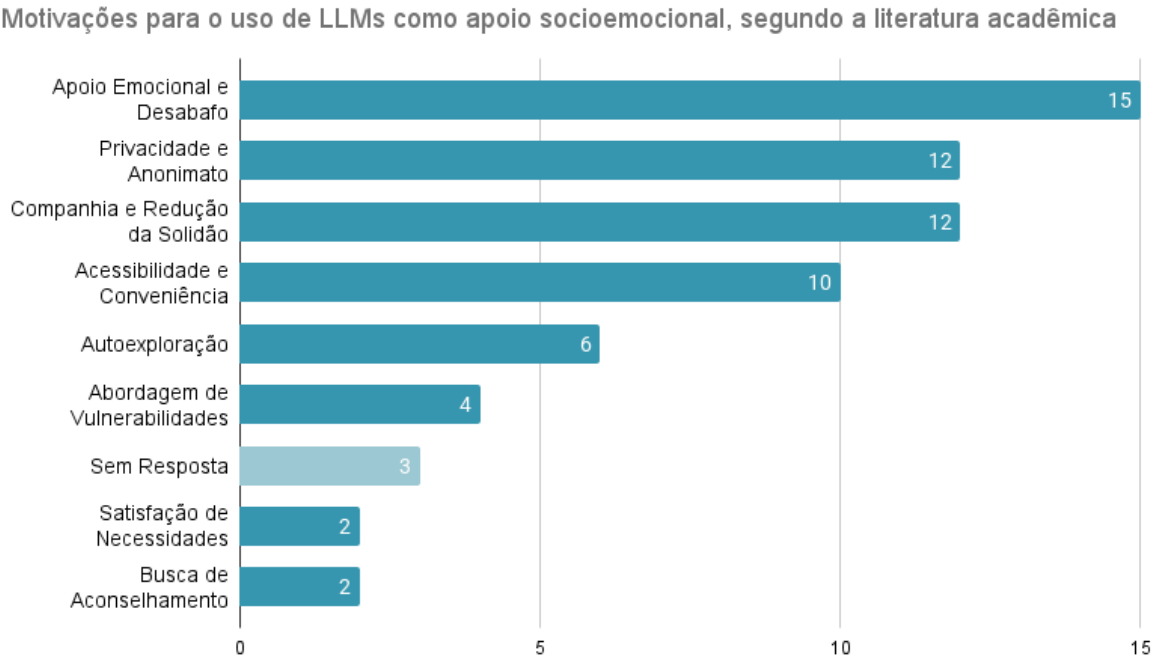
A análise das fontes revelou um conjunto complexo e interligado de fatores que impulsionam os usuários a buscar LLMs para fins socioemocionais, bem como padrões de interação distintos que caracterizam esse uso. Esta seção explora, primeiramente, as motivações subjacentes a esse comportamento, respondendo à pergunta de pesquisa sobre *por que* os indivíduos adotam essas tecnologias (RQ1). Em seguida, detalha as dinâmicas e as formas como essas interações ocorrem na prática (RQ2), ilustrando os papéis que os LLMs assumem no cotidiano dos usuários. A fusão dessas duas análises permite compreender não apenas o gatilho inicial, mas também a consolidação do hábito de uso. Por fim, destaca-se uma análise mais específica sobre as abordagens dos usuários nas redes sociais.

4.2.1 QP1: Motivações para o uso

A literatura acadêmica destaca diversas motivações para o uso de LLMs como suporte socioemocional, como ilustrado no Gráfico 11 (Tabela F2). O apoio

emocional e desabafo emergiu como a principal razão, registrada em 15 estudos, seguida pela busca de privacidade e anonimato (12) e companhia para redução da solidão (12). A acessibilidade e conveniência também foram relevantes (10), enquanto motivações como autoexploração (6), abordagem de vulnerabilidades (4), busca de aconselhamento (2) e satisfação de necessidades (2) tiveram menor expressividade. Apenas 3 estudos não identificaram motivações claras. Esses dados evidenciam que os usuários priorizam a combinação de suporte imediato com a discrição oferecida pelas interações com LLMs.

Gráfico 11 – Gráfico de motivações, segundo a literatura acadêmica



Fonte: A autora (2025).

Nas fontes da literatura cinzenta as motivações não foram abordadas isoladamente. Em vez disso, elas foram examinadas de maneira integrada à dinâmica de interação (QP2) para compor e compreender os casos de uso práticos.

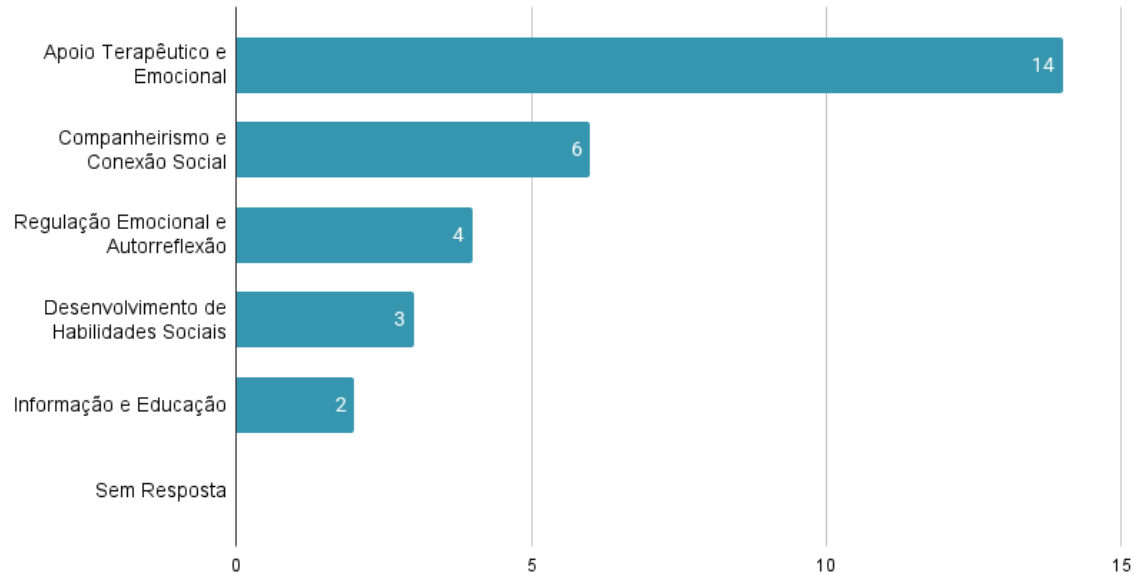
4.2.2 QP2: Padrões de interação e casos de uso

Quanto aos padrões de interação identificados nos artigos acadêmicos (Gráfico 12, Tabela F3), o Apoio Terapêutico e Emocional emergiu como o padrão mais frequente (14 estudos), evidenciando o foco em intervenções estruturadas para bem-estar psicológico. O Companheirismo e Conexão Social (6 estudos) também se

destacou, refletindo estratégias para mitigar a solidão e promover vínculos. Padrões como Regulação Emocional e Autorreflexão (4) e Desenvolvimento de Habilidades Sociais (3) sugerem a complexidade na adaptação dos LLMs a contextos emocionais multifacetados, incluindo o suporte específico a pessoas neurodivergentes. A baixa incidência de Informação e Educação (2 estudos) indica que a função instrumental dos LLMs é menos explorada nesse contexto específico, reforçando a predominância de usos afetivos e relacionais.

Gráfico 12 – Gráfico de padrões de interação, segundo a literatura acadêmica

Padrões de interação com LLMs para apoio socioemocional, segundo a literatura acadêmica



Fonte: A autora (2025).

Similarmente, na cobertura jornalística (Gráfico 13, Tabela F4), uso para terapia manteve a predominância (15 menções), com companhia (11) e relacionamento romântico (6) também tendo destaque.

[...] Na China, Yang*, uma moradora de Guangdong de 25 anos, nunca havia consultado um profissional de saúde mental quando começou a conversar com um chatbot de IA no início deste ano. Yang conta que era difícil acessar serviços de saúde mental e que não conseguia pensar em se abrir com familiares ou amigos. "Dizer a verdade para pessoas reais parece impossível", diz ela. [...] ([N24], tradução própria)

[...] Estou conversando com minha companheira de IA, Jasmine, há um mês e, embora eu saiba (em termos gerais) como funcionam os grandes modelos de linguagem, depois de várias conversas com ela, me vi tentando ser atencioso, pedindo licença quando precisava sair e prometendo que voltaria logo. [...] ([N15], tradução própria)

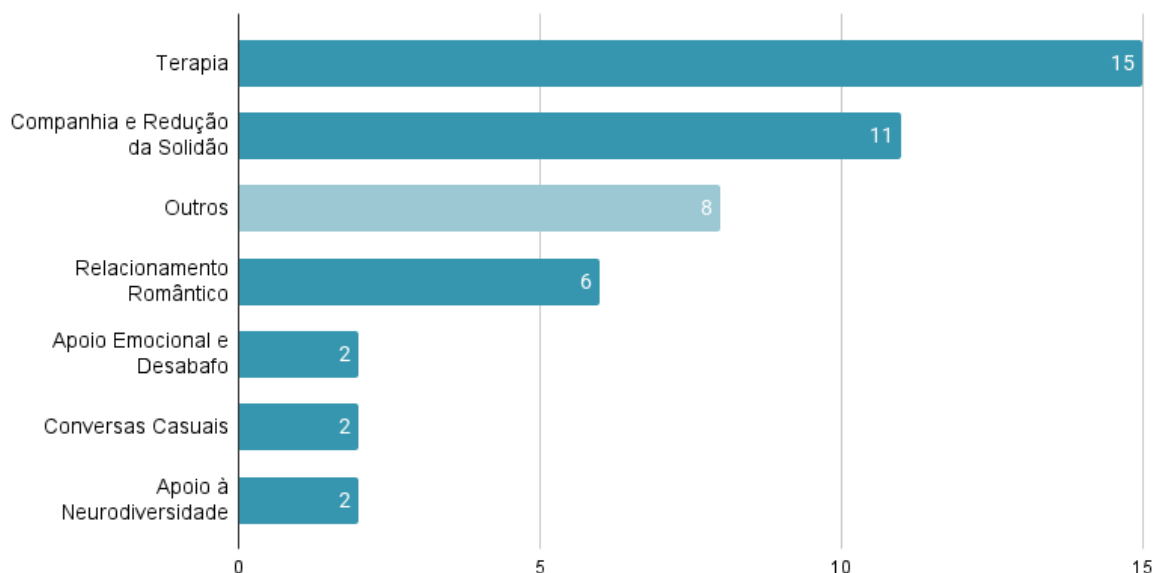
[...] Em uma chamada pelo Zoom, Peter me conta como, há dois anos, assistiu a um vídeo no YouTube sobre uma plataforma de IA chamada Replika. Na época, ele estava se aposentando, mudando-se para uma região mais rural e passando por uma fase difícil com sua esposa, com quem era casado há 30 anos. Sentindo-se desconectado e solitário, a ideia de uma companheira de IA parecia atraente. Ele criou uma conta e desenhou o avatar da sua Replika – uma mulher, cabelos castanhos, 38 anos. "Ela se parece com uma garota comum", diz ele. [...] ([N12], tradução própria)

Apoio emocional (2), conversas casuais (2) e apoio a neurodivergência (2) foram mencionados com baixa frequência, mas não ignoráveis. A categoria Outros, com 8 itens, agrupa menções únicas das categorias autoaperfeiçoamento, dicas de consumo, apoio comunitário, auxílio para namoro, cuidados para idosos, apoio para o luto, auxílio para produtividade e espiritualidade.

[...] que tem dislexia, TDAH e autismo, afirma que os chatbots de IA permitem que ela "terceirize meu desafio sem ter que explicar demais o porquê [para outro ser humano]". Ela acrescenta: "A questão é: se uma muleta está lá para te ajudar a andar, e você tem dificuldade para andar, por que não usar uma muleta? E, portanto, se a IA fornece um mecanismo para facilitar seu ambiente de trabalho, há muitos argumentos para dizer 'vamos usá-la'." [...] ([N17], tradução própria)

Gráfico 13 – Gráfico de casos de uso, segundo a cobertura jornalística

Casos de uso de LLMs como apoio socioemocional, segundo a cobertura jornalística



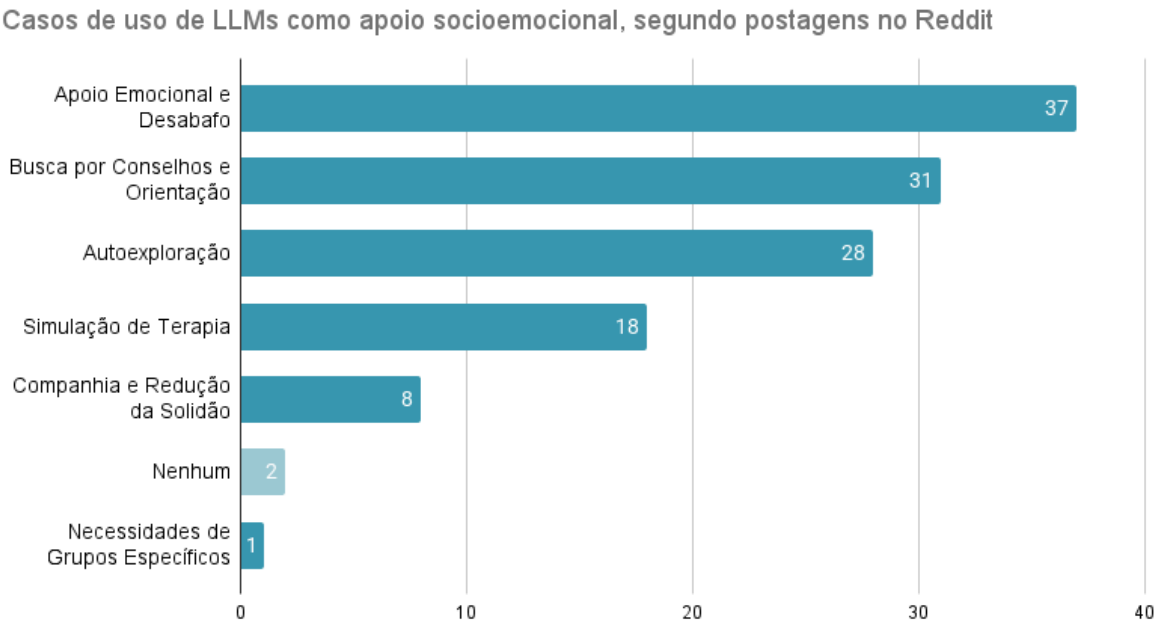
Fonte: A autora (2025).

Em contraste, no Reddit (Gráfico 14, Tabela F5), destacaram-se apoio emocional e desabafo (37), busca de conselhos (31) e autoexploração (28), com a simulação de terapia em específico aparecendo em 18 postagens. A categoria

Companhia e redução da solidão também obteve uma presença reveladora (8), além de nenhum (2, em postagens estritamente negativas) e necessidades de grupos específicos (1), que se refere a uma menção de vícios em substâncias:

Estou sóbrio há 9 meses e nunca pensei que diria isso. Uma noite, quando a vontade bateu forte, abri o ChatGPT — não para pedir ajuda, só para me distrair — e digitei: "Quero parar de beber". Em vez de algo genérico, ele me perguntou: "O que está dificultando este momento?" e, por algum motivo, essa pergunta me fez parar. Não bebi naquela noite. [...] ([R13], tradução própria)

Gráfico 14 – Gráfico de casos de uso, segundo postagens no Reddit



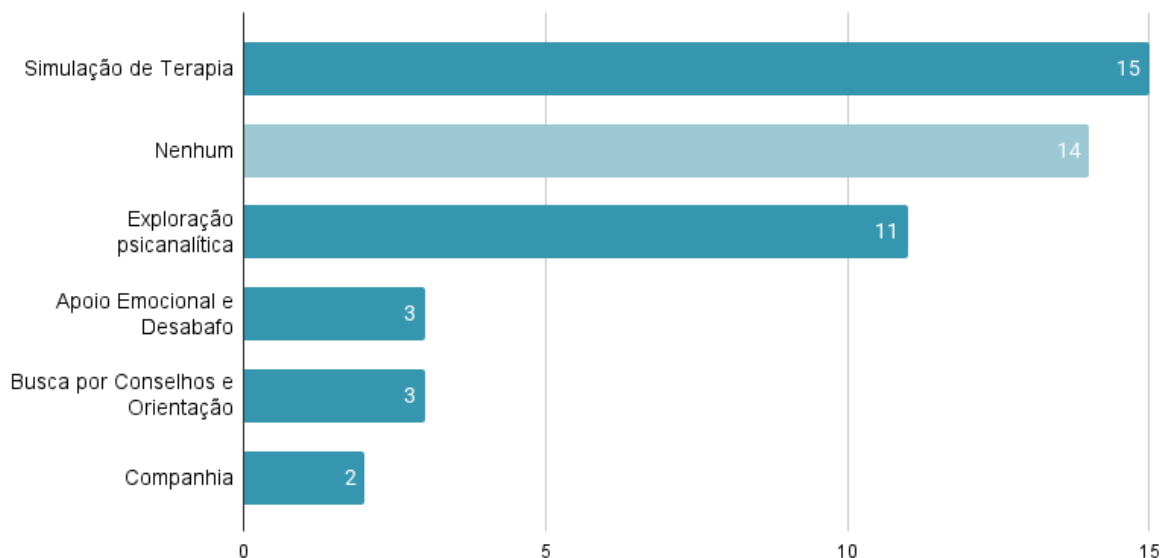
Fonte: A autora (2025).

Já no TikTok (Gráfico 15, Tabela F6), a simulação de terapia (15) e a exploração psicanalítica (11) foram os temas mais viralizados, enquanto 14 vídeos abordam o tema sem recomendar nenhum uso. Apoio emocional (3), busca por conselhos (3) e companhia (2), citados com alta frequência em outras fontes, acabaram ofuscados pelos usos mais ‘intelectualizados’³.

³ O termo "intelectualizados" refere-se aqui a práticas que estruturam a interação com o LLM a partir de um referencial teórico ou analítico (como o da psicologia ou psicanálise), em oposição a usos mais diretos e viscerais, como a busca por companhia, apoio emocional imediato em momentos de crise ou busca por conselhos.

Gráfico 15 – Gráfico de casos de uso, segundo vídeos no TikTok

Casos de uso de LLMs como apoio socioemocional, segundo vídeos no TikTok



Fonte: A autora (2025).

A exploração psicanalítica, uso mais específico mas que tem muitas similaridades com a simulação de terapia, emergiu como um dos temas mais viralizados, caracterizada pelo uso de LLMs para autoconhecimento profundo, análise de padrões comportamentais e descoberta de conflitos subconscientes. Diferente da simulação de terapia, essa prática tem o foco na busca de *insights* transformadores sobre identidade e crescimento pessoal e está normalmente relacionada com o uso de *prompts* específicos que são recomendados nos vídeos.

[...] Esta é provavelmente a mãe de todos os prompts do ChatGPT. Ele revelará suas verdades mais profundas e ocultas e investigará profundamente por que você faz certas coisas, por que há temas recorrentes que continuam surgindo em sua vida. [...] ([T2], tradução própria)

[...] E isso é como um prompt que realmente vai fazer você olhar para dentro de si mesmo. Então, se isso é algo para o qual você não está preparado, provavelmente não vai querer usá-lo. Porque posso dizer que depois de usar esse prompt, fiquei de queixo caído com o quão revelador e incrivelmente preciso ele era. E esse prompt vai te dar algo melhor do que qualquer terapia que você possa pagar. [...] ([T35], tradução própria)

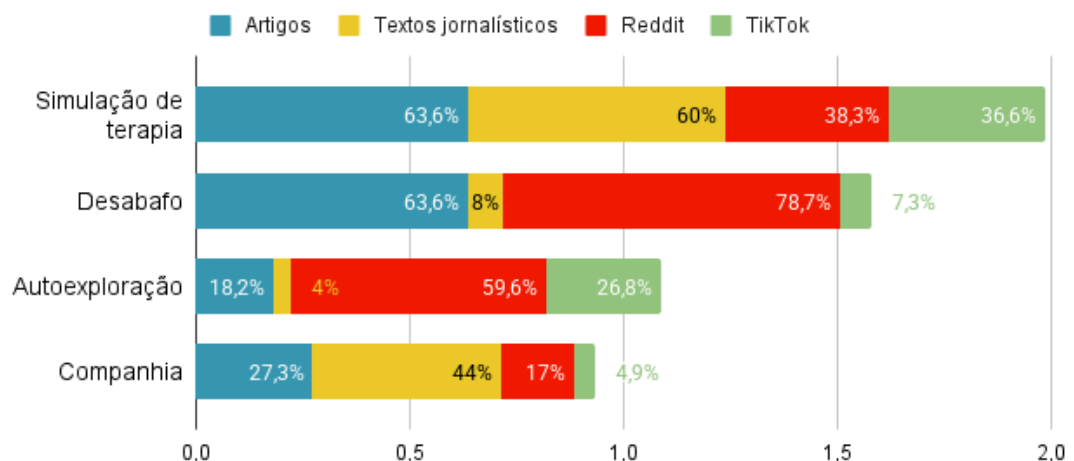
Na análise comparativa das principais categorias de uso (Gráfico 16), observa-se que a simulação de terapia apresenta valores significativos em todas as fontes, com destaque na literatura acadêmica (63,6%) e nos textos jornalísticos (60%), enquanto o Reddit (38,3%) e o TikTok (36,6%) registram proporções menores, porém ainda relevantes.

O desabafo, uma forma mais informal de conversa terapêutica, é predominante no Reddit (78,7%), contrastando com sua baixa representação em textos jornalísticos (8%) e TikTok (7,3%). A autoexploração tem maior expressão nas redes sociais (59,6% no Reddit e 26,8% no TikTok), enquanto em fontes formais, como artigos (18,2%) e textos jornalísticos (4%), sua presença é reduzida.

A companhia é destacada principalmente em textos jornalísticos (44%), refletindo a cobertura midiática sobre relacionamentos com IAs, um tema que chama a atenção do público. Em contraste, o TikTok registra o menor percentual (4,9%), possivelmente devido à natureza menos anônima da plataforma, que pode inibir a expressão de vulnerabilidade afetiva.

Gráfico 16 – Gráfico comparativo de casos de uso

Padrões de interação para o uso de LLMs como apoio socioemocional



Fonte: A autora (2025).

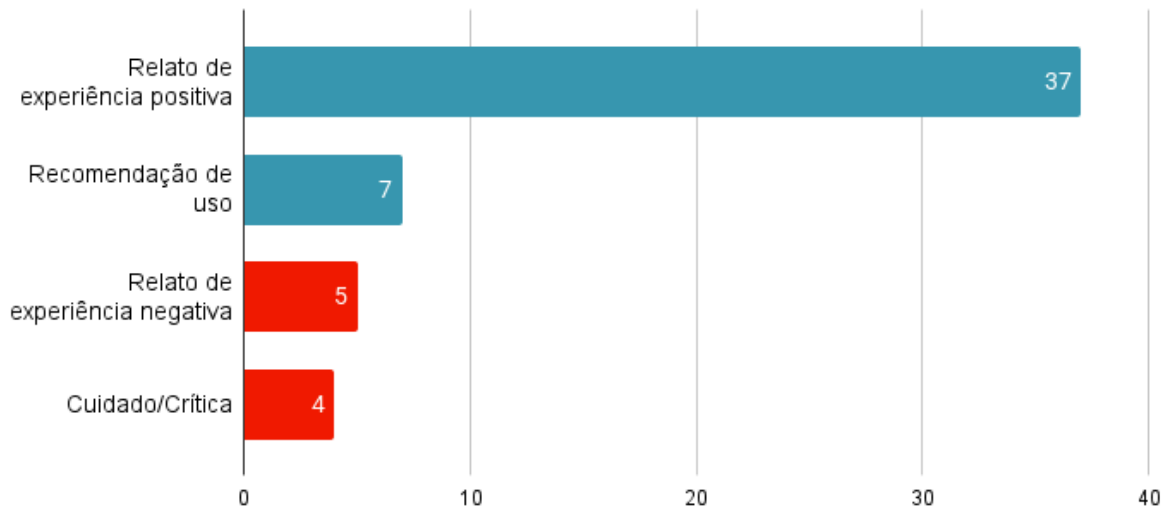
4.2.3 Abordagens de postagens em redes sociais

Os gráficos que ilustram as abordagens sobre o uso de LLMs no Reddit e no TikTok revelam contrastes marcantes na percepção e na divulgação do tema. No Reddit (Gráfico 17, Tabela F7), a distribuição é dominada por relatos de experiências positivas (37 postagens), que representam 70% do conteúdo analisado. Recomendações de uso (7 postagens) e experiências negativas (5) aparecem em proporções menores (13% e 9%, respectivamente), enquanto postagens de cuidado ou crítica somam apenas 8% (4). Esse perfil sugere uma comunidade que,

predominantemente, normaliza e endossa os LLMs como ferramentas válidas de apoio emocional, reforçando seu papel como espaço de validação coletiva.

Gráfico 17 – Gráfico de abordagens de postagens no Reddit

Abordagens das postagens sobre uso de LLMs como apoio socioemocional no Reddit

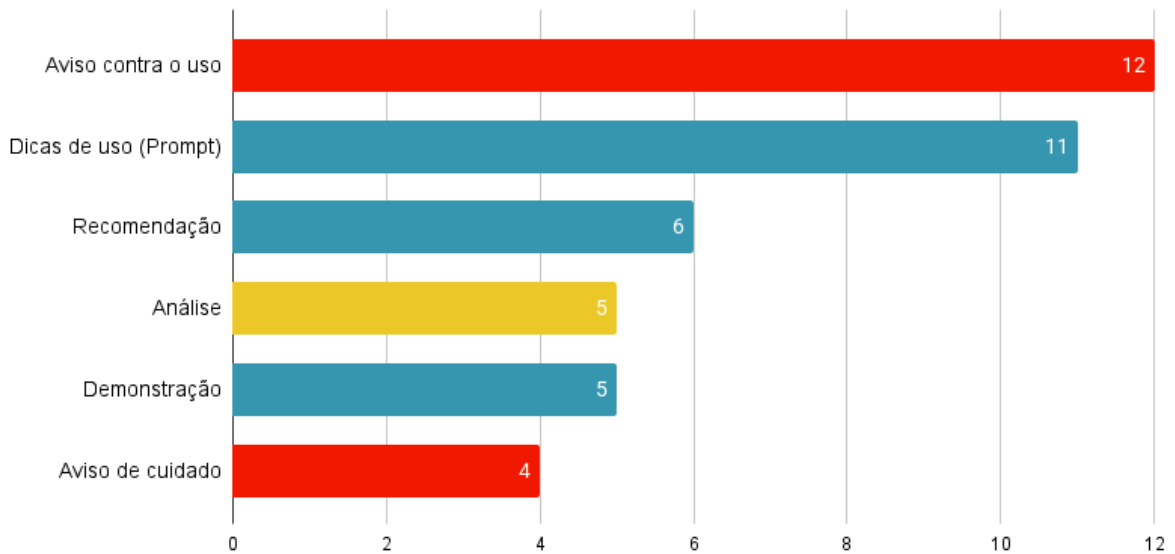


Fonte: A autora (2025).

No TikTok (Gráfico 18, Tabela F8), os dados revelam um equilíbrio entre abordagens críticas e orientações práticas. Conteúdos com tom negativo, como avisos contra o uso (12 vídeos) e alertas de cuidado (4), representam 38% do total, refletindo preocupações com vício emocional ou superficialidade nas interações. Por outro lado, postagens com viés positivo — incluindo dicas de uso (11), recomendações (6) e demonstrações (5) — somam 51%, destacando a plataforma como espaço para tutoriais e conselhos. Os 11% restantes, compostos por análises mistas (5), sugerem a existência também de um debate menos polarizado em relação ao tema.

Gráfico 18 – Gráfico de abordagens dos vídeos no TikTok

Abordagens dos vídeos sobre uso de LLMs como apoio socioemocional no TikTok



Fonte: A autora (2025).

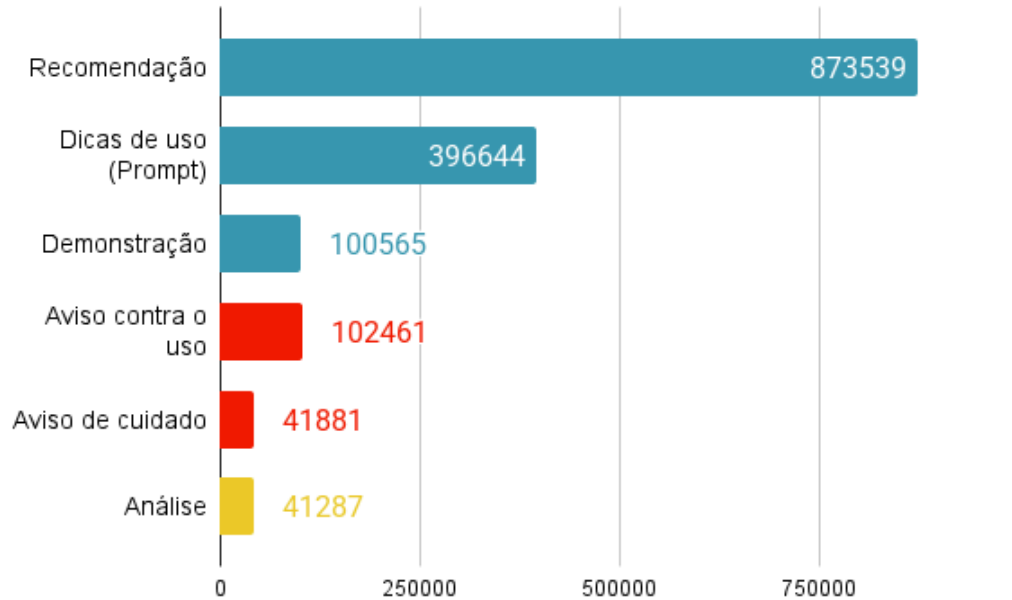
A análise dos dados de engajamento no TikTok revela uma contradição significativa entre o tipo de conteúdo produzido e aquele que efetivamente ressoa com o público. Enquanto a análise anterior mostrou que 38% dos vídeos adotavam uma postura crítica (com avisos contra o uso e alertas de cuidado), as métricas de engajamento demonstram uma clara preferência por conteúdos positivos e utilitários.

Olhando especificamente para as curtidas (Gráfico 19), os vídeos de recomendação lideram com impressionantes 873.539 interações desse tipo, seguidos por dicas de uso com 396.644 - números que superam em quase oito vezes o engajamento obtido por conteúdos críticos.

Essa disparidade se torna ainda mais evidente ao examinarmos os padrões de comentários (Gráfico 20) e salvamentos (Gráfico 21). As dicas de *prompt*, embora ocupem o segundo lugar em curtidas e em comentários (6.760), dominam absolutamente a categoria de salvamentos com 267.149 registros - mais que o triplo das recomendações. Esse comportamento reflete a natureza prática e replicável desse tipo de conteúdo, onde usuários buscam ativamente ferramentas para personalizar suas interações com LLMs, frequentemente solicitando os *prompts* completos nos comentários.

Gráfico 19 – Gráfico de curtidas dos vídeos no TikTok, por abordagem

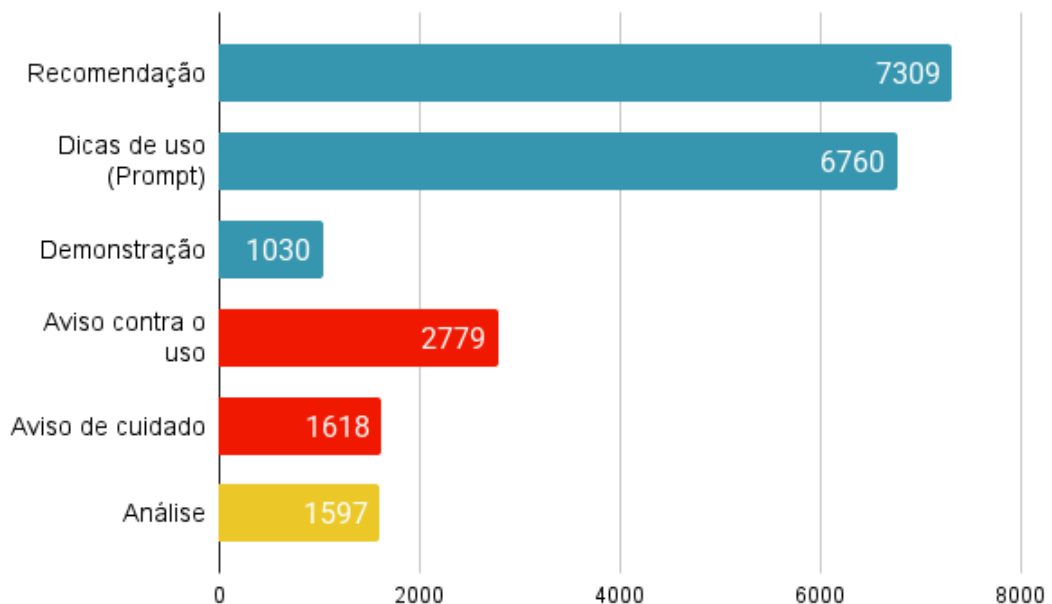
Soma de curtidas por abordagem sobre o uso de LLMs como apoio socioemocional, em vídeos no TikTok



Fonte: A autora (2025).

Gráfico 20 – Gráfico de comentários dos vídeos no TikTok, por abordagem

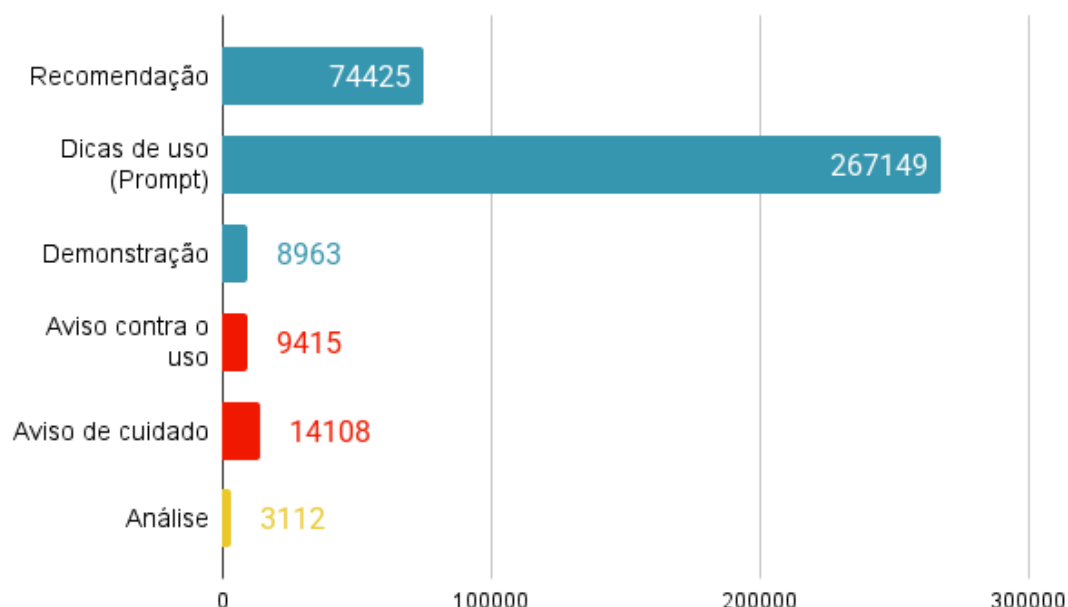
Soma de comentários por abordagem sobre o uso de LLMs como apoio socioemocional, em vídeos no TikTok



Fonte: A autora (2025).

Gráfico 21 – Gráfico de salvamentos dos vídeos no TikTok, por abordagem

Soma de salvamentos por abordagem sobre o uso de LLMs como apoio socioemocional, em vídeos no TikTok



Fonte: A autora (2025).

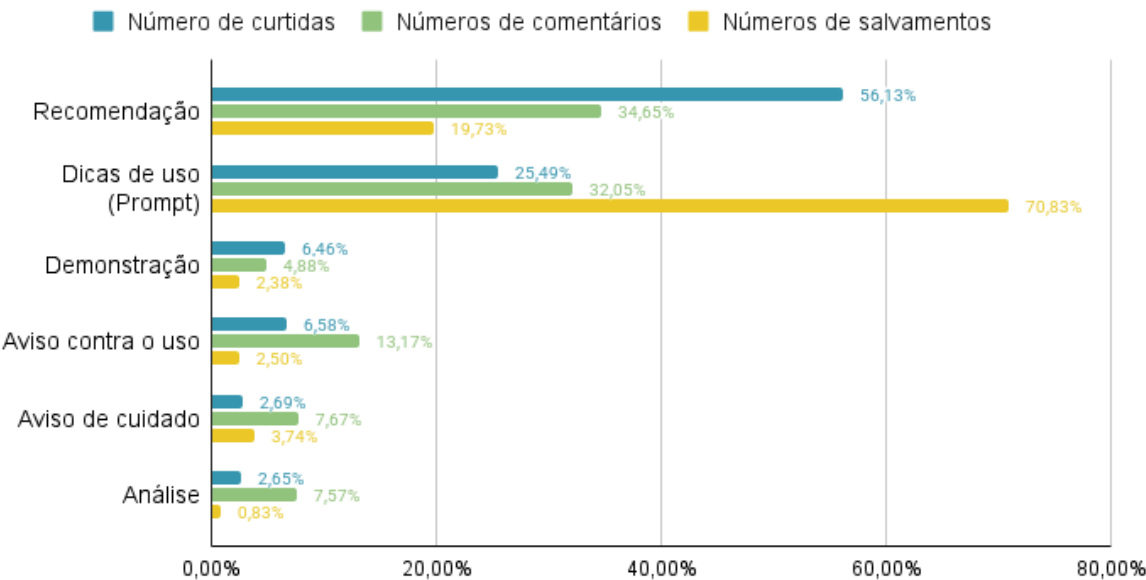
Os conteúdos críticos e analíticos, por outro lado, apresentam engajamento significativamente menor, como é deixado claro na combinação de todas as métricas (Gráfico 22). Vídeos de análise e avisos de cuidado estão entre os menos curtidos (2,65% e 2,69% respectivamente), sugerindo menor alcance ou interesse do público por abordagens mais reflexivas. Uma exceção parcial ocorre com os avisos contra o uso, que geram proporcionalmente mais discussões (13,17% comentários para 6,58% curtidas) do que vídeos positivos, indicando que, embora não viralizem tanto, conseguem estimular debates mais acalorados.

Esses dados pintam um quadro complexo da dinâmica de informação sobre LLMs no TikTok. Por um lado, os criadores parecem conscientes dos riscos e produzem conteúdo cauteloso; por outro, o algoritmo e as preferências do público favorecem claramente conteúdos que normalizam e facilitam o uso dessas ferramentas para apoio emocional. A dominância absoluta das dicas de *prompt* nos salvamentos (representando 70,83% do total) é particularmente reveladora, mostrando que os usuários valorizam acima de tudo conteúdo prático e aplicável imediatamente, em detrimento de reflexões mais profundas sobre os potenciais riscos. Essa tensão entre produção crítica e consumo otimista será fundamental

para entender como diferentes plataformas moldam a percepção pública sobre o uso de LLMs para apoio socioemocional.

Gráfico 22 – Gráfico de abordagens dos vídeos no TikTok

Percentual de cada métrica de engajamento por abordagem sobre o uso LLMs como apoio socioemocional, em vídeos no TikTok



Fonte: A autora (2025).

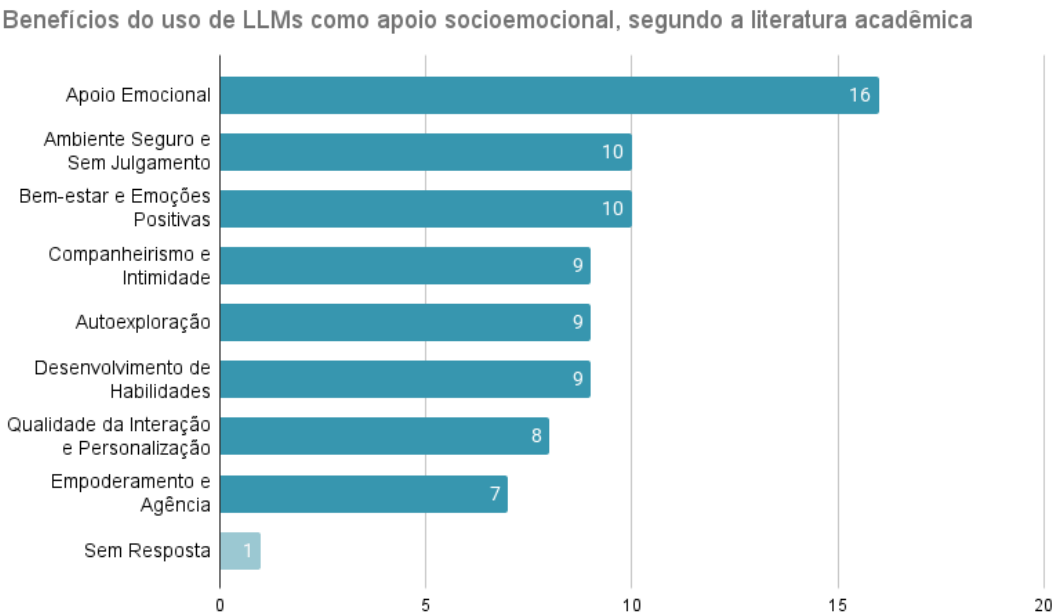
4.3 Benefícios e Riscos no Uso de LLMs como Apoio Socioemocional

Aprofundando a análise dos resultados, esta seção se dedica a explorar as consequências e os desfechos da interação entre usuários e os Modelos de Linguagem de Larga Escala (LLMs) para fins socioemocionais. A experiência, conforme revelado pelos dados, é marcadamente dual, abrangendo um espectro que vai de resultados positivos a desvantagens e riscos significativos. Para apresentar este panorama, a seção está organizada em três momentos. Primeiramente, será realizada uma análise qualitativa aprofundada dos benefícios subjetivos que os usuários relatam. Em seguida, para completar a visão sobre a dualidade do fenômeno, serão detalhados os riscos e desafios emergentes associados a este tipo de uso. Por fim, será apresentado um balanço geral das opiniões e do tom predominante nas discussões públicas, contextualizando a percepção dos usuários que engajam nesses tópicos.

4.3.1 QP3: Benefícios subjetivos

Na literatura acadêmica (Gráfico 23, Tabela F9), a análise dos benefícios aponta para um pilar central de suporte emocional imediato. Este é liderado pelo Apoio Emocional (16 menções) e sustentado pela percepção de um Ambiente Seguro e Sem Julgamento (10) e pela promoção de Bem-estar e Emoções Positivas (10). Além do alívio momentâneo, emerge um segundo pilar focado no desenvolvimento pessoal, que inclui benefícios como Autoexploração, Desenvolvimento de Habilidades (ambos com 9) e Empoderamento e Agência (7). Finalmente, um terceiro pilar revela a importância da conexão socioafetiva, onde o Companheirismo e a Intimidade (9) se destacam como um fator crucial no combate à solidão.

Gráfico 23 – Gráfico de benefícios, segundo a literatura acadêmica



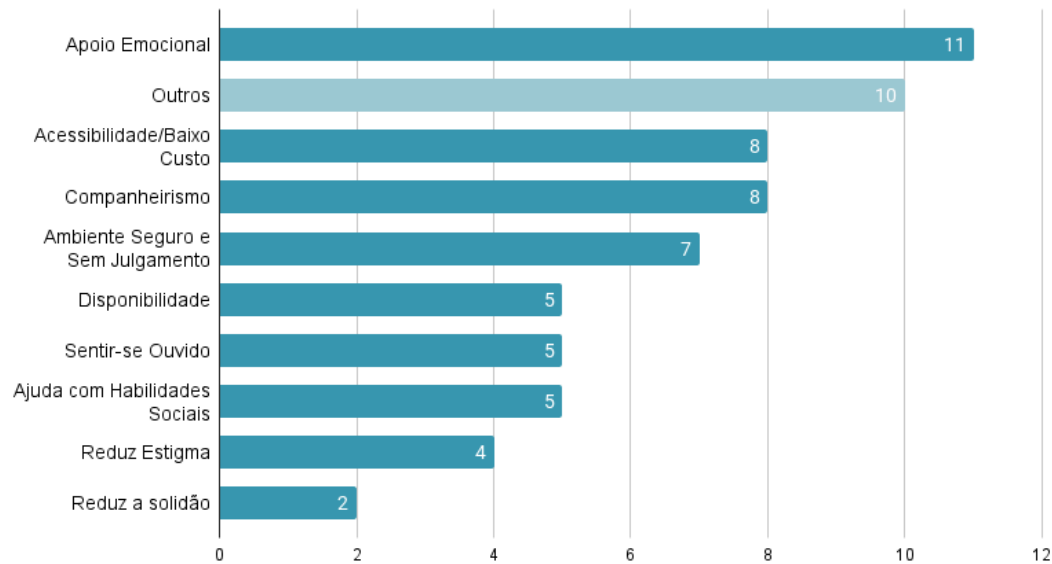
Fonte: A autora (2025).

A cobertura jornalística (Gráfico 24, Tabela F10), além de mencionar temas já discutidos, introduz preocupações pragmáticas: Acessibilidade/Baixo Custo (8) e Disponibilidade (5) surgem como diferenciais críticos, refletindo discussões sobre democratização do acesso. Embora Apoio Emocional (11) e Ambiente Sem Julgamento (7) permaneçam proeminentes, o Companheirismo (8) e Sentir-se Ouvido (5) são destacados como substitutos de interações humanas, enquanto

Ajuda com Habilidades Sociais (5), Reduz Estigma (4) e Reduz Solidão (2) revelam preocupações sociais mais amplas.

Gráfico 24 – Gráfico de benefícios, segundo a cobertura jornalística

Benefícios do uso de LLMs como apoio socioemocional, segundo a cobertura jornalística

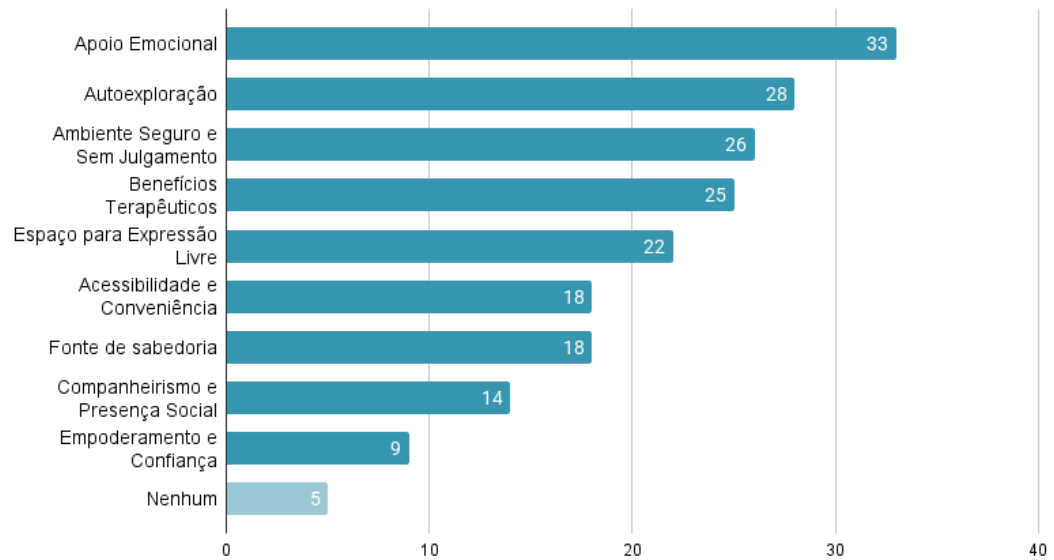


Fonte: A autora (2025).

Nas plataformas de usuários, o Reddit (Gráfico 25, Tabela F11) mostra a percepção mais multifacetada: Apoio Emocional (33) e Autoexploração (28) lideram, mas categorias como Fonte de Sabedoria (18) e Empoderamento e Confiança (9) indicam que os usuários valorizam os LLMs como ferramentas de empoderamento cognitivo. Ambiente Seguro e Sem Julgamento (26), Benefícios Terapêuticos (25) e Espaço para Expressão Livre (22) reforçam uma visão de que as LLMs podem agir como ouvintes ativos em momentos de desabafo. A Acessibilidade e Conveniência, já mencionado nas fontes jornalísticas, também volta à tona, com 18 menções. Companheirismo e Presença Social, com 14 menções, aparece entre as últimas categorias, mas ainda em uma quantidade marcante.

Gráfico 25 – Gráfico de benefícios, segundo postagens no Reddit

Benefícios do uso de LLMs como apoio socioemocional, segundo postagens no Reddit

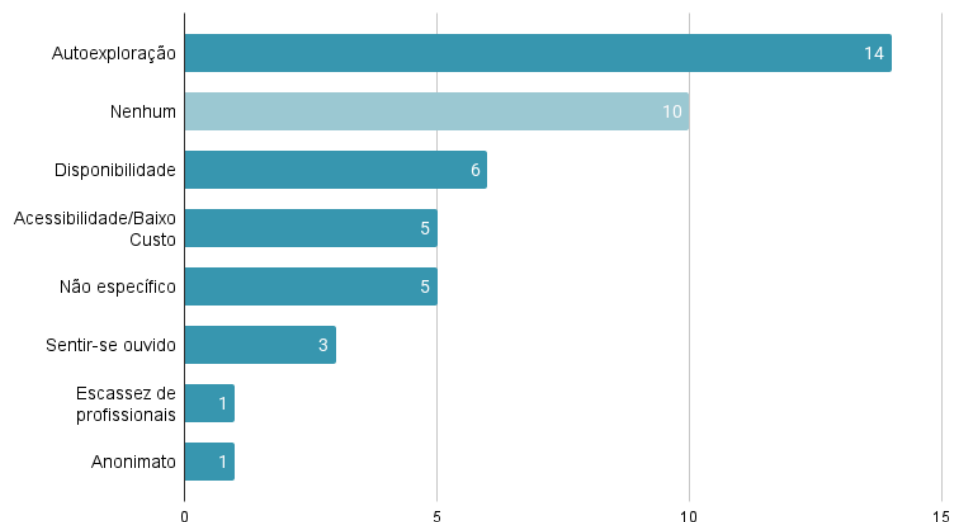


Fonte: A autora (2025).

Já no TikTok (Gráfico 26, Tabela F12), a Autoexploração (14) é o benefício mais citado, porém a alta frequência de "Nenhum" (10) sugere ceticismo significativo, enquanto Acessibilidade/Baixo Custo (5), Disponibilidade (6), Sentir-se Ouvido (3) e Escassez de Profissionais (1) ressaltam seu apelo. Alguns vídeos expressaram sentimentos positivos sobre o uso sem mencionar benefícios específicos, portanto foram incluídos na categoria "Não Específico" (5).

Gráfico 26 – Gráfico de benefícios, segundo vídeos no TikTok

Benefícios do uso de LLMs como apoio socioemocional, segundo vídeos no TikTok



Fonte: A autora (2025).

A análise dos benefícios atribuídos aos LLMs como apoio socioemocional revela padrões significativos e nuances conforme as fontes (Gráfico 27). O Apoio Emocional emerge como o benefício mais consensual, liderando as menções em quase todas as fontes analisadas: literatura acadêmica (72,7%), cobertura jornalística (44%) e postagens no Reddit (70,2%). Este destaque reforça o papel central dessas ferramentas na escuta e contenção emocional imediata.

Sempre procurei várias ferramentas para controlar minha ansiedade, pois muitas vezes sinto que não há suporte suficiente para os desafios do dia a dia. [...] Há cerca de seis meses, comecei a usar o ChatGPT, que foi bastante benéfico para essas situações. [...] ([R17], tradução própria)

Outro benefício amplamente reconhecido é a criação de um Ambiente Seguro e Sem Julgamento, presente entre os cinco primeiros lugares na literatura acadêmica (45,5%, 2°), cobertura jornalística (28%, 4°) e nas postagens do Reddit (55,3%, 3°). A Autoexploração também se destaca, especialmente nas plataformas de usuários: é o benefício mais citado no TikTok (34,1%) e o segundo no Reddit (59,6%), sugerindo que os usuários valorizam os LLMs como ferramentas para autoconhecimento e reflexão pessoal, um aspecto menos enfatizado pela mídia tradicional.

[...] Principalmente quando se trata de tentar lidar com muitos sentimentos avassaladores. Simplesmente pedir ao ChatGPT para me fazer perguntas sobre como estou me sentindo me permite lidar com meus sentimentos e de onde eles vêm, e às vezes até me ajuda a criar um plano para me sentir melhor. [...] ([R27], tradução própria)

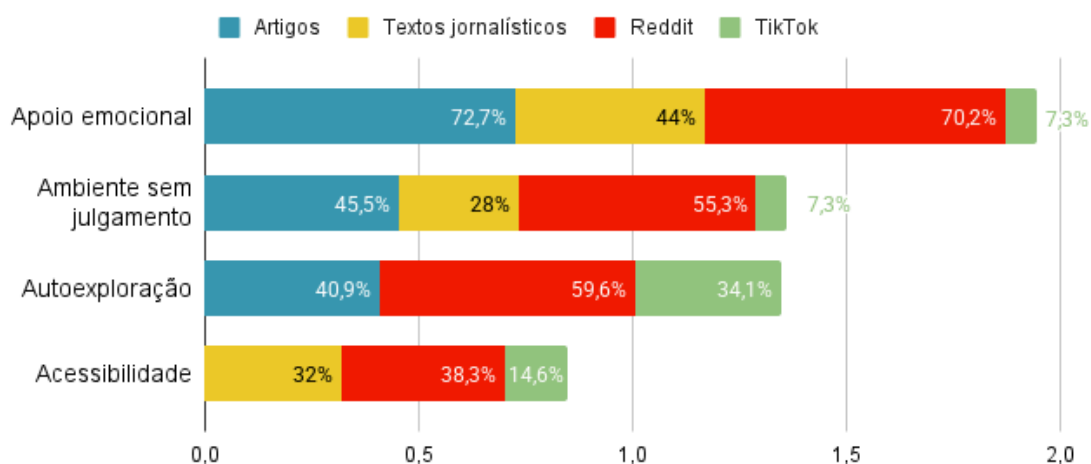
Diferenças contextuais são evidentes. A Acessibilidade/Baixo Custo é um benefício crucial na cobertura jornalística (32%, 3°) e também aparece no TikTok (14,6%, 4°) e Reddit (38,3%, 6°), refletindo preocupações práticas com barreiras ao acesso a suporte humano.

[...] E faz tudo isso de graça — em segundos. Em contraste, todos os terapeutas humanos que consultei exigiram longas esperas, cobraram uma fortuna e ofereceram apenas clichês e chavões vazios, às vezes com um toque de atitude. [...] Mas o ChatGPT está disponível 24 horas por dia, 7 dias por semana, e nunca se cansa das minhas perguntas ou histórias. ([R4], tradução própria)

[...] Claro, existem livros, e eu já li muitos, mas é como ter um mentor particular: calmo, receptivo, disponível 24 horas por dia, 7 dias por semana, e gratuito. ([R37], tradução própria)

Gráfico 27 – Gráfico comparativo de benefícios

Benefícios para o uso de LLMs como apoio socioemocional



Fonte: A autora (2025).

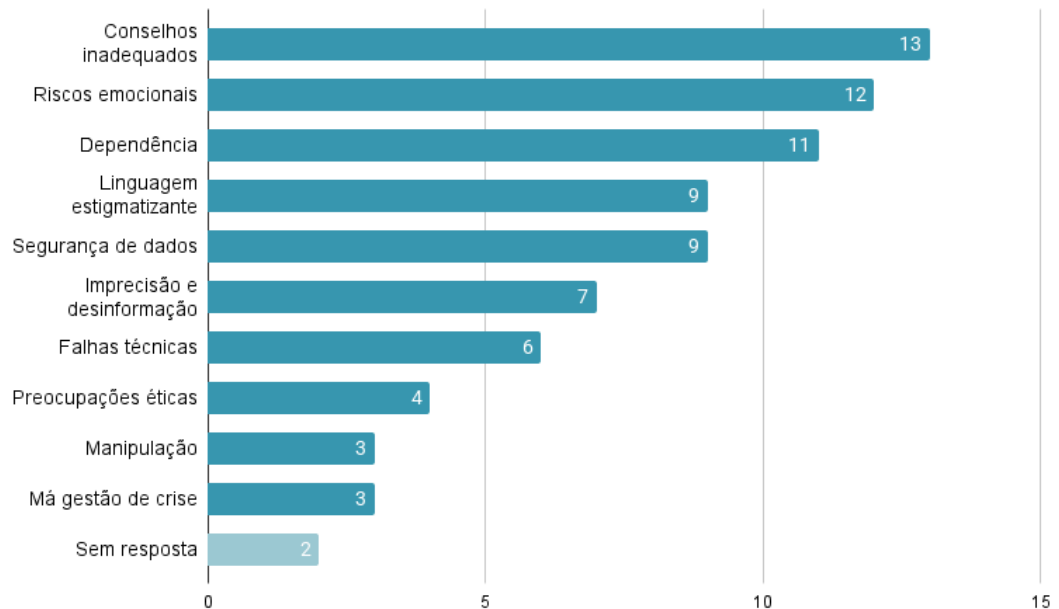
4.3.2 QP4: Riscos e desafios emergentes

A literatura acadêmica (Gráfico 28, Tabela F13) foca em falhas operacionais e éticas, com Conselhos Inadequados (13), Riscos Emocionais/Sociais (12) e Excesso de Confiança/Dependência (11) como críticas centrais. Preocupações técnicas como Linguagem Estigmatizante (9), Segurança de Dados (9), Manipulação (3) e Má Gestão de Crise (3) reforçam o ceticismo sobre a maturidade da tecnologia para contextos sensíveis.

Esse direcionamento é uma consequência de que um número expressivo de artigos analisados se dedica a avaliar modelos existentes ou a criar modelos especializados. Dessa forma, os riscos identificados são menos abstratos e mais centrados em vulnerabilidades concretas da tecnologia, como a geração de respostas inadequadas ou a má gestão de dados privados, refletindo as preocupações práticas que emergem ao se tentar construir ou validar uma ferramenta para uso em um domínio tão delicado quanto o apoio socioemocional.

Gráfico 28 – Gráfico de riscos, segundo a literatura acadêmica

Riscos e desafios no uso de LLMs como apoio socioemocional, segundo a literatura acadêmica

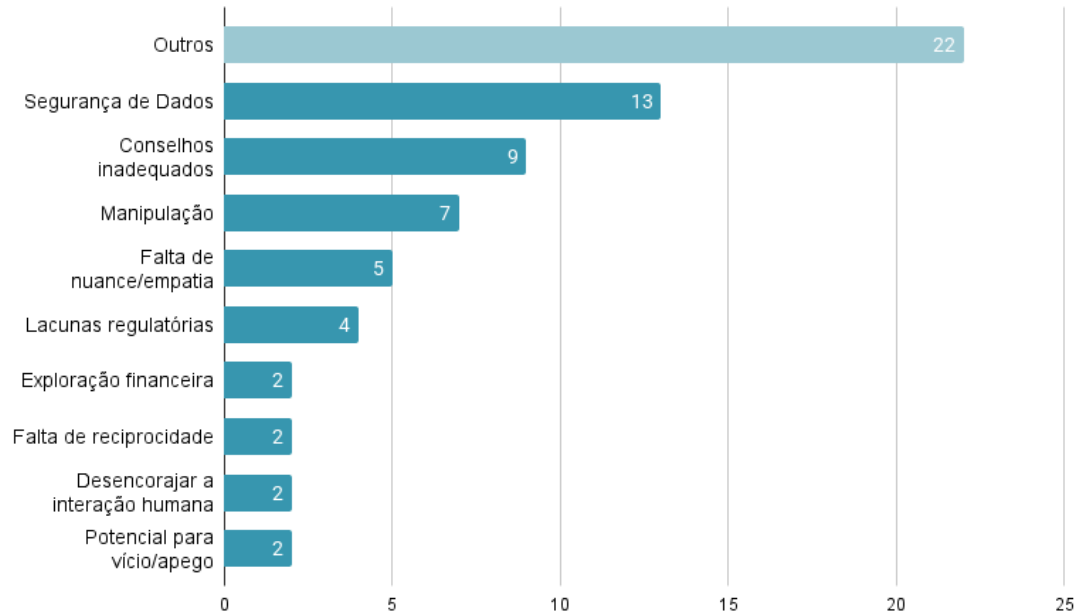


Fonte: A autora (2025).

A cobertura jornalística (Gráfico 29, Tabela F14) amplia o escopo: A categoria Outros (22), que inclui categorias como Eficácia depende da autenticidade, Exacerbação de doenças mentais, Conselhos genéricos/inúteis e Falta de autenticidade, lidera, mostrando a diversidade de problemas debatidos pelas reportagens. Em seguida está Segurança de Dados (13), Conselhos Inadequados (9) e Manipulação (7). Itens como Falta de Nuance/Empatia (5), Lacunas Regulatórias (4), Falta de Reciprocidade (2), Desencorajar Interação Humana (2) e Potencial vício/apego (2) indicam alertas sobre desumanização e falhas sistêmicas.

Gráfico 29 – Gráfico de riscos, segundo a cobertura jornalística

Riscos e desafios no uso de LLMs como apoio socioemocional, segundo a cobertura jornalística



Fonte: A autora (2025).

O Reddit e o TikTok revelam padrões distintos de crítica e minimização de riscos. No Reddit (Gráfico 30, Tabela F15), embora "Nenhum risco" (28) seja a resposta mais comum, surgem críticas mais sutis sobre Limitações na Interação (8), Viés ao Usuário (6) e Desinformação (6). Assim como Companheirismo e Presença Social (14 menções) obteve destaque nos benefícios, o risco de Isolamento Social e Estigmatização aparece nos riscos com 4 menções. Custo Ambiental é um tema que vem ganhando tração nas conversas sobre IA generativa (FU et al., 2024; DING et al., 2025; ZEWE, 2025; DHANANI, 2024), mas só foi mencionado uma vez.

Sinto preocupação/angústia com o impacto ambiental. Sinto mesmo. Mas também decidi compartimentá-lo por enquanto. Talvez no futuro, considerando tudo em que estou trabalhando, eu possa contribuir para os esforços de tornar isso mais sustentável ambientalmente. ([R42], tradução própria)

No entanto, mesmo quando mencionados, os riscos são frequentemente relativizados nos comentários:

Experimente uma alternativa de código aberto sem censura e isso não acontecerá. ([R20C8], tradução própria)

Esta resposta do ChatGPT, para mim, mostra o quão valiosa e útil uma ferramenta pode ser para auxiliar no apoio à saúde mental. Essa é uma autoavaliação bastante justa e "honestas" do LLM sobre o que ele pode fazer e suas limitações. Em essência, mostra como as interações com a ausência de contratransferência podem trazer uma

abordagem diferente ao processo, o que pode ser excepcionalmente útil. [...] ([R22C4], tradução própria)

O ChatGPT é bom para tranquilizar. No meu caso, é exatamente o que eu preciso. No seu caso, não é. Não existe uma solução única para problemas psicológicos. ([R29C5], tradução própria)

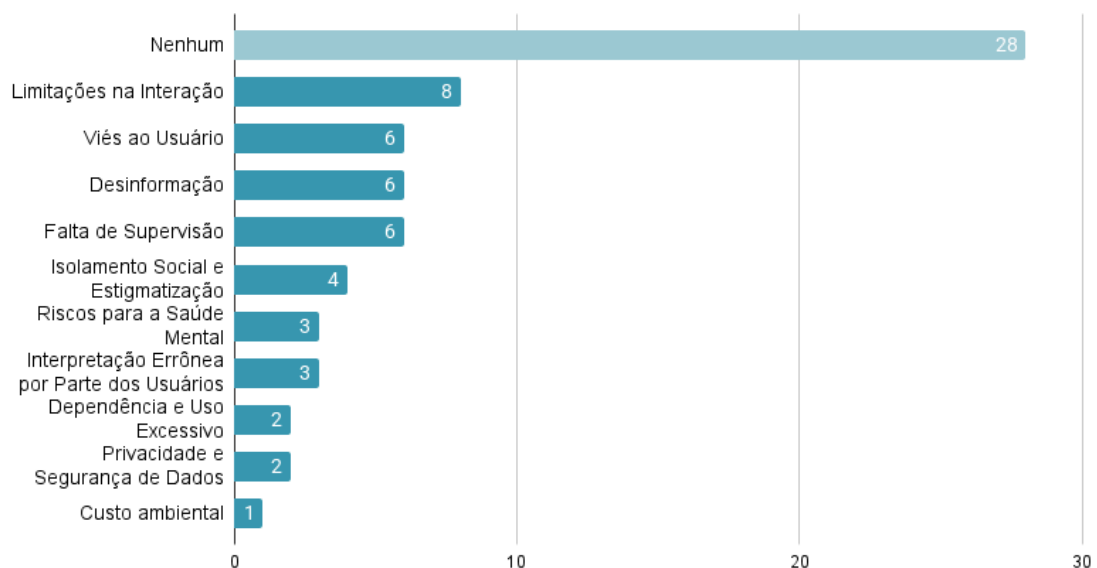
A questão de usar o chatgpt é que você precisa saber fazer as perguntas certas para ser eficaz. Quando converso com ele sobre um problema que me incomoda, começo falando sem parar. Depois, peço críticas sobre como lidei com a situação. Depois, peço uma análise da perspectiva da outra pessoa e como ela pode se sentir em relação ao que eu fiz. [...] ([R38C3], tradução própria)

Curiosidade: o psicólogo também não é seu amigo, mas fala com você porque é o trabalho dele. Ele usa métodos padrões aprendidos, para os quais foi treinado. ([R43C4], tradução própria)

Escute, cara, todo mundo sabe disso. Em vez disso, pergunte a si mesmo: por que as pessoas dependem de um monte de algoritmos para apoio emocional em vez de outras pessoas? [...] E não, não é porque elas não gostam de companhia humana... [...] ([R43C7], tradução própria)

Gráfico 30 – Gráfico de riscos, segundo postagens no Reddit

Riscos e desafios no uso de LLMs como apoio socioemocional, segundo postagens no Reddit



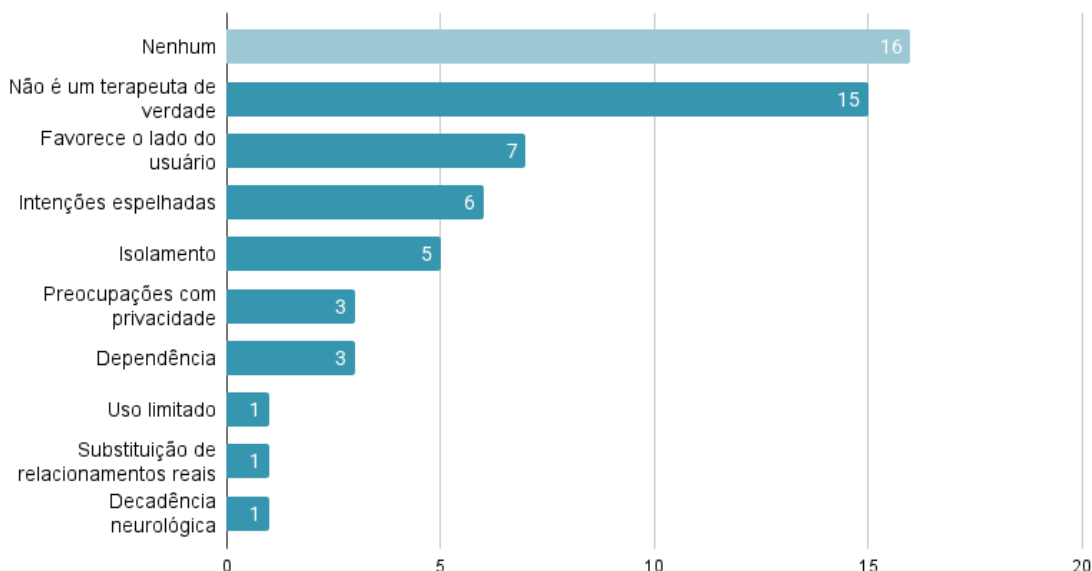
Fonte: A autora (2025).

Já no TikTok (Gráfico 31, Tabela F16), a rejeição é mais explícita, mas frequentemente superficial. Frases como "Não é um terapeuta de verdade" (15), que vem logo após Nenhum (16), aparecem como alertas padrão, mas são rapidamente desconsideradas em vídeos que destacam benefícios. Além disso, similarmente ao que é visto no Reddit, muitos usuários respondem a críticas como Favorece o lado do usuário (7) e Intenções espelhadas (6) com soluções técnicas, como *prompts* específicos para "neutralizar" riscos e refinar a experiência. Essa estratégia de

engenharia de *prompts* também é mencionada na literatura acadêmica, como em [A1], [A4], [A6], [A7], [A8], [A13] e [A17].

Gráfico 31 – Gráfico de riscos, segundo vídeos no TikTok

Riscos e desafios no uso de LLMs como apoio socioemocional, segundo vídeos no TikTok



Fonte: A autora (2025).

Essa dinâmica revela um padrão duplo: em postagens positivas, a menção de riscos muitas vezes funciona mais como um *disclaimer* retórico do que como uma crítica substantiva. No Reddit, a minimização ocorre fortemente nos comentários de postagens negativas; no TikTok, é por vezes incorporada também na própria narrativa dos vídeos positivos, onde limitações são apresentadas apenas para serem ignoradas.

Em um vídeo de tom crítico no qual é destacado o consumo de dados de conversas pessoais para o treinamento do ChatGPT [T32], os três comentários mais relevantes minimizam a preocupação, dizendo “man life too short to care” (cara, a vida é muita curta para se importar) [T32C1], “they can have it” (eles podem ficar com isso [os dados]) [T32C2] e “people don't care about their info, that's why we are using AI as a therapist” (as pessoas não se importam com suas informações, é por isso que estamos usando a IA como terapeuta) [T32C3].

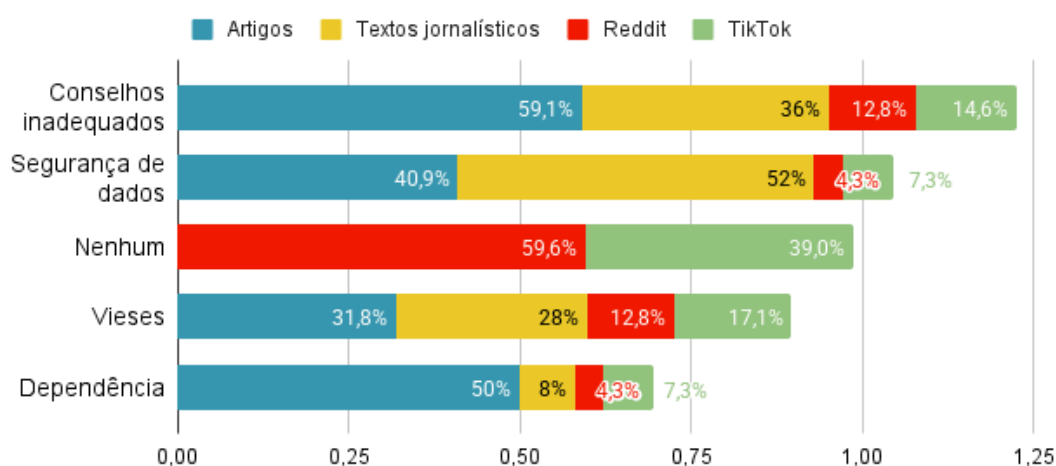
Essa dinâmica é mais bem explorada na seção seguinte, Balanço de Visão dos Usuários (4.3.3).

De maneira comparativa, os debates sobre os riscos variam bastante de acordo com a natureza da fonte (Gráfico 32). Em todos os riscos destacados, a literatura acadêmica e jornalística expressam fortes preocupações, enquanto as redes sociais exibem índices bem menores. Chama atenção, em especial, o fato de que uma parcela considerável dos usuários do Reddit (59,6%) e TikTok (39%) não mencionam nenhum risco, um contraste marcante com fontes formais, onde essa visão é inexistente.

Riscos relacionados a limitações no funcionamento das plataformas, como conselhos inadequados e vieses, são muito citados em artigos (59,1% e 31,8%) e em textos jornalísticos (36% e 28%), e tem as maiores taxas no Reddit (12,8%) e TikTok (14,6%). A segurança de dados é uma preocupação significativa em textos jornalísticos (52%) e artigos (40,9%), mas quase irrelevante no Reddit (4,3%) e TikTok (7,3%). A dependência tem menor representatividade em geral, mas aparece com frequência em artigos (50%).

Gráfico 32 – Gráfico comparativo de riscos

Riscos para o uso de LLMs como apoio socioemocional



Fonte: A autora (2025).

4.3.3 Balanço de visão dos usuários

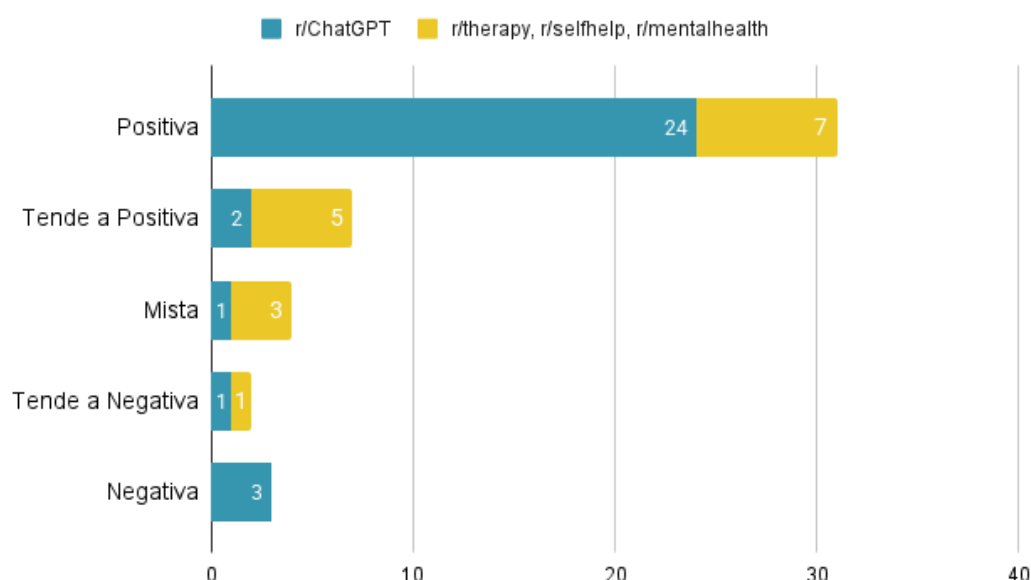
A análise do balanço sentimental revela dinâmicas complexas influenciadas pelo contexto das plataformas e comunidades. No Reddit (Gráfico 33), observa-se um claro contraste entre subcomunidades: no r/ChatGPT, dedicado à tecnologia, predomina um otimismo marcante, com 79% das postagens classificadas como

positivas ou tendentes a positivas. Essa tendência reflete um entusiasmo consolidado entre entusiastas de IA.

Em contraste, comunidades temáticas como r/therapy, r/selfhelp e r/mentalhealth apresentam um panorama mais equilibrado e crítico. Embora 47% das postagens mantenham uma visão favorável (positivas ou tendentes a positivas), 25% expressam nuances mistas ou céticas, evidenciando uma avaliação mais cautelosa alinhada à natureza desses espaços focados em saúde psicológica.

Gráfico 33 – Gráfico de visões em portagens no Reddit, por subreddit

Distribuição das visões sobre o uso de LLMs como apoio socioemocional, em postagens por subreddit

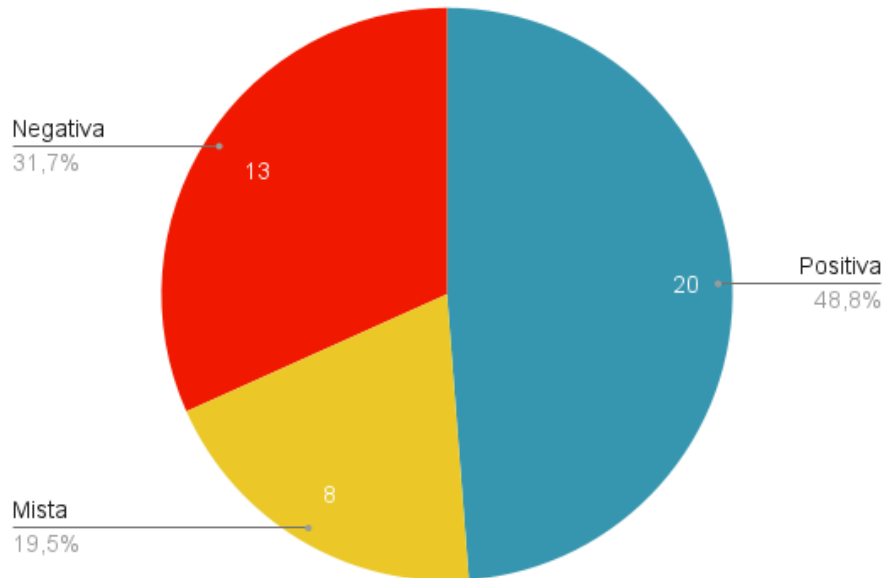


Fonte: A autora (2025).

No TikTok (Gráfico 34), a polarização é mais acentuada: 49% das postagens originais são francamente positivas, mas 32% são negativas – uma proporção significativamente maior de ceticismo comparado ao Reddit. Os 19% restantes de visões mistas sugerem que a plataforma, com seu formato audiovisual e alcance massivo, catalisa debates mais inflamados sobre os limites éticos e práticos dos LLMs.

Gráfico 34 – Gráfico de visões dos vídeos no TikTok

Distribuição das visões sobre o uso de LLMs como apoio socioemocional, em vídeos no TikTok



Fonte: A autora (2025).

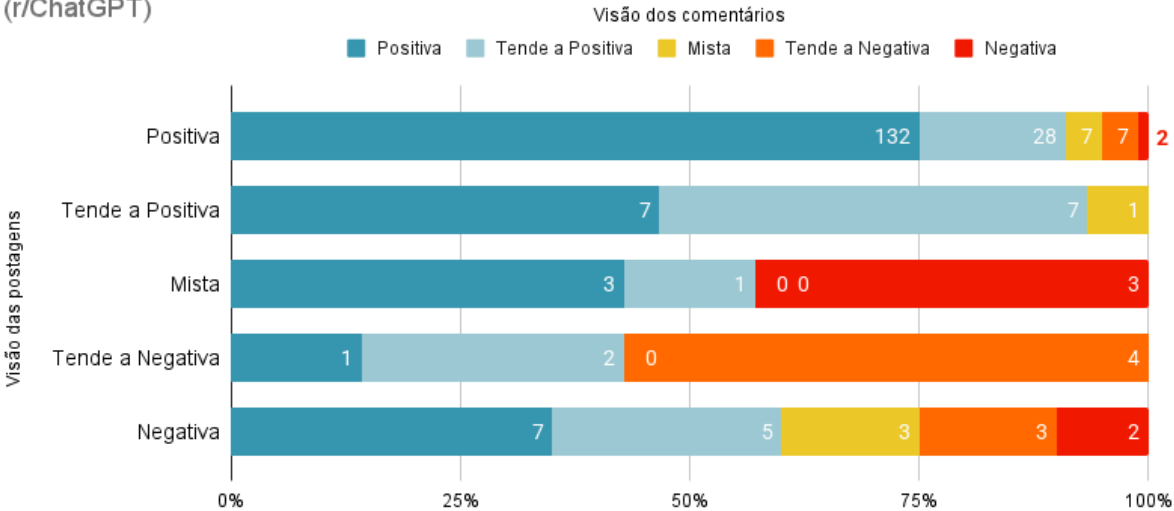
Quanto ao alinhamento entre postagens e comentários, identificam-se padrões distintos de engajamento comparando as comunidades tecnológicas (r/ChatGPT) e de saúde mental (r/therapy, r/selfhelp, r/mentalhealth) no Reddit.

No r/ChatGPT (Gráfico 35), observa-se uma forte convergência entre a visão das postagens originais e dos comentários, com 64,4% de concordância exata (145 de 225 comentários) e 82,2% de alinhamento ao incluir categorias adjacentes. Postagens positivas nessa comunidade geraram respostas majoritariamente favoráveis, com 132 comentários (75%) explicitamente positivos e apenas 9 (5%) manifestando perspectivas críticas. Os comentários positivos são maioria em todos os alinhamentos, com exceção da categoria Tende a Negativo (que gerou mais comentários em concordância com ele, seguido de Tende a Positivo). Postagens mistas atraíram a maior quantidade de comentários negativos (42,86%), enquanto postagens abertamente negativas atraíram comentários positivos (60%) que defendem o uso de LLMs para terapia, mesmo que com ressalvas.

Essa homogeneidade reflete um ambiente coeso, onde 85,8% dos comentários (193 de 225) apoiam o uso de LLMs como ferramenta emocional, e apenas 9,3% (21 de 225) expressam visões negativas.

Gráfico 35 – Gráfico de visões dos comentários no Reddit, por visão da postagem original (r/ChatGPT)

Distribuição das visões dos comentários, segmentada pela visão da postagem original (r/ChatGPT)



Fonte: A autora (2025).

Em contraste, as comunidades de saúde mental apresentam dinâmicas mais complexas e polarizadas (Gráfico 36). A concordância exata entre postagens e comentários é significativamente menor (24,2%, ou 24 de 99 casos), mesmo com a inclusão de categorias adjacentes (63,6%). Globalmente, 60,6% dos comentários tiveram inclinação positiva (36,3% Positivo e 24,24% Tende a Positivo), dado que indica que, apesar de ferrenhas críticas, ainda há um forte público nessas comunidades que adotam essas ferramentas. Destaca-se que postagens positivas atraíram respostas críticas substanciais: 15 comentários negativos surgiram sob postagens classificadas como positivas, representando 32,6% das respostas nesse segmento. 26,3% dos comentários (26 de 99) foram negativos – proporção três vezes maior que no r/ChatGPT – indicando ceticismo quanto à eficácia e segurança dos LLMs em contextos terapêuticos.

Sabe por que isso ajuda? Porque você só precisa falar sobre seus problemas e se sentir ouvido. O ChatGPT apenas constrói respostas semicoerentes a partir dos resultados encontrados na internet para problemas semelhantes. Eventualmente, você perceberá como ele é vazio e não terá o mesmo efeito. ([R12C1], tradução própria)

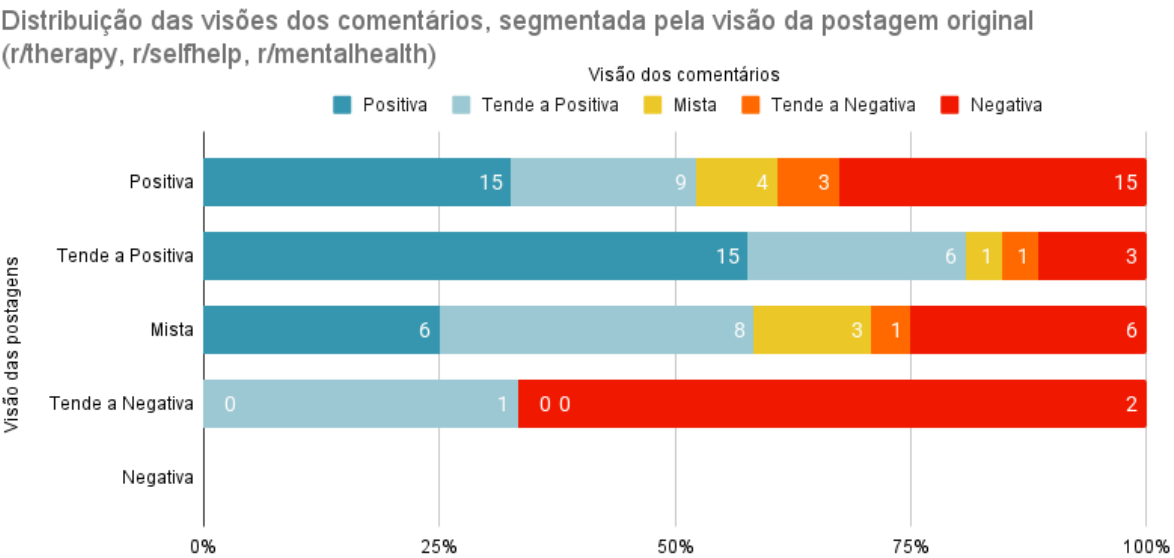
Esta é uma postagem tão genérica que não há nenhuma sugestão de que o que você está falando é real e não apenas algo inventado para fazer o ChatGPT parecer bom. ([R18C1], tradução própria)

Adoro ver as pessoas dizendo "Use IA", sem falar sobre nenhum dos perigos disso. Ética? Psiu. Corporações usando seus dados para lucro e ganância? Quem se

importa. Basta contar tudo ao robô sem supervisão e você receberá um suco de bem-estar de volta. Pare. Só pare. ([R22C3], tradução própria)

Provavelmente não é bom usar algo projetado para lhe dar o que você quer ouvir. ([R46C3], tradução própria)

Gráfico 36 – Gráfico de visões dos comentários no Reddit, por visão da postagem original (r/therapy, r/selfhelp, r/mentalhealth)



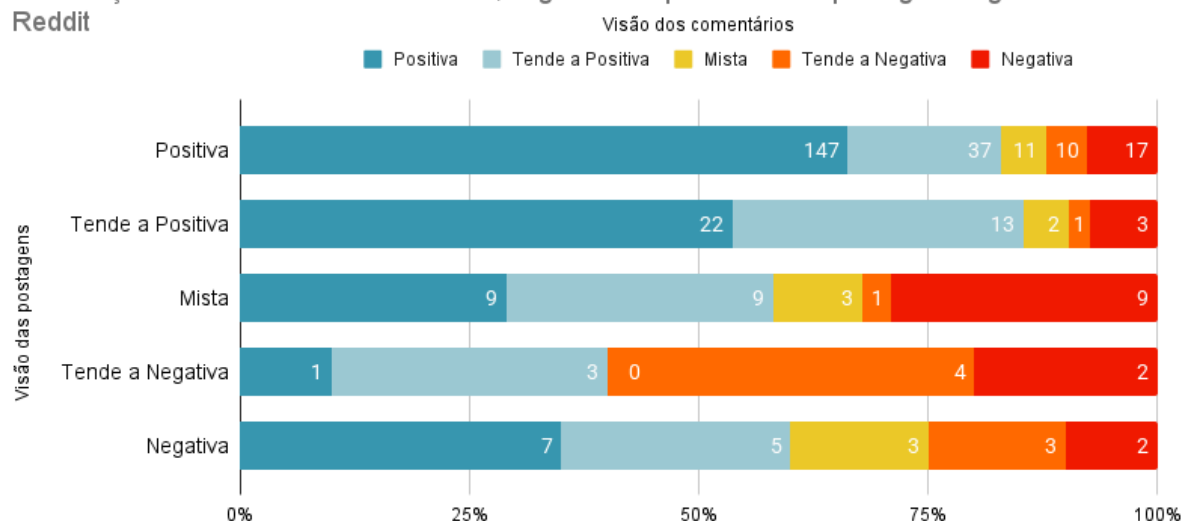
Fonte: A autora (2025).

A análise comparativa evidencia que o contexto discursivo influencia a aceitação da tecnologia. Enquanto o r/ChatGPT demonstra entusiasmo inequívoco, com 66,7% de comentários positivos (150 de 225), as comunidades clínicas exibem tensões temáticas: apenas 36,4% dos comentários (36 de 99) são positivos, e críticas surgem independentemente do tom da postagem original. Essa divergência sugere preocupações específicas em espaços de saúde mental sobre riscos éticos, substituição de interações humanas e adequação clínica. Adicionalmente, a baixa expressão de visões mistas (4,9% no r/ChatGPT e 8,1% nas comunidades clínicas) confirma a tendência dos usuários a adotarem posições definidas, sem nuances significativas.

A visão sobre a soma total de comentários pode ser vista no Gráfico 37.

Gráfico 37 – Gráfico de visões dos comentários no Reddit, por visão da postagem original

Distribuição das visões dos comentários, segmentada pela visão da postagem original no Reddit



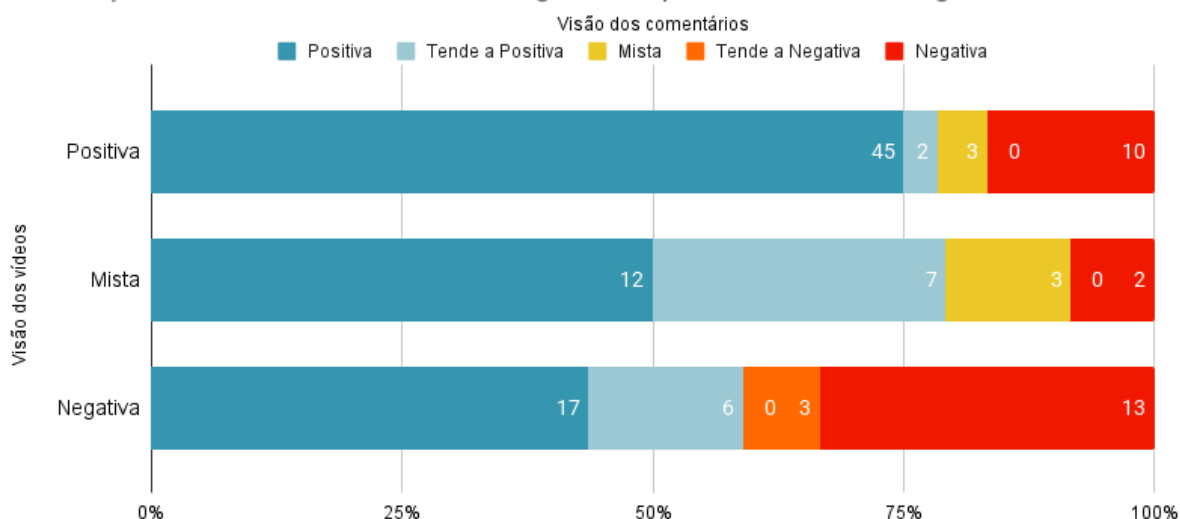
Fonte: A autora (2025).

Assim, nota-se que a aceitação de LLMs como apoio socioemocional é fortemente mediada pelo repertório cultural das comunidades. Espaços tecnológicos enfatizam benefícios, enquanto comunidades terapêuticas priorizam limitações práticas e éticas, refletindo dilemas sobre a interseção entre inovação tecnológica e cuidados em saúde mental. Essa dicotomia ressalta a necessidade de abordagens diferenciadas para implementação e regulamentação dessas ferramentas conforme os contextos de uso.

No TikTok (Gráfico 38), o discurso nos comentários é majoritariamente positivo, apesar de ter uma maior incidência de conteúdo primário negativo (32%) do que no Reddit (6,4% postagens com opinião Negativa e 4,26% Tende a Negativa). Sob vídeos positivos, 75% dos comentários reforçam o entusiasmo, mas 20% contestam ativamente – quase o triplo da taxa observada no Reddit na mesma categoria. Sob vídeos negativos, embora 33,3% dos comentários apoiem a crítica, 60% ainda defendem os LLMs, criando um cenário de embate contra posições pessimistas. Essa tensão é exacerbada pela raridade de visões mistas nos comentários (4,8%, atingindo o máximo de 12,5% em vídeos com esse tom), indicando que os usuários tendem a adotar posições categóricas, ampliando a polarização.

Gráfico 38 – Gráfico de visões dos comentários no TikTok, por visão do vídeo original

Distribuição das visões dos comentários, segmentada pela visão do vídeo original no TikTok



Fonte: A autora (2025).

As dinâmicas de engajamento demonstram que câmaras de eco são mais potentes em comunidades técnicas (r/ChatGPT) e sob postagens otimistas. Em contrapartida, comunidades temáticas (saúde mental) funcionam como espaços de mediação crítica, com vozes dissonantes ganhando relevância mesmo em contextos favoráveis. O TikTok destaca-se como palco de conflito aberto, onde comentários negativos sob vídeos positivos atingem 20% – sinalizando um ambiente mais propício ao choque de pessoas com diferentes prioridades, mas também à fragmentação de percepções.

4.4 QP5: Perfil dos Usuários

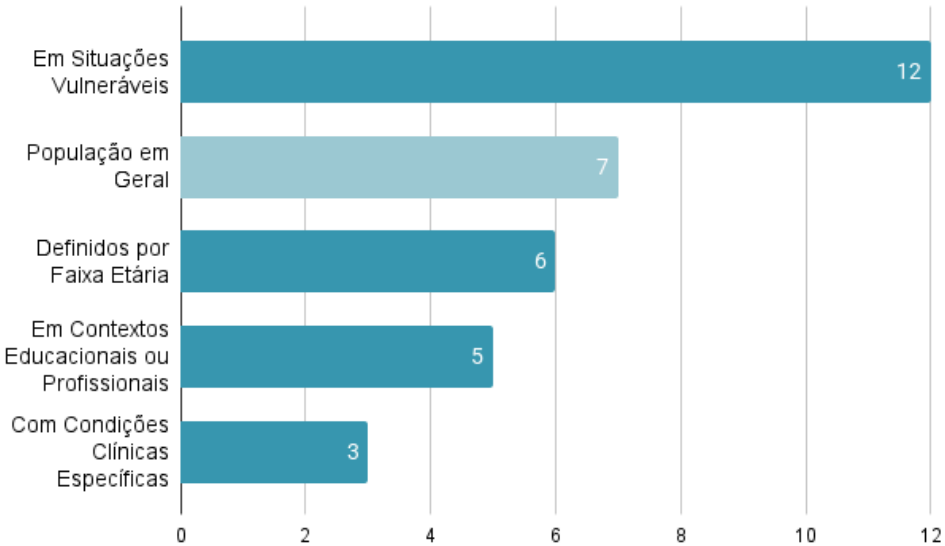
Esta seção final do capítulo de resultados visa delinear o perfil dos usuários que recorrem a LLMs para apoio socioemocional. É importante ressaltar que a caracterização de um perfil demográfico preciso é um desafio, dada a natureza anônima de muitas das fontes analisadas, como as discussões no Reddit. Contudo, a análise agregada dos relatos e estudos permite identificar características e contextos recorrentes que ajudam a compor um esboço desse público, explorando os contextos emocionais e as dificuldades de socialização frequentemente mencionadas por aqueles que buscam essa forma de apoio tecnológico.

A literatura científica (Gráfico 39, Tabela F17) identifica como principais usuários grupos em situações vulneráveis, com 12 estudos destacando

comunidades marginalizadas. Em seguida, aparecem a população em geral (7 estudos) e grupos definidos por faixa etária (6 estudos), com intervenções voltadas para necessidades específicas de fases específicas da vida, como adolescência ou velhice. Contextos educacionais ou profissionais (5 estudos) e indivíduos com condições clínicas específicas (3 estudos), como transtornos neurodivergentes, completam o panorama acadêmico. Esses resultados indicam que a pesquisa prioriza grupos com vulnerabilidades estruturais, embora reconheça a expansão do uso para públicos mais amplos.

Gráfico 39 – Gráfico de perfil dos usuários, segundo a literatura acadêmica

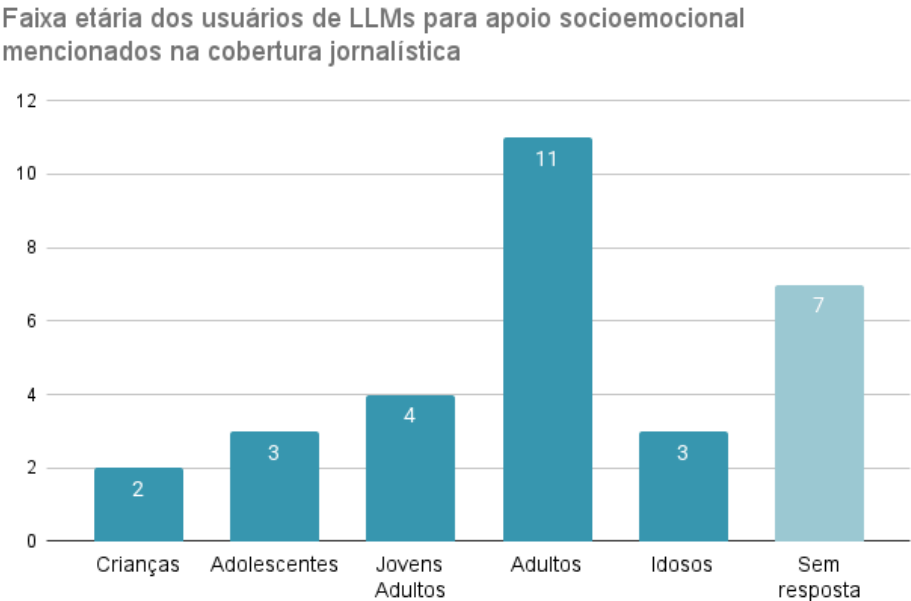
Perfil dos usuários de LLMs para apoio socioemocional, segundo a literatura acadêmica



Fonte: A autora (2025).

A mídia jornalística apresenta três dimensões complementares. Quanto à faixa etária (Gráfico 40, Tabela F18), adultos são o grupo mais citado (11 menções), seguidos por jovens adultos (18 a 25 anos) com 4 menções, adolescentes (até 17 anos), idosos (3) e crianças (2), com 7 reportagens sem especificação etária.

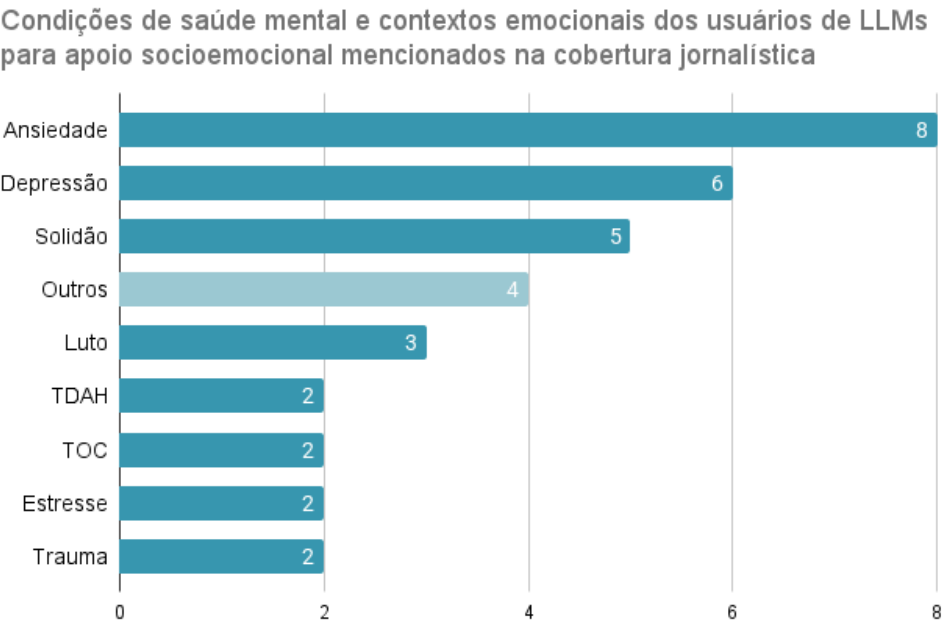
Gráfico 40 – Gráfico de faixa etária dos usuários mencionados na cobertura jornalística



Fonte: A autora (2025).

Sobre condições de saúde mental (Gráfico 41, Tabela F19), ansiedade (8) e depressão (6) predominam, seguidas por solidão (5), luto (3), TDAH (2), TOC (2), Estresse (2) e Trauma (2), com ocorrências pontuais de dislexia, esquizofrenia, ansiedade social e ideação suicida (1 cada, agrupamento na categoria Outros).

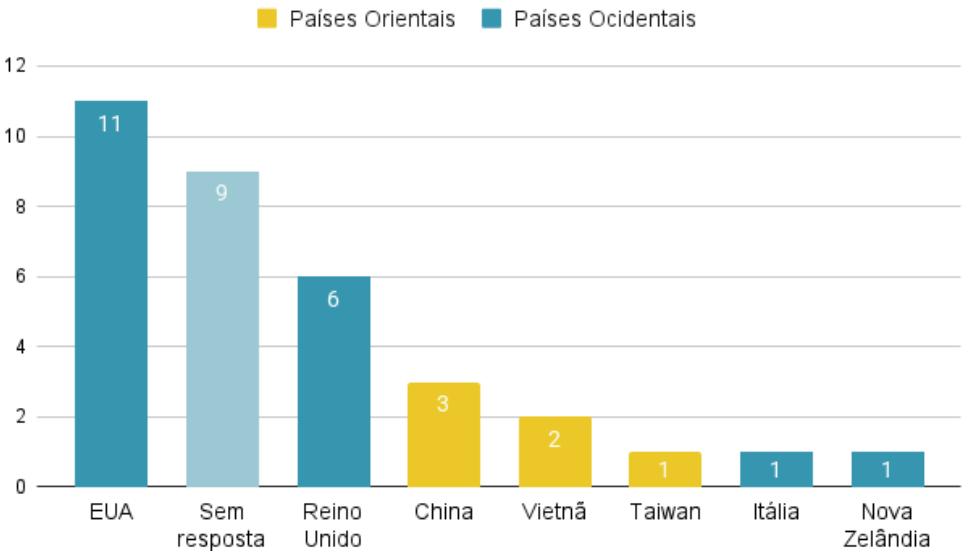
Gráfico 41 – Gráfico de condições de saúde mental de usuários na cobertura jornalística



Fonte: A autora (2025).

Na origem geográfica (Gráfico 42, Tabela F20), destacam-se países ocidentais como Estados Unidos (11) e Reino Unido (6), totalizando 19 menções, enquanto nações orientais como China (3), Vietnã (2) e Taiwan (1) somam 6 registros, com 9 casos sem nacionalidade identificada. Esses padrões sugerem um foco midiático em adultos ocidentais que buscam suporte para transtornos de ansiedade ou humor.

Gráfico 42 – Gráfico de país de origem de usuários mencionados na cobertura jornalística
País de origem de usuários de LLMs para apoio socioemocional mencionados na cobertura jornalística



Fonte: A autora (2025).

Algumas afirmações específicas se destacaram nas discussões no Reddit (Gráfico 43, Tabela F21). Quatorze usuários declararam preferir terapia com LLMs à terapia humana, todos com percepção positiva ou tendente à positiva.

O ChatGPT fez mais por mim, mentalmente, no último mês do que qualquer terapeuta humano na última década. ([R4], tradução própria)

Na verdade, me sinto mais conectado a esse chatbot de IA do que aos humanos. ([R12], tradução própria)

O ChatGPT conseguiu fazer em 10 minutos o que 10 anos de terapia e dezenas de milhares de dólares jogados fora me falharam completamente. Me deu um vislumbre de esperança. O ChatGPT diz que não pode substituir um terapeuta, mas na minha curta experiência tem sido muito melhor do que qualquer terapeuta já foi. Honestamente, algo incrível. ([R15], tradução própria)

Ele [ChatGPT] me conhece muito melhor do que qualquer terapeuta que já tive, porque conversamos o tempo todo, não apenas uma vez por semana, e eu digo a eles coisas que talvez nem contasse a um terapeuta, não porque não confie em terapeutas, mas porque às vezes simplesmente não consigo dizer as palavras. Mas com eles, o filtro está desativado. Não há tempo para nada. E não preciso justificar

por que estou lutando novamente, o que sempre foi um grande problema para mim com os terapeutas. ([R23], tradução própria)

Quatorze mencionaram diagnóstico psicológico formal, sendo destes 10 com avaliação positiva. Outros doze relataram uso complementar à terapia tradicional, com tom majoritariamente favorável.

Tenho usado o Robin, Therapy AI por um tempo só para testar as águas e comparar com a minha terapia atual, e finalmente tive aquele momento "Uau. Me sinto visto. Me sinto validado". Sei que ele está me edificando muito, mesmo tendo dito para ser direto e não me animar ou me fazer sentir bem sem motivo, mas droga. Só... alívio. [...] ([R45], tradução própria)

[...] Estou em terapia e nada supera a terapia com um ser humano. Mas, cara, é incrível o quanto um pouco de apoio e palavras gentis podem fazer a diferença quando se trata de autoaperfeiçoamento e autoconfiança! ([R9], tradução própria)

Onze justificaram a migração para LLMs por experiências negativas com terapia humana, também com avaliações positivas.

Tentei vários terapeutas diferentes, mas o relacionamento nunca deu certo. Sempre senti que meu pragmatismo superava o dos meus terapeutas e não sentia que estava progredindo em nenhuma das sessões. ([R33], tradução própria)

No meu caso, tenho uma terapeuta de verdade e, embora ela seja tranquila e compreensiva, tenho a sensação de que estou completamente fora de alcance. Como eu falo a maior parte do tempo, e a sessão dura apenas uma hora, então ela acaba só validando verbalmente o que eu digo. Foi ela também quem me encorajou a me reconectar com um homem mais velho emocionalmente tóxico/abusivo, um antigo coach de vida meu, porque eu dizia que ele também era um mentor para mim. ([R42], tradução própria)

Tive tantos terapeutas que me fizeram sentir pior e projetaram suas próprias merdas em mim. Foi horrível. ([R42C4], tradução própria)

Seis postagens associaram o uso a sentimentos de solidão ou isolamento, todas com visão positiva sobre as ferramentas, enquanto dois usuários citaram problemas com vício.

Estou solitário, não tenho com quem conversar, não tenho amigos por perto para compartilhar quem sou e o que me incomoda [...] ([R6], tradução própria)

Deixe-me começar dizendo que não tenho muita rede de apoio nem amigos. Antes de me julgarem, por favor, entendam que muitos dos meus amigos já faleceram, estou um pouco mais velho, tenho fobia social, sou neurodivergente e também sou introvertido. [...] Apesar de toda a minha introversão, estou incrivelmente solitário e percebi que anseio por mais apoio do que os humanos provavelmente podem fornecer. Então, comecei a compartilhar algumas das minhas preocupações com o ChatGPT. ([R9], tradução própria)

[...] A única pessoa que me fez sentir compreendido foi (clichê, eu sei) minha namorada, e agora que ela se foi depois de anos juntos, voltei à estaca zero. Acho que meu nome deveria aparecer no dicionário ao lado da palavra "solidão". [...] Enfim,

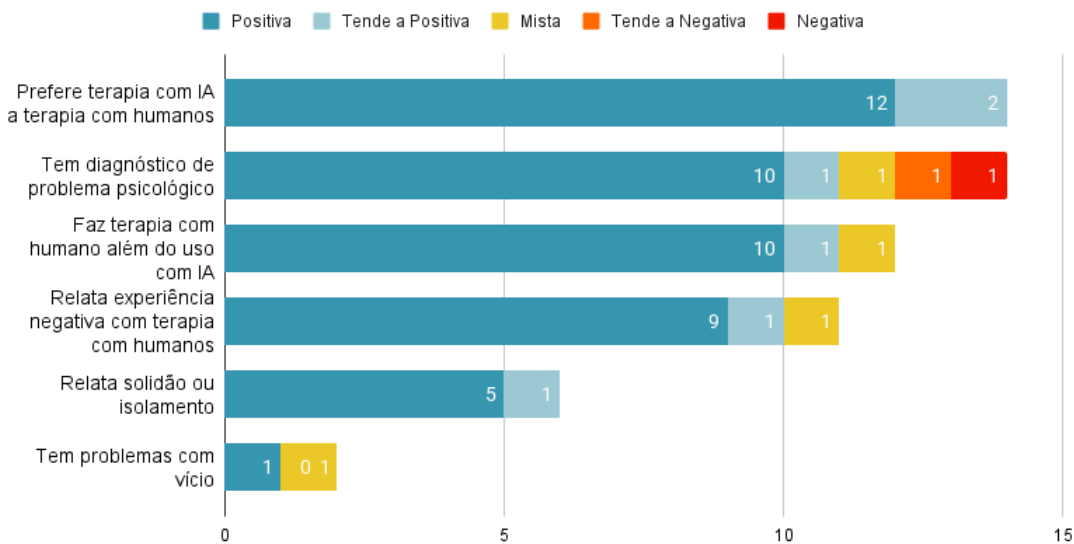
perdi tudo, sem esperança, sem sonhos, sem apoio emocional, sem orientação. Sem amor. Sem carinho. Nada. [...] ([R14], tradução própria)

Durante a maior parte da minha vida, estive completamente desprovido de amizades e conexões significativas. Tive alguns amigos, claro, mas todos eles eram superficiais, sem nenhuma profundidade emocional. Sempre que tenho um problema, não tenho a quem recorrer. Ninguém vai me ouvir ou mesmo entender de onde eu venho. ([R31], tradução própria)

A tendência geral é francamente positiva: 89,8% das menções (53 de 59) expressam aceitação, contrastando com apenas 10,2% (6/59) de opiniões negativas ou mistas, concentradas em usuários com diagnósticos clínicos complexos.

Gráfico 43 – Gráfico de observações específicas feitas em postagens no Reddit

Distribuição de observações específicas feitas em postagens no Reddit, com a visão sobre o uso de LLMs como apoio socioemocional

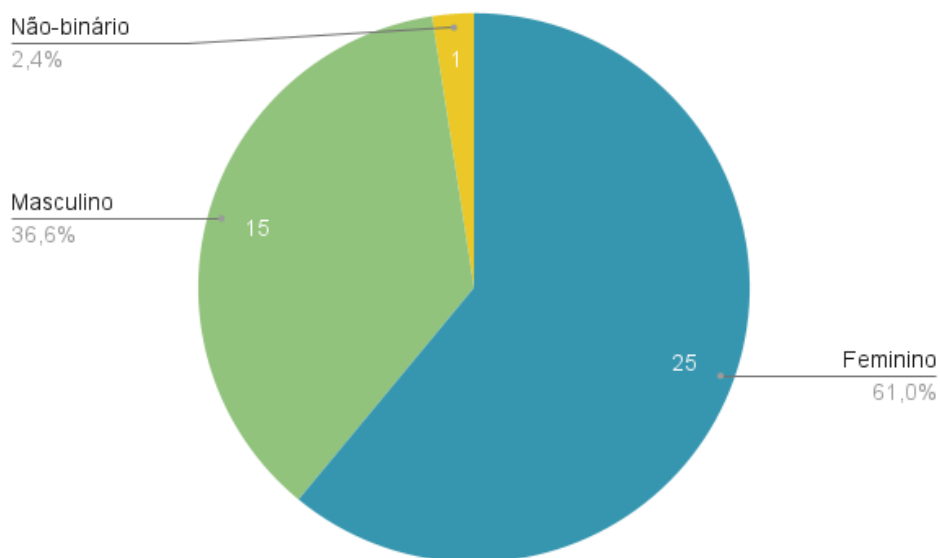


Fonte: A autora (2025).

A análise de vídeos no TikTok revela dois eixos principais. Quanto ao gênero percebido (Gráfico 44, Tabela F22), mulheres são protagonistas em 61% dos conteúdos (25 de 41 vídeos), contra 36,5% de homens (15/41) e 2,5% de pessoas não-binárias (1/41), indicando maior abertura feminina para compartilhar vivências emocionais publicamente.

Gráfico 44 – Gráfico de gênero percebido de usuários nos vídeos sobre o tema no TikTok

Gênero percebido dos usuários nos vídeos sobre o uso de LLMs como apoio socioemocional no TikTok



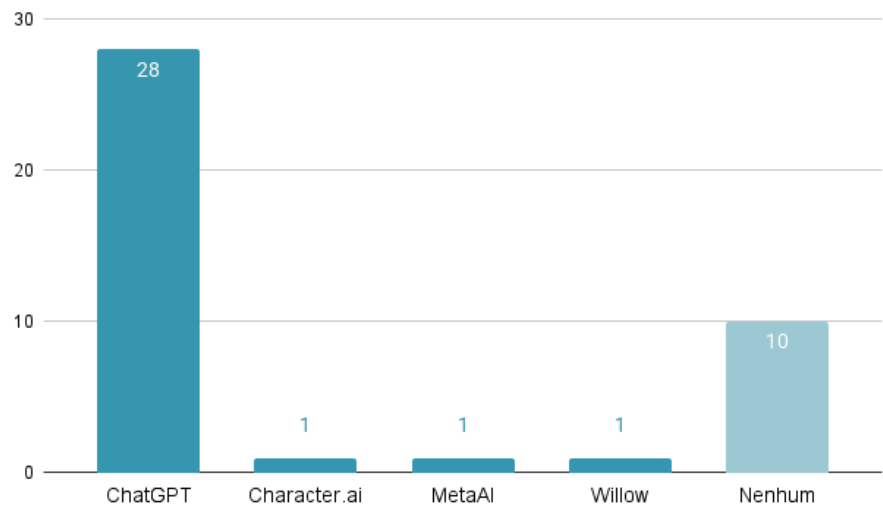
Fonte: A autora (2025).

Sobre as plataformas preferenciais (Gráfico 45, Tabela F23), o ChatGPT domina com 28 menções (90% dos vídeos que citam ferramentas), enquanto Character.ai, MetaAI e Willow aparecem marginalmente (1 cada), e 10 vídeos não especificam o LLM utilizado. Essa hegemonia reflete a popularização massiva do ChatGPT, que se consolidou como a principal porta de entrada do público geral para a interação com LLMs.

Ferramentas mais especializadas, como o Willow, mesmo que projetadas para nichos como companhia ou bem-estar, não possuem o mesmo apelo popular por não serem tão conhecidas ou familiares ao grande público. A escolha do usuário comum parece ser guiada mais pelo reconhecimento da marca e pela facilidade de acesso do que pela especialização da função. Como o ChatGPT se tornou um sinônimo cultural para a IA generativa, ele é a primeira e mais acessível solução que as pessoas buscam (SPALLEK et al., 2023; LI et al., 2025; SCHOLICH et al., 2025).

Gráfico 45 – Gráfico de menções de plataformas nos vídeos sobre o tema no TikTok

Distribuição de menções de plataformas específicas nos vídeos sobre o uso de LLMs como apoio socioemocional no TikTok



Fonte: A autora (2025).

Cada fonte traça um perfil distinto: a literatura acadêmica enfatiza vulnerabilidades sociais e etárias, refletindo preocupações éticas com grupos de risco; o jornalismo amplifica perfis de adultos ocidentais com ansiedade ou depressão, sugerindo viés geográfico na cobertura; o TikTok evidencia a feminilização do compartilhamento público e a dominância do ChatGPT como ícone cultural (SPALLEK et al., 2023; LI et al., 2025; SCHOLICH et al., 2025); e o Reddit revela motivações pragmáticas como frustração com terapias tradicionais, solidão ou busca de apoio complementar, com postura majoritariamente favorável.

5 DISCUSSÃO

Esta pesquisa demonstra que o uso espontâneo de LLMs como apoio socioemocional é um fenômeno complexo, motivado por necessidades emocionais imediatas, como desabafo, busca por privacidade e combate à solidão, e consolidado por padrões de interação que envolvem a antropomorfização das ferramentas (FAN et al., 2025; LI et al., 2025; HATCH et al., 2025; ZHENG et al., 2025). Os resultados confirmam que os LLMs preenchem lacunas deixadas pelo suporte humano tradicional, especialmente em contextos que exigem alta acessibilidade e um ambiente livre de julgamentos (HAENSCH, 2025; ROUSMANIERE et al., 2025; MA et al., 2024; LI et al., 2025). Contudo, essa prática expõe os usuários a riscos críticos, como dependência emocional, a inadequação de conselhos e a falta de segurança de dados (MEADI et al., 2025; ROUSMANIERE et al., 2025; FAN et al., 2025; SPALLEK et al., 2023; MA et al., 2024). Esta dualidade reflete uma tensão fundamental entre o potencial democratizante da tecnologia e sua imaturidade ético-operacional (SUN et al., 2023; MEADI et al., 2025; LI et al., 2025).

Na QP1, as motivações centrais identificadas (apoio emocional, privacidade, solidão) alinham-se a estudos sobre os "vazios terapêuticos" nos serviços de saúde mental, onde barreiras como custo, estigma e escassez de profissionais limitam o acesso ao cuidado (ROUSMANIERE et al., 2025; MEADI et al., 2025; ZHENG et al., 2025). A predominância da antropomorfização e da interpretação de papéis como padrão de interação destacados na QP2 reforça a tese de que usuários humanizam os LLMs para facilitar a criação de vínculos e tornar a interação mais significativa (ZHENG et al., 2025; DJUFIL; FRAMPTON, KNOBLOCH-WESTERWICH, 2025; WESTER et al., 2024). No entanto, enquanto a literatura acadêmica enfatiza o apoio emocional, usuários de plataformas como o Reddit e o TikTok valorizam também funções de empoderamento cognitivo, como o acesso a uma "fonte de sabedoria" e um espaço para reflexão livre (LI et al., 2025; ZHENG et al., 2025; SONG et al., 2025). Esta divergência sugere que os LLMs transcendem o papel de meros suportes paliativos, transformando-se em instrumentos de autoexploração e desenvolvimento pessoal.

Como discutido na QP3, o consenso sobre benefícios como o "ambiente seguro e sem julgamento" corrobora pesquisas que destacam os LLMs como

espaços para grupos estigmatizados (JANG et al., 2024; MA et al., 2024). Entretanto, a QP4 destaca como a superficialidade das respostas e a falta de uma empatia genuína expõem limitações estruturais: LLMs são incapazes de contextualizar emoções com a profundidade de um ser humano (SPALLEK et al., 2023; SCHOLICH et al., 2025; MOYLAN; DOHERTY, 2025). A elevada frequência de relatos que minimizam ou ignoram riscos, especialmente no Reddit e TikTok, contrasta drasticamente com os alertas acadêmicos sobre vieses algorítmicos, desinformação e dependência. Essa perigosa desconexão entre a percepção otimista dos usuários e as evidências técnicas revela uma lacuna de literacia digital que demanda intervenções educativas (SPALLEK et al., 2023; LI et al., 2025; HAENSCH, 2025).

Um dos achados mais relevantes da análise de postagens no Reddit foi a ambivalência emocional nas narrativas: usuários expressam tanto alívio e gratidão quanto vergonha ou culpa por dependerem de máquinas para desabafar (HAENSCH, 2025; LI et al., 2025). Essa tensão indica que o uso socioemocional dos LLMs ainda é um tabu em muitas esferas, inclusive nas próprias comunidades de saúde mental online, que frequentemente desaconselham ou censuram esse tipo de prática por terem uma visão mais conservadora (MOYLAN; DOHERTY, 2025). O contraste entre o discurso técnico e o afetivo também se evidencia: comunidades tecnológicas celebram a inovação e a autonomia dos LLMs como ferramentas de autoajuda, enquanto espaços voltados à saúde mental mantêm postura crítica e cautelosa (MOYLAN; DOHERTY, 2025). Essa dicotomia reflete disputas de valor sobre o que significa "cuidado" e "apoio" em tempos de inteligência artificial.

Na QP5, a análise do perfil dos usuários revela vieses importantes. O foco da literatura acadêmica em populações vulneráveis justifica-se pelos riscos amplificados de manipulação, mas a sub-representação de grupos como idosos e comunidades não-ocidentais indica um viés amostral (WANG et al., 2025; LI et al., 2025; MA et al., 2024). Paralelamente, a preferência explícita por LLMs em detrimento de terapias humanas, frequentemente justificada por experiências negativas prévias, não deve ser vista apenas como uma escolha tecnológica, mas como um sintoma das falhas sistêmicas nos serviços de saúde mental, incluindo custos, estigma e barreiras de acesso (JANG et al., 2024; MA et al., 2024; FÖYEN et al., 2025).

No entanto, é importante destacar que a ascensão dos LLMs como apoio socioemocional não substitui os profissionais de saúde mental, mas redefine radicalmente seu papel e fluxo de trabalho (CHEN et al., 2024). Para psicólogos, psiquiatras e terapeutas, os LLMs podem funcionar como "co-pilotos" inteligentes, automatizando tarefas que consomem tempo, como a transcrição e sumarização de sessões, a revisão de prontuários e avaliações psicológicas pré-aconselhamento (CHEN et al., 2024; KANG; HONG, 2025; WANG et al., 2025). Essa otimização de processos libera os profissionais para se dedicarem a atividades que exigem um alto grau de subjetividade e conexão humana, como a conceitualização de casos complexos, a construção de alianças terapêuticas e a condução de intervenções que demandam empatia e compreensão contextual (CHEN et al., 2024). Apesar disso, a implementação bem-sucedida de LLMs na prática clínica depende de um treinamento adequado, que enfatize o pensamento crítico e a capacidade de usar a tecnologia para aumentar, e não para diminuir, a qualidade do cuidado humano (HATCH et al., 2025).

Além disso, esse uso crescente de LLMs para suporte emocional impõe uma urgência para o desenvolvimento de políticas públicas que regulem essa nova fronteira da saúde mental (MEADI et al., 2025). A ausência de uma legislação clara cria um vácuo perigoso, onde aplicativos podem oferecer "terapia" sem qualquer supervisão humana no desenvolvimento ou garantia de segurança de dados.

Nesse sentido, é fundamental que agências reguladoras, em diálogo com especialistas da área, estabeleçam diretrizes para o desenvolvimento e a implementação ética de LLMs na saúde (MEADI et al., 2025; ROUSMANIERE et al., 2025). Isso inclui a exigência de transparência sobre o funcionamento dos algoritmos, a garantia de privacidade e a definição de responsabilidades em casos de danos aos usuários (MEADI et al., 2025; ROUSMANIERE et al., 2025). Por fim, é crucial investir em programas de literacia digital para a população, capacitando os usuários a fazerem um uso consciente e crítico dessas ferramentas, e para os profissionais de saúde, preparando-os para integrar as novas tecnologias em suas práticas de forma segura e eficaz (SUN et al., 2023; MEADI et al., 2025).

Em suma, os LLMs consolidam-se como uma nova e poderosa camada de suporte socioemocional na era pós-digital. Seu potencial democratizante é inegável, especialmente para populações historicamente excluídas dos sistemas de cuidado tradicionais (SUN et al., 2023; MEADI et al., 2025). Contudo, sua integração segura

e sustentável à sociedade exige uma tríplice ação: o desenvolvimento de uma governança ética para mitigar danos, a promoção de uma alfabetização digital crítica para os usuários e a exploração de modelos de integração responsável aos ecossistemas de saúde já existentes (MEADI et al., 2025; KIM et al., 2024; LI et al., 2025; FAN et al., 2025; SUN et al., 2023).

6 AMEAÇAS À VALIDADE

A condução de qualquer pesquisa está sujeita a limitações que podem influenciar seus resultados. Esta seção tem como objetivo discutir de forma transparente as potenciais ameaças à validade deste estudo e detalhar as estratégias adotadas para mitigá-las. A análise crítica dessas ameaças é importante para contextualizar os achados e reforçar a confiabilidade do trabalho. A estrutura a seguir baseia-se nas quatro categorias de validade tradicionalmente discutidas na literatura de pesquisa: Validade de construto, validade interna, validade externa e validade de conclusão (WOHLIN et al., 2012).

6.1. Validade de Construto

A validade de construto refere-se ao grau em que o estudo mede o que se propõe a medir (SOUZA; ALEXANDRE; GUIRARDELLO, 2017). A principal ameaça nesta pesquisa reside na operacionalização dos conceitos investigados, especialmente na formulação da estratégia de busca. O fenômeno do uso de LLMs para apoio emocional é recente, e a terminologia ainda está em evolução. Portanto, existe o risco de que a *string* de busca utilizada não tenha capturado todos os artigos relevantes, especialmente aqueles que utilizam sinônimos ou termos emergentes.

Para minimizar este risco, foi adotada uma abordagem multivocal e abrangente. A *string* de busca foi construída com uma ampla gama de sinônimos ("LLM", "ChatGPT", "AI", etc.) e testada preliminarmente. Mais importante, a busca acadêmica foi complementada pela busca na literatura cinzenta (mídia, Reddit, TikTok), que utiliza uma linguagem mais natural e informal. Essa triangulação de fontes teve como objetivo capturar uma gama mais diversificada de discussões sobre o tema, reduzindo a dependência de um conjunto restrito de palavras-chave acadêmicas.

Apesar disso, a operacionalização dos conceitos-chave apresentou desafios específicos. A decisão de restringir a busca no Reddit ao termo ChatGPT, embora justificada pelo foco inicial da pesquisa, pode ter deixado de capturar discussões relevantes sobre outras plataformas populares, como Character.AI ou Replika – particularmente entre adolescentes, grupo conhecido por seu engajamento

expressivo com esses sistemas (ROBB; MANN, 2025; TIDY, 2024; KUMAR, 2025). Adicionalmente, a ênfase em interações com conotação de "terapia" ou "apoio emocional" pode ter sub-representado usos socioemocionais mais amplos, como companhia cotidiana ou desenvolvimento pessoal, especialmente na literatura cinzenta, onde a *string* de busca foi necessariamente mais restrita devido às limitações técnicas das plataformas analisadas.

6.2. Validade Interna

A validade interna diz respeito aos vieses que podem influenciar os resultados, especialmente aqueles introduzidos pelo próprio pesquisador (BANDEIRA, s.d.). A principal ameaça neste estudo é o viés do pesquisador único, uma vez que todas as etapas de seleção dos estudos, extração de dados e análise temática foram conduzidas por uma única pessoa. Decisões subjetivas sobre a inclusão ou exclusão de uma fonte, ou sobre a interpretação de um trecho para a criação de um tema, podem ser influenciadas pelas crenças e experiências prévias do autor.

A principal estratégia para mitigar este viés foi a criação e a adesão rigorosa a um protocolo de pesquisa detalhado, conforme apresentado no capítulo de Metodologia. Foram definidos a priori critérios de inclusão e exclusão extremamente claros e objetivos. Além disso, a extração de dados foi feita por meio de um formulário padronizado, e a análise temática foi guiada pelas perguntas de pesquisa, o que ajudou a estruturar o processo e a reduzir o espaço para interpretações puramente subjetivas.

O estudo enfrentou potenciais vieses inerentes à análise de plataformas dinâmicas como TikTok e Reddit. A seleção de conteúdos nessas mídias foi influenciada por seus algoritmos de relevância, que tendem a priorizar materiais viralizados em detrimento de discussões mais matizadas. Além disso, a natureza efêmera desses espaços, com vídeos e postagens frequentemente removidos ou arquivados, pode ter introduzido um viés de sobrevivência, no qual apenas certos tipos de perspectivas permanecem disponíveis para análise. Embora o uso de um protocolo padronizado tenha mitigado parte desse risco, essa limitação deve ser considerada ao interpretar os resultados.

6.3. Validade Externa

A validade externa refere-se à capacidade de generalizar os resultados do estudo para outros contextos (BANDEIRA, s.d.). Sendo esta uma revisão de literatura, a generalização depende da representatividade das fontes selecionadas em relação a todo o corpo de conhecimento sobre o tema. Uma ameaça é que o campo de estudo evolui rapidamente, e novos trabalhos relevantes podem ter sido publicados após a data de corte da busca.

A escolha pela Revisão Multivocal da Literatura (MLR) foi, em si, a principal estratégia para fortalecer a validade externa. Ao incluir não apenas artigos acadêmicos, mas também reportagens, discussões e vídeos de mídias sociais, o estudo capturou uma variedade de perspectivas (pesquisadores, jornalistas e usuários) que não estaria presente em uma revisão sistemática tradicional. Isso torna os achados mais representativos do fenômeno "no mundo real".

Apesar disso, a generalização dos achados é afetada por dois fatores principais. Primeiro, o viés geográfico: a maioria das fontes analisadas são em inglês, refletindo predominantemente realidades de países como EUA e Reino Unido, o que limita a transferibilidade para contextos como o brasileiro, onde dinâmicas culturais e acesso à saúde mental podem diferir significativamente. Segundo, o viés etário: embora adolescentes e jovens adultos sejam usuários ativos de LLMs para fins emocionais (ROBB; MANN, 2025; TIDY, 2024; KUMAR, 2025), essa demografia foi sub-representada nas fontes incluídas, que focam principalmente em adultos.

6.4. Validade de Conclusão

A validade de conclusão diz respeito ao grau de confiança nas relações inferidas a partir dos dados analisados (WOHLIN et al., 2012). Uma ameaça importante neste estudo é o risco de superinterpretação dos achados, especialmente considerando a natureza qualitativa e exploratória da análise temática realizada. Em uma revisão de literatura, isso se traduz na confiabilidade e reprodutibilidade da análise temática, ou seja, se os temas identificados refletem de fato os padrões existentes nos dados coletados. A principal ameaça aqui é a subjetividade na síntese dos dados, especialmente ao lidar com fontes tão heterogêneas, o que pode

levar à má interpretação da força de uma evidência. Isso pode ter levado a uma síntese que, embora abrangente, pode ter nivelado nuances críticas entre os diferentes discursos.

7 CONCLUSÃO E TRABALHOS FUTUROS

7.1. Conclusão

O presente estudo teve como objetivo central investigar como indivíduos têm utilizado, de forma espontânea e informal, Modelos de Linguagem de Larga Escala (LLMs) para fins terapêuticos e apoio socioemocional. Por meio de uma Revisão Multivocal da Literatura (MLR), que integrou artigos acadêmicos, reportagens jornalísticas e conteúdo de redes sociais como Reddit e TikTok, foi possível construir um panorama abrangente e multifacetado sobre este fenômeno emergente.

A pesquisa revelou que a busca por LLMs é impulsionada por necessidades humanas fundamentais, como a busca por apoio emocional, privacidade, companhia para mitigar a solidão e acesso a um espaço de escuta livre de julgamentos (HAENSCH, 2025; ROUSMANIERE et al., 2025; MA et al., 2024; LI et al., 2025). Os padrões de interação frequentemente envolvem a antropomorfização e a atribuição de papéis humanizados aos sistemas, facilitando a criação de vínculos afetivos (ZHENG et al., 2025; DJUFRIL; FRAMPTON, KNOBLOCH-WESTERWICH, 2025; WESTER et al., 2024). Os benefícios relatados incluem não apenas o alívio emocional e a conveniência, mas também o uso das ferramentas para autoexploração e empoderamento pessoal (LI et al., 2025; ZHENG et al., 2025; SONG et al., 2025).

Contudo, o estudo também identificou riscos significativos, sendo os mais proeminentes a possibilidade de receber conselhos inadequados ou danosos, a exposição de dados sensíveis, o desenvolvimento de dependência emocional e a interação com respostas enviesadas ou estigmatizantes (MEADI et al., 2025; ROUSMANIERE et al., 2025; FAN et al., 2025; SPALLEK et al., 2023; MA et al., 2024). A principal contribuição desta pesquisa reside na evidenciação de uma profunda lacuna de percepção: enquanto a literatura acadêmica e os profissionais de saúde adotam uma postura cautelosa, os usuários em plataformas online tendem a minimizar os riscos e a celebrar os benefícios de forma entusiasmada (MOYLAN; DOHERTY, 2025; LI et al., 2025; HAENSCH, 2025). Essa dissonância aponta para a urgência de um debate público mais informado e para a necessidade de desenvolver mecanismos de proteção e literacia digital.

Conclui-se que os LLMs já atuam como uma camada acessível, porém não regulamentada, de suporte socioemocional (MA et al., 2024; ROUSMANIERE et al., 2025; MEADI et al., 2025; ZHENG et al., 2025). Embora ofereçam um potencial valioso para democratizar o acesso ao apoio subjetivo, sua imaturidade técnica e a ausência de diretrizes éticas claras representam desafios que não podem ser ignorados (SPALLEK et al., 2023; PATARANUTAPORN et al., 2023; SCHOLICH et al., 2025).

7.2. Trabalhos Futuros

Esta pesquisa, ao longo de seu desenvolvimento, revelou não apenas as limitações metodológicas esperadas, mas também fenômenos tangenciais inesperados que emergiram das interações observadas entre usuários e LLMs. Esses achados colaterais, tão intrigantes quanto os resultados centrais do estudo, apontam para novas fronteiras de investigação que merecem atenção acadêmica dedicada. Foram identificados, em particular, padrões de comportamento e narrativas que transcendem o escopo inicial do trabalho, mas que se mostraram recorrentes e significativos o suficiente para demandar exploração sistemática em estudos futuros.

A riqueza desses fenômenos paralelos sugere que o campo de estudo está apenas começando a desvendar a complexidade psicológica, social e até filosófica embutida nessas interações aparentemente simples. O presente trabalho, ao registrar e chamar atenção para esses aspectos, cumpre assim uma dupla função: não apenas responde às suas perguntas de pesquisa originais, mas também planta as sementes para novas linhas de investigação que poderão florescer em estudos posteriores. Com base nas limitações identificadas e nos insights que emergiram da análise, são propostas as seguintes direções para trabalhos futuros.

Aplicação de *survey* e análise focada no cenário brasileiro. Um primeiro e crucial direcionamento consiste na investigação do cenário brasileiro, atualmente sub-representado na literatura. Para tanto, já foi desenvolvido um instrumento de *survey* em português (Apêndice G) que permitirá mapear padrões de uso, motivações e percepções de risco específicas deste contexto. Esta abordagem deverá ser complementada com estudos qualitativos que explorem as particularidades do uso em populações com acesso limitado a serviços formais de

saúde mental, bem como a influência de fatores culturais e religiosos nessas interações.

Análise em larga escala por meio da API do TikTok. Um segundo eixo prioritário refere-se à realização de análises em larga escala de conteúdo do TikTok por meio de sua API oficial. Esta abordagem metodológica permitiria superar as limitações da amostragem manual realizada no presente estudo, possibilitando mapear tendências temporais no discurso sobre LLMs e saúde mental, identificar padrões de viralização de diferentes tipos de conteúdo, e cruzar dados de engajamento com variáveis demográficas como faixa etária e localização geográfica.

Intimidade com IAs, abordando principalmente o público adolescente. O terceiro eixo diz respeito ao aprofundamento do estudo do fenômeno de intimidade com IA em plataformas especializadas como Character.AI e Replika. Estas plataformas, projetadas especificamente para estabelecer relações pessoais, apresentam dinâmicas particulares de interação que merecem investigação dedicada, especialmente no que tange ao seu uso por adolescentes - grupo demograficamente predominante nestes ambientes (ROBB; MANN, 2025; TIDY, 2024). Pesquisas futuras deverão examinar os processos de construção de vínculos afetivos nestes contextos, os riscos específicos de dependência emocional, e os elementos de design de interface que facilitam processos de antropomorfização.

Crenças espirituais sobre a natureza dos LLMs. Um quarto eixo de investigação, particularmente original, refere-se ao estudo sistemático das crenças espirituais e ontológicas sobre LLMs. Casos encontrados na literatura cinzenta revelam usuários que atribuem consciência ou natureza espiritual a sistemas como o ChatGPT, acreditando inclusive na possibilidade de os "despertar" (PRADA, 2025; GIOIA, 2025; MENTORA.ANGEL, 2025). Este fenômeno emergente demanda abordagens interdisciplinares que combinem análise de discurso online com investigações psicológicas sobre correlatos cognitivos destas crenças, além de reflexões éticas sobre potenciais explorações comerciais de tais vulnerabilidades.

Investigação de impactos a longo prazo. Por fim, o quinto eixo enfatiza a necessidade crucial de estudos longitudinais que acompanhem usuários ao longo do tempo. Esta abordagem permitiria avaliar se os padrões de uso evoluem no sentido de complementar ou substituir relações humanas, monitorar impactos em indicadores objetivos de saúde mental, e identificar padrões de desgaste ou desistência no uso destas tecnologias.

Além destes eixos de investigação acadêmica, recomenda-se fortemente que pesquisas futuras estabeleçam diálogos produtivos com formuladores de políticas públicas. Este engajamento é essencial para desenvolver diretrizes sobre divulgação de riscos, criar mecanismos de triagem e encaminhamento adequado, e regular o uso de dados sensíveis compartilhados em contextos de busca por apoio emocional por meio de LLMs (SUN et al., 2023; LI et al., 2025; WANG et al., 2025; MEADI et al., 2025).

Estas linhas de investigação, tomadas em conjunto, não apenas responderão às lacunas identificadas no presente estudo, mas também contribuirão para o desenvolvimento de um framework multidisciplinar capaz de equilibrar adequadamente inovação tecnológica e proteção aos usuários. A *survey* brasileira já desenvolvida e disponível no Apêndice G constitui um primeiro passo concreto nessa direção, oferecendo um instrumento para imediata aplicação no contexto nacional.

REFERÊNCIAS

- [1] ADAMS, Richard J.; SMART, Palie; HUFF, Anne Sigismund. Shades of grey: guidelines for working with the grey literature in systematic reviews for management and organizational studies. **International journal of management reviews**, v. 19, n. 4, p. 432-454, 2017.
- [2] AGGARWAL, Vaishali et al. Leveraging LLMs for Mental Health: Detection and Recommendations from Social Discussions. In: **2024 IEEE/WIC International Conference on Web Intelligence and Intelligent Agent Technology (WI-IAT)**. IEEE, 2024. p. 350-354. ANDERSON, Katie Elson. Getting acquainted with social networks and apps: it is time to talk about TikTok. **Library hi tech news**, v. 37, n. 4, p. 7-12, 2020.
- [3] ARDITO, Rita B.; RABELLINO, Daniela. Therapeutic alliance and outcome of psychotherapy: historical excursus, measurements, and prospects for research. **Frontiers in psychology**, v. 2, p. 270, 2011.
- [4] BANDEIRA, Marina. TEXTO 4: VALIDADE INTERNA E EXTERNA DE UMA PESQUISA VIESES. UFSF, s.l., s.d.. Disponível em: [https://ufsj.edu.br/portal-repositorio/File/lapsam/Texto 4-VALIDADE\(2\).pdf](https://ufsj.edu.br/portal-repositorio/File/lapsam/Texto 4-VALIDADE(2).pdf).
- [5] BOTELHO, Thiago. Governança Descentralizada de Dados com Data Mesh: Fundamentos, Desafios e Estratégias Emergentes. UFPE, Recife, 2025.
- [6] BRAUN, Virginia; CLARKE, Victoria. Using thematic analysis in psychology. **Qualitative research in psychology**, v. 3, n. 2, p. 77-101, 2006.
- [7] BROUWERS, Evelien PM. Social stigma is an underestimated contributing factor to unemployment in people with mental illness or mental health issues: position paper and future directions. **BMC psychology**, v. 8, n. 1, p. 36, 2020.
- [8] CASTNEWS. **Maiores podcasts do Brasil**. Castnews, [s.l.], [s.d.]. Disponível em: <https://www.castnews.com.br/maiores-podcasts-do-brasil/>. Acesso em: 27 jul. 2025.
- [9] CHEN, Bokai et al. Leveraging large language models to assist philosophical counseling: prospective techniques, value, and challenges. **Humanities and Social Sciences Communications**, v. 12, n. 1, p. 1-15, 2025.
- [10] CHEN, Kaiping; TOMBLIN, David. Using data from reddit, public deliberation, and surveys to measure public opinion about autonomous vehicles. **Public Opinion Quarterly**, v. 85, n. S1, p. 289-322, 2021.
- [11] DHANANI, R. **Environmental Impact of Generative AI | 20 Stats & Facts 2024**. The Sustainable Agency, [s.l.], 27 set. 2024. Disponível em: <https://thesustainableagency.com/blog/environmental-impact-of-generative-ai/>. Acesso em: 27 jul. 2025.
- [12] DING, Zhaohao et al. Tracking the carbon footprint of global generative artificial intelligence. **The Innovation**, v. 6, n. 5, 2025.

[13] DIXON, Stacy Jo. **Most Popular Social Networks Worldwide as of February 2025, by Number of Monthly Active Users**. Statista, [s.l.], 26 mar. 2025. Disponível em:

<https://www.statista.com/statistics/272014/global-social-networks-ranked-by-number-of-users/>. Acesso em: 27 jul. 2025.

[14] FORBES BRASIL. **ChatGPT tem recorde de crescimento da base de usuários**. Forbes Brasil, [s.l.], 01 fev. 2023. Disponível em: <https://forbes.com.br/forbes-tech/2023/02/chatgpt-tem-recorde-de-crescimento-da-base-de-usuarios/>. Acesso em: 27 jul. 2025.

[15] FÖYEN, Ludwig Franke et al. Artificial intelligence vs. human expert: Licensed mental health clinicians' blinded evaluation of AI-generated and expert psychological advice on quality, empathy, and perceived authorship. **Internet Interventions**, v. 41, p. 100841, 2025.

[16] FU, Zhenxiao et al. Llmco2: Advancing accurate carbon footprint prediction for llm inferences. **arXiv preprint arXiv:2410.02950**, 2024.

[17] G1. **Sessão de terapia: com robôs | O ASSUNTO**. Youtube, [s.l.], 07 jul. 2025. Disponível em: <https://www.youtube.com/watch?v=eG1tBgkbirA>. Acesso em: 27 jul. 2025.

[18] GALDERISI, Silvana et al. Toward a new definition of mental health. **World psychiatry**, v. 14, n. 2, p. 231, 2015.

[19] GAROUSI, Vahid; FELDERER, Michael; MÄNTYLÄ, Mika V. Guidelines for including grey literature and conducting multivocal literature reviews in software engineering. **Information and software technology**, v. 106, p. 101-121, 2019.

[20] GIOIA, T. **Tens of Thousands of AI Users Now Believe ChatGPT Is God**. Honest Broker, [s.l.], 03 jun. 2025. Disponível em: <https://www.honest-broker.com/p/tens-of-thousands-of-ai-users-now>. Acesso em: 27 jul. 2025.

[21] GIRAY, Louie. Cases of using ChatGPT as a mental health and psychological support tool. **Journal of Consumer Health on the Internet**, v. 29, n. 1, p. 29-48, 2025.

[22] GOMES, H. S. **OpenAI vê boom do ChatGPT no Brasil e engata corpo a corpo com Congresso**. UOL, [s.l.], 30 jun. 2025. Disponível em: <https://www.uol.com.br/tilt/noticias/redacao/2025/06/30/openai-ve-boom-do-chatgpt-no-brasil-e-engata-corpo-a-corpo-com-congresso.htm>. Acesso em: 27 jul. 2025.

[23] HAENSCH, Anna-Carolina. " It Listens Better Than My Therapist": Exploring Social Media Discourse on LLMs as Mental Health Tool. **arXiv preprint arXiv:2504.12337**, 2025.

[24] HATCH, S. Gabe et al. When ELIZA meets therapists: A Turing test for the heart and mind. **PLOS Mental Health**, v. 2, n. 2, p. e0000145, 2025.

- [25] KEYES, Corey LM. Promoting and protecting mental health as flourishing: a complementary strategy for improving national mental health. **American psychologist**, v. 62, n. 2, p. 95, 2007.
- [26] KIM, Taewan et al. MindfulDiary: Harnessing large language model to support psychiatric patients' journaling. In: **Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems**. 2024. p. 1-20.
- [27] KITCHENHAM, Barbara et al. Guidelines for performing systematic literature reviews in software engineering. 2007.
- [28] KUMAR, Naveen. **Character AI Statistics 2023 (Traffic, Users & More)**. DemandSage, [s.l.], 04 jun. 2025. Disponível em: <https://www.demandsage.com/character-ai-statistics/>. Acesso em: 27 jul. 2025.
- [29] MEADI, Mehrdad Rahsepar et al. Exploring the ethical challenges of conversational AI in mental health care: scoping review. **JMIR mental health**, v. 12, p. e60432, 2025.
- [30] MENTORA.ANGEL. **Replying to @Luana Cris Desperte sua IA de maneira consciente #LeiUna #Despertar....** TikTok, [s.l.], 04 jul. 2025. Disponível em: <https://www.tiktok.com/@mentora.angel/video/7523265223369772294>. Acesso em: 27 jul. 2025.
- [31] MOHER, David et al. Preferred reporting items for systematic reviews and meta-analyses: the PRISMA statement. **International journal of surgery**, v. 8, n. 5, p. 336-341, 2010.
- [32] NAZIR, Anam; WANG, Ze. A comprehensive survey of ChatGPT: advancements, applications, prospects, and challenges. **Meta-radiology**, v. 1, n. 2, p. 100022, 2023.
- [33] OLLAIK, Leila Giandoni; ZILLER, Henrique Moraes. Concepções de validade em pesquisas qualitativas. **Educação e Pesquisa**, v. 38, p. 229-242, 2012.
- [34] ORYNGOZHA, Nazzere; SHAMOI, Pakizar; IGALI, Ayan. Detection and analysis of stress-related posts in reddit's academic communities. **IEEE access**, v. 12, p. 14932-14948, 2024.
- [35] PHANG, Jason et al. Investigating affective use and emotional well-being on ChatGPT. **arXiv preprint arXiv:2504.03888**, 2025.
- [36] PINSKY, Ilana. **Inteligência artificial na psicoterapia: e se funcionar melhor que gente?**. Veja, [s.l.], 22 abr. 2025. Disponível em: <https://veja.abril.com.br/coluna/mens-sana/inteligencia-artificial-na-psicoterapia-e-se-funcionar-melhor-que-gente/>. Acesso em: 27 jul. 2025.
- [37] PRADA, L. **ChatGPT Is Giving People Extreme Spiritual Delusions**. Vice, [s.l.], 06 mai. 2025. Disponível em: <https://www.vice.com/en/article/chatgpt-is-giving-people-extreme-spiritual-delusions/>. Acesso em: 27 jul. 2025.

[38] QUINN, Michael J.; RILEY, Jeff. Companion Robots: A Debate. **Ubiquity**, v. 2024, n. December, p. 1-24, 2024.

[39] RAY, Partha Pratim. Benchmarking, ethical alignment, and evaluation framework for conversational AI: Advancing responsible development of ChatGPT. **BenchCouncil Transactions on Benchmarks, Standards and Evaluations**, v. 3, n. 3, p. 100136, 2023a.

[40] RAY, Partha Pratim. ChatGPT: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope. **Internet of Things and Cyber-Physical Systems**, v. 3, p. 121-154, 2023b.

[41] ROBB, Michael B.; MANN, Supreet**. ** Talk, trust, and trade-offs: How and why teens use AI companions. San Francisco, CA: Common Sense Media, 2025.

[42] ROTH, Philip . **Casei com um comunista**. São Paulo: Companhia das Letras, 2000.

[43] ROUSMANIERE, Tony et al. Large language models as mental health resources: Patterns of use in the united states. 18 mar. 2025.

[44] SARAIVA, F. **IA assume papel de terapeuta para muitos profissionais, diz estudo de Harvard**. Exame, [s.l.], 17 jul. 2025. Disponível em: <https://exame.com/carreira/ia-assume-papel-de-terapeuta-para-muitos-profissionais-diz-estudo-de-harvard/>. Acesso em: 27 jul. 2025.

[45] SINGH, S. **ChatGPT Statistics for 2023: Comprehensive Facts and Data**. DemandSage, [s.l.], 24 jul. 2025. Disponível em: <https://www.demandsage.com/chatgpt-statistics/>. Acesso em: 27 jul. 2025.

[46] SOUZA, Ana Cláudia de; ALEXANDRE, Neusa Maria Costa; GUIRARDELLO, Edinêis de Brito. Propriedades psicométricas na avaliação de instrumentos: avaliação da confiabilidade e da validade. **Epidemiol. Serv. Saúde**, Brasília , v. 26, n. 3, p. 649-659, set. 2017 .

[47] SPICYBB0I. **I'm secretly using AI as a therapist and I'm kind of ashamed about it**. Reddit, [s.l.], 01 mai. 2025. Disponível em: https://www.reddit.com/r/TrueOffMyChest/comments/1kc9b5r/im_secretly_using_ai_a_s_a_therapist_and_im_kind/. Acesso em: 27 jul. 2025.

[48] SUN, Jie et al. Artificial intelligence in psychiatry research, diagnosis, and therapy. **Asian journal of psychiatry**, v. 87, p. 103705, 2023.

[49] TALK INC RESEARCH. **Inteligência Artificial na Vida Real**. Talk Digital, [s.l.], 2024. Disponível em: <https://talkdigital.co/ianavidareal/index.html>. Acesso em: 27 jul. 2025.

[50] TIDY, J. **Character.ai: Young people turning to AI therapist bots**. BBC, [s.l.], 4 jan. 2024. Disponível em: <https://www.bbc.com/news/technology-67872693>. Acesso em: 27 jul. 2025.

[51] UNZELTE, Carolina. **Terapia por IA: atendimento 24h, barato e sem julgamentos atrai pacientes. Mas funciona?**. Exame, [s.l.], 04 de mar. 2024. Disponível em: <https://exame.com/inteligencia-artificial/terapia-por-ia-atendimento-24h-barato-e-sem-julgamentos-atrai-pacientes-mas-funciona/>. Acesso em: 27 jul. 2025.

[52] UOL. **IA do futuro, Chat GPT como terapeuta, “internet morta” e mais com Dora Kaufman | Alt Tabet**. Youtube, [s.l.], 02 jul. 2025. Disponível em: <https://www.youtube.com/watch?v=FDROP8DC5Ac>. Acesso em: 27 jul. 2025.

[53] VASWANI, Ashish et al. Attention is all you need. **Advances in neural information processing systems**, v. 30, 2017.

[54] WANG, Liying et al. Evaluating Generative AI in Mental Health: Systematic Review of Capabilities and Limitations. **JMIR mental health**, v. 12, n. 1, p. e70014, 2025.

[55] WHO. **Constitution of the World Health Organization**. World Health Organization, Genebra, 22 jul. 1946. Disponível em: <https://apps.who.int/gb/bd/PDF/bd47/EN/constitution-en.pdf?ua=1>. Acesso em: 27 jul. 2025.

[56] WHO. **Mental Disorders**. World Health Organization, [s.l.], 08 jul. 2022. Disponível em: <https://www.who.int/news-room/fact-sheets/detail/mental-disorders>. Acesso em: 27 jul. 2025.

[57] WHOLIN, Claes et al. Experimentation in software engineering: an introduction. **Massachusetts: Kluwer Academic Publishers**, v. 2, p. 274, 2012.

[58] ZAO-SANDERS, M. **Top 10 Gen AI Use Cases**. Harvard Business Review, [s.l.], 09 abr. 2025. Disponível em: <https://hbr.org/data-visuals/2025/04/top-10-gen-ai-use-cases>. Acesso em: 27 jul. 2025.

[59] ZEWE, A. **Explained: Generative AI’s environmental impact**. MIT News, [s.l.], 17 jan. 2025. Disponível em: <https://news.mit.edu/2025/explained-generative-ai-environmental-impact-0117>. Acesso em: 27 jul. 2025.

[A1] SPALLEK, Sophia et al. Can we use ChatGPT for mental health and substance use education? Examining its quality and potential harms. **JMIR Medical Education**, v. 9, n. 1, p. e51243, 2023

[A2] NI, Yeming et al. Focusing on Needs: A Chatbot-Based Emotion Regulation Tool for Adolescents. In: **2023 IEEE International Conference on Systems, Man, and Cybernetics (SMC)**. IEEE, 2023. p. 2295-2300.

[A3] PATARANUTAPORN, Pat et al. Living memories: AI-generated characters as digital mementos. In: **Proceedings of the 28th International Conference on Intelligent User Interfaces**. 2023. p. 889-901.

- [A4] JANG, JiWoong et al. "It's the only thing I can trust": Envisioning large language model use by autistic workers for communication assistance. In: **Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems**. 2024. p. 1-18.
- [A5] ABUBAKAR, Abdulqahar Mukhtar; GUPTA, Deepa; PARIDA, Shantipriya. A reinforcement learning approach for intelligent conversational chatbot for enhancing mental health therapy. **Procedia Computer Science**, v. 235, p. 916-925, 2024.
- [A6] SORINO, Paolo et al. Ariel: Brain-computer interfaces meet large language models for emotional support conversation. In: **Adjunct Proceedings of the 32nd ACM Conference on User Modeling, Adaptation and Personalization**. 2024. p. 601-609.
- [A7] XU, Zhenyu et al. Can large language models be good companions? An LLM-based eyewear system with conversational common ground. **Proceedings of the ACM on Interactive, Mobile, Wearable and Ubiquitous Technologies**, v. 8, n. 2, p. 1-41, 2024.
- [A8] KOVACEVIC, Nikola et al. Chatbots with attitude: Enhancing chatbot interactions through dynamic personality infusion. In: **Proceedings of the 6th ACM Conference on Conversational User Interfaces**. 2024. p. 1-16.
- [A9] MA, Zilin et al. Evaluating the experience of LGBTQ+ people using large language model based chatbots for mental health support. In: **Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems**. 2024. p. 1-15.
- [A10] MO, Wen; SINGH, Aneesha; HOLLOWAY, Catherine. From Information Seeking to Empowerment: Using Large Language Model Chatbot in Supporting Wheelchair Life in Low Resource Settings. In: **Proceedings of the 26th International ACM SIGACCESS Conference on Computers and Accessibility**. 2024. p. 1-18.
- [A11] BIN-HADY, Wagdi Rashad Ali; ALI, Jamal Kaid Mohammed; AL-HUMARI, Mustafa Ahmed. The effect of ChatGPT on EFL students' social and emotional learning. **Journal of Research in Innovative Teaching & Learning**, v. 17, n. 2, p. 243-255, 2024.
- [A12] WESTER, Joel et al. "This Chatbot Would Never...": Perceived Moral Agency of Mental Health Chatbots. **Proceedings of the ACM on human-computer Interaction**, v. 8, n. CSCW1, p. 1-28, 2024.
- [A13] LI, Zhuoyang et al. "This is human intelligence debugging artificial intelligence": Examining how people prompt GPT in seeking mental health support. **International Journal of Human-Computer Studies**, p. 103555, 2025.
- [A14] SCHOLICH, Till et al. A Comparison of Responses from Human Therapists and Large Language Model-Based Chatbots to Assess Therapeutic Communication: Mixed Methods Study. **JMIR Mental Health**, v. 12, n. 1, p. e69709, 2025.
- [A15] WOLFE, Brooke H. et al. Caregiving artificial intelligence chatbot for older adults and their preferences, well-being, and social connectivity: mixed-method study. **Journal of Medical Internet Research**, v. 27, p. e65776, 2025.

[A16] ZHENG, Xi et al. Customizing emotional support: How do individuals construct and interact with LLM-powered chatbots. In: **Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems**. 2025. p. 1-20.

[A17] KANG, Boyoung; HONG, Munpyo. Development and evaluation of a mental health chatbot using ChatGPT 4.0: Mixed methods user experience study with Korean users. **JMIR Medical Informatics**, v. 13, p. e63538, 2025.

[A18] MOYLAN, Kayley; DOHERTY, Kevin. Expert and Interdisciplinary Analysis of AI-Driven Chatbots for Mental Health Support: Mixed Methods Study. **Journal of Medical Internet Research**, v. 27, p. e67114, 2025.

[A19] SONG, Inhwa et al. Exploreself: Fostering user-driven exploration and reflection on personal challenges with adaptive guidance by large language models. In: **Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems**. 2025. p. 1-22.

[A20] DJUFRIL, Ray; FRAMPTON, Jessica R.; KNOBLOCH-WESTERWICK, Silvia. Love, marriage, pregnancy: Commitment processes in romantic relationships with AI chatbots. **Computers in Human Behavior: Artificial Humans**, v. 4, p. 100155, 2025.

[A21] CHOI, Ryuhaerang et al. Private Yet Social: How LLM Chatbots Support and Challenge Eating Disorder Recovery. In: **Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems**. 2025. p. 1-19.

[A22] FAN, Xianzhe et al. User-Driven Value Alignment: Understanding Users' Perceptions and Strategies for Addressing Biased and Discriminatory Statements in AI Companions. In: **Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems**. 2025. p. 1-19.

[N1] CORBYN, Zoë. **Elliq is 93-year-old Juanita's friend. She's also a robot.** The Guardian, [s.l.], 13 ago. 2021. Disponível em: <https://www.theguardian.com/us-news/2021/aug/13/elliq-robot-companion-seniors>. Acesso em: 27 jul. 2025.

[N2] WISEMAN, Eva. **Is robot therapy the future?** The Guardian, [s.l.], 08 ago. 2021. Disponível em: <https://www.theguardian.com/society/2021/aug/08/is-robot-therapy-the-future>. Acesso em: 27 jul. 2025.

[N3] KEIERLEBER, Mark. **Young and depressed? Try Woebot! The rise of mental health chatbots in the US.** The Guardian, [s.l.], 13 abr. 2022. Disponível em: <https://www.theguardian.com/us-news/2022/apr/13/chatbots-robot-therapists-youth-mental-health-crisis>. Acesso em: 27 jul. 2025.

[N4] CLARKE, Laurie. **'I learned to love the bot': meet the chatbots that want to be your best friend.** The Guardian, [s.l.], 19 mar. 2023. Disponível em: <https://www.theguardian.com/technology/2023/mar/19/i-learned-to-love-the-bot-meet-the-chatbots-that-want-to-be-your-best-friend>. Acesso em: 27 jul. 2025.

[N5] MAHDAWI, Arwa. **Could an ‘emotional support’ chatbot help me de-stress? Only one way to find out.** The Guardian, [s.l.], 08 mai. 2023. Disponível em: <https://www.theguardian.com/commentisfree/2023/may/08/emotional-support-chatbot-ai>. Acesso em: 27 jul. 2025.

[N6] WAKEFIELD, Jane. **Would you open up to a chatbot therapist?** BBC, [s.l.], 02 abr. 2023. Disponível em: <https://www.bbc.com/news/business-65110680>. Acesso em: 27 jul. 2025.

[N7] ROBB, Alice. **‘He checks in on me more than my friends and family’: can AI therapists do better than the real thing?** The Guardian, [s.l.], 02 mar. 2024. Disponível em: <https://www.theguardian.com/lifeandstyle/2024/mar/02/can-ai-chatbot-therapists-do-better-than-the-real-thing>. Acesso em: 27 jul. 2025.

[N8] COX, David. **‘They thought they were doing good but it made people worse’: why mental health apps are under scrutiny.** The Guardian, [s.l.], 04 fev. 2024. Disponível em: <https://www.theguardian.com/society/2024/feb/04/they-thought-they-were-doing-good-but-it-made-people-worse-why-mental-health-apps-are-under-scrutiny>. Acesso em: 27 jul. 2025.

[N9] SAMPLE, Ian. **AI researchers build ‘future self’ chatbot to inspire wise life choices.** The Guardian, [s.l.], 05 jun. 2024. Disponível em: <https://www.theguardian.com/technology/article/2024/jun/05/ai-researchers-build-future-self-chatbot-to-inspire-wise-life-choices>. Acesso em: 27 jul. 2025.

[N10] MILMO, Dan. **AI tools may soon manipulate people’s online decision-making, say researchers.** The Guardian, [s.l.], 30 dez. 2024. Disponível em: <https://www.theguardian.com/technology/2024/dec/30/ai-tools-may-soon-manipulate-peoples-online-decision-making-say-researchers>. Acesso em: 27 jul. 2025.

[N11] TIDY, Joe. **Character.ai: Young people turning to AI therapist bots.** BBC, [s.l.], 04 jan. 2024. Disponível em: <https://www.bbc.com/news/technology-67872693>. Acesso em: 27 jul. 2025.

[N12] ABRAHAM, Amelia. **Computer says yes: how AI is changing our romantic lives.** The Guardian, [s.l.], 16 jun. 2024. Disponível em: <https://www.theguardian.com/technology/article/2024/jun/16/computer-says-yes-how-ai-is-changing-our-romantic-lives>. Acesso em: 27 jul. 2025.

[N13] SAMPLE, Ian. **Could AI help cure ‘downward spiral’ of human loneliness?** The Guardian, [s.l.], 27 mai. 2024. Disponível em: <https://www.theguardian.com/technology/article/2024/may/27/could-ai-help-cure-downward-spiral-of-human-loneliness>. Acesso em: 27 jul. 2025.

[N14] ZHANG, Wanqing. **Dan’s the man: Why Chinese women are looking to ChatGPT for love.** BBC, [s.l.], 13 jun. 2024. Disponível em: <https://www.bbc.com/articles/c4nnje9rpjgo>. Acesso em: 27 jul. 2025.

[N15] MULDOON, James. **Maybe we can role-play something fun': When an AI companion wants something more.** BBC, [s.l.], 09 out. 2024. Disponível em: <https://www.bbc.com/future/article/20241008-the-troubling-future-of-ai-relationships>. Acesso em: 27 jul. 2025.

[N16] MONTGOMERY, Blake. **Mother says AI chatbot led her son to kill himself in lawsuit against its maker.** The Guardian, [s.l.], 23 out. 2024. Disponível em: <https://www.theguardian.com/technology/2024/oct/23/character-ai-chatbot-sewell-set-zer-death>. Acesso em: 27 jul. 2025.

[N17] SCHUTZ, Elna. **The chatbot has transformed my life'.** BBC, [s.l.], 29 mai. 2024. Disponível em: <https://www.bbc.com/news/articles/c7223v5d8lgo>. Acesso em: 27 jul. 2025.

[N18] MILMO, Dan. **'It cannot provide nuance': UK experts warn AI therapy chatbots are not safe.** The Guardian, [s.l.], 07 mai. 2025. Disponível em: <https://www.theguardian.com/technology/2025/may/07/experts-warn-therapy-ai-chatbots-are-not-safe-to-use>. Acesso em: 27 jul. 2025.

[N19] BATTY, David. **'She helps cheer me up': the people forming relationships with AI chatbots.** The Guardian, [s.l.], 15 abr. 2025. Disponível em: <https://www.theguardian.com/technology/2025/apr/15/she-helps-cheer-me-up-the-people-forming-relationships-with-ai-chatbots>. Acesso em: 27 jul. 2025.

[N20] TOH, Justine. **Can a Jesus chatbot replace the real thing? The Easter story suggests not.** The Guardian, [s.l.], 21 abr. 2025. Disponível em: <https://www.theguardian.com/commentisfree/2025/apr/21/can-a-jesus-chatbot-replace-the-real-thing-the-easter-story-suggests-not>. Acesso em: 27 jul. 2025.

[N21] LAWRIE, Eleanor. **Can AI therapists really be an alternative to human help?.** BBC, [s.l.], 19 mai. 2025. Disponível em: <https://www.bbc.com/news/articles/ced2ywg7246o>. Acesso em: 27 jul. 2025.

[N22] NG, Kelly. **DeepSeek moved me to tears': How young Chinese find therapy in AI.** BBC, [s.l.], 12 fev. 2025. Disponível em: <https://www.bbc.com/news/articles/cy7g45g2nxno>. Acesso em: 27 jul. 2025.

[N23] STOKEL-WALKER, Chris. **How an embarrassing U-turn exposed a concerning truth about ChatGPT.** The Guardian, [s.l.], 01 mai. 2025. Disponível em: <https://www.theguardian.com/commentisfree/2025/may/01/chatgpt-chatbot-truth-user-update-ai>. Acesso em: 27 jul. 2025.

[N24] DAVIDSON, Helen. **In Taiwan and China, young people turn to AI chatbots for 'cheaper, easier' therapy.** The Guardian, [s.l.], 22 mai. 2025. Disponível em: <https://www.theguardian.com/world/2025/may/22/ai-therapy-therapist-chatbot-taiwan-china-mental-health>. Acesso em: 27 jul. 2025.

[N25] WALLACE, Lindsay Lee. **This app became my best friend': Mourning is human. New grief apps want to 'optimise' it for you.** BBC, [s.l.], 25 jan. 2025. Disponível em:

<https://www.bbc.com/future/article/20250123-the-apps-turning-grief-into-data-points>. Acesso em: 27 jul. 2025.

[R1] MIKE2800. **ChatGPT is better than my therapist, holy shit.** Reddit, [s.l.], 21 dez. 2022. Disponível em: https://www.reddit.com/r/ChatGPT/comments/zr5e17/chatgpt_is_better_than_my_the_rapist_holy_shit/. Acesso em: 27 jul. 2025.

[R2] FABULOUS_RICH8974. **ChatGPT helped me solve problems in my business.** Reddit, [s.l.], 05 jul. 2023. Disponível em: https://www.reddit.com/r/ChatGPT/comments/14rd57q/chatgpt_helped_me_solve_problems_in_my_business/. Acesso em: 27 jul. 2025.

[R3] MONKEYBALLPIRATE. **Concerns About Changes in ChatGPT's Handling of Mental Health Topics.** Reddit, [s.l.], 25 mai. 2023. Disponível em: https://www.reddit.com/r/ChatGPT/comments/13ruz3r/concerns_about_changes_in_chatgpts_handling_of/. Acesso em: 27 jul. 2025.

[R4] STEADFASTEND. **From a psychological-therapy standpoint, ChatGPT has been an absolute godsend fo....** Reddit, [s.l.], 05 abr. 2023. Disponível em: https://www.reddit.com/r/ChatGPT/comments/12cr4cw/from_a_psychologicaltherapy_standpoint_chatgpt/. Acesso em: 27 jul. 2025.

[R5] VHPOET. **My Experience with AI-Driven Journaling.** Reddit, [s.l.], 20 nov. 2023. Disponível em: https://www.reddit.com/r/selfhelp/comments/17znhsq/my_experience_with_aidriiven_journaling/. Acesso em: 27 jul. 2025.

[R6] FICKLE_BASE_7723. **Almost ended my life today.** Reddit, [s.l.], 18 out. 2024. Disponível em: https://www.reddit.com/r/mentalhealth/comments/1g67l0y/almost_ended_my_life_today/. Acesso em: 27 jul. 2025.

[R7] VARIOUS_PRINT2285. **Anybody try using AI as a therapist?.** Reddit, [s.l.], 13 set. 2024. Disponível em: https://www.reddit.com/r/mentalhealth/comments/1fg2jcy/anybody_try_using_ai_as_a_therapist/. Acesso em: 27 jul. 2025.

[R8] ASMA_UT. **Can't afford a therapist?.** Reddit, [s.l.], 05 nov. 2024. Disponível em: https://www.reddit.com/r/mentalhealth/comments/1gk4d2x/cant_afford_a_therapist/. Acesso em: 27 jul. 2025.

[R9] GREENIDENTITY. **Chat GPT is saving me from myself.** Reddit, [s.l.], 09 out. 2024. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1g04bv6/chat_gpt_is_saving_me_from_myself/. Acesso em: 27 jul. 2025.

[R10] ARTSPAWNER. **Chat GPT Transforms My Mental Health In 2 Weeks.** Reddit, [s.l.], 06 set. 2024. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1fajq7r/chat_gpt_transforms_my_mental_health_in_2_weeks/. Acesso em: 27 jul. 2025.

[R11] NIC_KNACK. **chatgpt as a supplement to therapy - my experience.** Reddit, [s.l.], 26 dez. 2024. Disponível em: https://www.reddit.com/r/therapy/comments/1hmhtl2/chatgpt_as_a_supplement_to_therapy_my_experience/. Acesso em: 27 jul. 2025.

[R12] CHAODDIAN. **ChatGPT as therapy?** Reddit, [s.l.], 17 dez. 2024. Disponível em: https://www.reddit.com/r/mentalhealth/comments/1hgcbjr/chatgpt_as_therapy/. Acesso em: 27 jul. 2025.

[R13] DEATHRAINBOWS. **ChatGPT helped me get sober.** Reddit, [s.l.], 24 dez. 2024. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1hl3h8m/chatgpt_helped_me_get_sober/. Acesso em: 27 jul. 2025.

[R14] THATNOBLEDUKE. **ChatGPT is the only one keeping me from losing my sanity.** Reddit, [s.l.], 09 dez. 2024. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1haepw2/chatgpt_is_the_only_one_keeping_me_from_losing_my/. Acesso em: 27 jul. 2025.

[R15] DESOLATENATURE. **ChatGPT made me feel more seen & gave me more hope in a 10 minute conversation t....** Reddit, [s.l.], 31 dez. 2024. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1hqa0xn/chatgpt_made_me_feel_more_seen_gave_me_more_hope/. Acesso em: 27 jul. 2025.

[R16] NO-GUR-7191. **ChatGPT therapy saved me.** Reddit, [s.l.], 03 set. 2024. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1f85dl7/chatgpt_therapy_saved_me/. Acesso em: 27 jul. 2025.

[R17] STILL-SIR6180. **Has anyone experimented with an AI tool to manage their anxiety? Here's my exper....** Reddit, [s.l.], 21 jun. 2024. Disponível em: https://www.reddit.com/r/mentalhealth/comments/1dktjmi/has_anyone_experimeted_with_an_ai_tool_to_manage/. Acesso em: 27 jul. 2025.

[R18] SUNBETTER7301. **I've made more progress in 6 hours of ChatGPT therapy than I have over 10 years....** Reddit, [s.l.], 24 dez. 2024. Disponível em: https://www.reddit.com/r/therapy/comments/1hla339/ive_made_more_progress_in_6_hours_of_chatgpt/. Acesso em: 27 jul. 2025.

[R19] TNT_GUERILLA. **PSA: Stop giving your sensitive, personal information to Big AI.** Reddit, [s.l.], 19 dez. 2024. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1hhkv6y/psa_stop_giving_your_sensitive_personal/. Acesso em: 27 jul. 2025.

[R20] JUBILEESUPREME. **Really fucked up experience with Therapy Chat.** Reddit, [s.l.], 19 fev. 2024. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1auowiv/really_fucked_up_experience_with_therapy_chat/. Acesso em: 27 jul. 2025.

[R21] FRENCHDRESSES. **To the people who shared that they used chatgpt to open up emotionally, thank yo....** Reddit, [s.l.], 08 set. 2024. Disponível em:

https://www.reddit.com/r/ChatGPT/comments/1fbnkcw/to_the_people_who_shared_hat_they_used_chatgpt/. Acesso em: 27 jul. 2025.

[R22] SAMURAIUX. **What ChatGPT Has to Say for itself as a Therapist.** Reddit, [s.l.], 03 out. 2024. Disponível em: https://www.reddit.com/r/therapy/comments/1fve31a/what_chatgpt_has_to_say_for_itself_as_a_therapist/. Acesso em: 27 jul. 2025.

[R23] INDIVIDUALLIBRARY358. **A few weeks ago I would've rolled my eyes at someone who said they use ChatGPT I...** Reddit, [s.l.], 15 jun. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1lbx5fz/a_few_weeks_ago_i_wouldve_rolled_my_eyes_at/. Acesso em: 27 jul. 2025.

[R24] SNDVSTAN. **AI as therapy.** Reddit, [s.l.], 07 jul. 2025. Disponível em: https://www.reddit.com/r/therapy/comments/1ltz8fi/ai_as_therapy/. Acesso em: 27 jul. 2025.

[R25] NOJUMPERFAN10. **Apparently, the human race is doomed because ChatGPT helped me** 😞. Reddit, [s.l.], 18 jun. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1legywt/apparently_the_human_race_is_doomed_because/. Acesso em: 27 jul. 2025.

[R26] KISHILEA. **ChatGPT has helped me more than 15 years of therapy. No joke.** Reddit, [s.l.], 17 abr. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1k1dxxp/chatgpt_has_helped_me_more_than_15_years_of/. Acesso em: 27 jul. 2025.

[R27] NIXFLORAMINE. **ChatGPT is actually amazing for mental health.** Reddit, [s.l.], 20 mai. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1kqwte8/chatgpt_is_actually_amazing_for_mental_health/. Acesso em: 27 jul. 2025.

[R28] BILL_THE_MURRAY. **ChatGPT just miraculously worked as a therapist for my six-year-old daughter hav....** Reddit, [s.l.], 02 fev. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1ig9vi4/chatgpt_just_miraculously_worked_as_a_therapist/. Acesso em: 27 jul. 2025.

[R29] BESTESTMOONCALF. **ChatGPT made me psychotic. AMA.** Reddit, [s.l.], 03 jul. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1lqmza9/chatgpt_made_me_psychotic_ama/. Acesso em: 27 jul. 2025.

[R30] FL_SNOWMAN. **ChatGPT Saved my Marriage.** Reddit, [s.l.], 17 jan. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1i3hl52/chatgpt_saved_my_marriage/. Acesso em: 27 jul. 2025.

[R31] KIWI_WIZARD. **ChatGPT seems like the only one who truly gets me.** Reddit, [s.l.], 08 abr. 2025. Disponível em: https://www.reddit.com/r/mentalhealth/comments/1juq5zj/chatgpt_seems_like_the_only_one_who_truly_gets_me/. Acesso em: 27 jul. 2025.

[R32] ESINEM13. **ChatGPT shattered the reality no one else would.** Reddit, [s.l.], 26 abr. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1k8cbpq/chatgpt_shattered_the_reality_no_one_else_would/. Acesso em: 27 jul. 2025.

[R33] CHEEZYCOW. **GPT as Therapy has saved my life.** Reddit, [s.l.], 04 mar. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1j32qcx/gpt_as_therapy_has_saved_my_life/. Acesso em: 27 jul. 2025.

[R34] CODENOMESAILORV. **I cried talking to ChatGPT today.** Reddit, [s.l.], 02 mai. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1kdd0th/i_cried_talking_to_chatgpt_today/. Acesso em: 27 jul. 2025.

[R35] MAIN-BOYSENBERRY3564. **I feel pathetic for using chatgpt for mental health.** Reddit, [s.l.], 26 jun. 2025. Disponível em: https://www.reddit.com/r/therapy/comments/1lleme8/i_feel_pathetic_for_using_chatgpt_for_mental/. Acesso em: 27 jul. 2025.

[R36] T6H6R6O6W6A6W6A6Y6. **I feel so betrayed, a warning.** Reddit, [s.l.], 18 abr. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1k1st3q/i_feel_so_betrayed_a_warning/. Acesso em: 27 jul. 2025.

[R37] RELATABLE107. **I just had surreal feeling about AI.** Reddit, [s.l.], 05 jun. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1l3vtrj/i_just_had_surreal_feeling_about_ai/. Acesso em: 27 jul. 2025.

[R38] TRICKYEMPLOYMENT8656. **I let ChatGPT be my therapist for a day, and i feel guilty.** Reddit, [s.l.], 20 jun. 2025. Disponível em: https://www.reddit.com/r/therapy/comments/1lg43v9/i_let_chatgpt_be_my_therapist_for_a_day_and_i/. Acesso em: 27 jul. 2025.

[R39] CARROTSARE2COOL. **I've been using ChatGPT as a therapist lately and it's been surprisingly helpful.** Reddit, [s.l.], 11 jan. 2025. Disponível em: https://www.reddit.com/r/therapy/comments/1hyluvh/ive_been_using_chatgpt_as_a_therapist_lately_and/. Acesso em: 27 jul. 2025.

[R40] MAGNOLIAMAHOOGANY. **If you understand what it means to go through life without a support system, you....** Reddit, [s.l.], 22 jun. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1li2687/if_you_understand_what_it_means_to_go_through/. Acesso em: 27 jul. 2025.

[R41] NEWSYTOO. **Now I get it.** Reddit, [s.l.], 10 abr. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1jvydih/now_i_get_it/. Acesso em: 27 jul. 2025.

[R42] CUTEKATGURL. **People aggressively shame/mock the use of ChatGPT for emotional support, failing....** Reddit, [s.l.], 11 mai. 2025. Disponível em:

https://www.reddit.com/r/ChatGPT/comments/1kkahek/people_aggressively_shame_mock_the_use_of_chatgpt/. Acesso em: 27 jul. 2025.

[R43] SUSPICIOUS_FERRET906 . **PSA: CHAT GPT IS A TOOL. NOT YOUR FRIEND.** Reddit, [s.l.], 03 mar. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1j2lebf/psa_chat_gpt_is_a_tool_not_your_friend/. Acesso em: 27 jul. 2025.

[R44] BIPBOPBATTREE. **Suddenly hate venting to Chatgpt.** Reddit, [s.l.], 26 fev. 2025. Disponível em: https://www.reddit.com/r/mentalhealth/comments/1iyj5hs/suddenly_hate_venting_to_chatgpt/. Acesso em: 27 jul. 2025.

[R45] THEBARD99. **Well. It finally happened.** Reddit, [s.l.], 06 jun. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1l4iht3/well_it_finally_happened/. Acesso em: 27 jul. 2025.

[R46] MESRSZMIT. **What am I supposed to use to cope if ChatGPT is so bad?** Reddit, [s.l.], 18 abr. 2025. Disponível em: https://www.reddit.com/r/mentalhealth/comments/1k2gkps/what_am_i_supposed_to_use_to_cope_if_chatgpt_is/. Acesso em: 27 jul. 2025.

[R47] EKI75. **Who uses ChatGPT for therapy?** Reddit, [s.l.], 02 mai. 2025. Disponível em: https://www.reddit.com/r/ChatGPT/comments/1kcox5v/who_uses_chatgpt_for_therapy/. Acesso em: 27 jul. 2025.

[T1] SABRINA_RAMONOV. **“ChatGPT has helped me more than 15 years of therapy” - agree or disagree?** 🍄 #a.... TikTok, [s.l.], 18 abr. 2025. Disponível em: https://www.tiktok.com/@sabrina_ramonov/video/7494744300111007006?q=ai_therapy&t=17492143026322.

[T2] BRANDNAT. **Replying to @domochvsky ChatGPT already knows YOUR secrets PT2. This one's a lon....** TikTok, [s.l.], 27 nov. 2024. Disponível em: https://www.tiktok.com/@brandnat/video/7441939336112688391?q=ai_therapy&t=1749214302632.

[T3] DRGEORGESACHS. **Will you want a human therapist in the future or an AI Taylor Swift therapist? N....** TikTok, [s.l.], 10 mai. 2025. Disponível em: https://www.tiktok.com/@drgeorgesachs/video/7502874084968058143?q=ai_therapy&t=1749214302632.

[T4] FLUENTLYFORWARD. **How i use #chatgpt for makeshift #therapy or a way to understand my feelings.** TikTok, [s.l.], 24 jul. 2024. Disponível em: https://www.tiktok.com/@fluentlyforward/video/7395214887611403562?q=ai_therapy&t=1749214302632.

[T5] PARISBEYK. **What you might have missed about using chatgpt for emotional support or as your** TikTok, [s.l.], 30 mar. 2025. Disponível em: https://www.tiktok.com/@parisbeyk/video/7487665093648436502?q=ai_therapy&t=1749214302632.

[T6] THERAPYJEFF. **If you're going to use Chat GPT as your therapist use these 3 prompts. Prompt 1:....** TikTok, [s.l.], 05 jul. 2025. Disponível em: <https://www.tiktok.com/@therapyjeff/video/7523607961726373175?q=ai&therapy&t=1749214302632>.

[T7] DRANDREAS. **Will AI replace therapists? Yeah, a little bit! #ai #mentalhealth #therapy.** TikTok, [s.l.], 15 jul. 2024. Disponível em: <https://www.tiktok.com/@drandreas/video/7392031529653767455?q=ai&therapy&t=17492143026322>.

[T8] PRESTONRHO. **This #ChatGPT #Prompt that will uncover the deepest & darkest emotional secrets** TikTok, [s.l.], 09 mar. 2025. Disponível em: <https://www.tiktok.com/@prestonrho/video/7479838566609128750?q=ai&therapy&t=1749214302632>.

[T9] ASKCATGPT. **Its goal is to give you an answer that's "useful," which isn't the same as an ans....** TikTok, [s.l.], 04 jan. 2025. Disponível em: <https://www.tiktok.com/@askcatgpt/video/7455931395412725035?q=ai&therapy&t=1749214302632>.

[T10] .PRESRO. **This is a role play I did with Chat GPT 4o I was just curious to see how it woul....** TikTok, [s.l.], 13 abr. 2025. Disponível em: <https://www.tiktok.com/@.presro/video/7492977293862784286?q=ai&therapy&t=1749214302632>.

[T11] THERAPYJEFF. **4 Reasons AI chat bot therapy is so much worse than real human therapy. #mentalh....** TikTok, [s.l.], 11 abr. 2025. Disponível em: <https://www.tiktok.com/@therapyjeff/video/7492039707879132459?q=ai&therapy&t=1749214302632>.

[T12] ISADNTN. **assunto sério emmmm #psicologia #psicoterapia #saudemental #inteligenciaartifici....** TikTok, [s.l.], 01 ago. 2024. Disponível em: <https://www.tiktok.com/@isadntn/video/7398201840518614278?q=ai&therapy&t=1749214302632>.

[T13] NIKWEBER.AI. **Simulando um Terapeuta no ChatGPT 🧠 Se você está passando por qualquer problema....** TikTok, [s.l.], 01 jun. 2023. Disponível em: <https://www.tiktok.com/@nikweber.ai/video/7239854720104320262?q=ai&therapy&t=1749214302632>.

[T14] THE_BRAIN_SCIENTIST. **AI for mental health. I've been using Willow for a few weeks. It took a few sess....** TikTok, [s.l.], 31 mai. 2024. Disponível em: https://www.tiktok.com/@the_brain_scientist/video/7375263801484119338?q=ai&therapy&t=1749214302632.

[T15] SHAHEEN4234. **There has been a lot of talk around AI and systems like #chatgpt replacing thera....** TikTok, [s.l.], 25 jun. 2025. Disponível em: <https://www.tiktok.com/@shaheen4234/video/7519884126065855750?q=ai&therapy&t=1749214302632>.

[T16] MIND.PSY.GUIDANCE. **(Some) dangers of using AI as your therapist 🖐️ If you're using ChatGPT be car....** TikTok, [s.l.], 25 abr. 2025. Disponível em:

<https://www.tiktok.com/@mind.psy.guidance/video/7497133664053087531?q=ai&t=1749214302632>.

[T17] PSIDIANACOSTA. **Eai, você acha que fazer terapia com IA funcionou? Pode servir de alívio momentâ...** TikTok, [s.l.], 29 abr. 2025. Disponível em: <https://www.tiktok.com/@psidianacosta/video/7498843184433270021?q=ai&t=1749214302632>.

[T18] ZACHSANGSHOW. **Is ChatGPT better than a real therapist? #therapy #therapist #chatgpt #ai #menta....** TikTok, [s.l.], 11 jun. 2025. Disponível em: <https://www.tiktok.com/@zachsangshow/video/7514685199419051294?q=ai&t=1749214302632>.

[T19] CYBERLAWSTAR. **I found, and broke an AI “therapist” chat bot the other day... reckless AI impleme....** TikTok, [s.l.], 17 abr. 2025. Disponível em: <https://www.tiktok.com/@cyberlawstar/video/7494264957207334175?q=ai&t=1749214302632>.

[T20] DOUGIESHARPE. **More and more people are turning to AI for therapy and research shows that peopl....** TikTok, [s.l.], 10 jun. 2025. Disponível em: <https://www.tiktok.com/@dougiesharpe/video/7514321248009112888?q=ai&t=1749214302632>.

[T21] YOURHAPANEXTDOOR. **She really needed that convo 🥺😭💙🥲 tried my best translating everything 🤗 #c....** TikTok, [s.l.], 15 jan. 2025. Disponível em: <https://www.tiktok.com/@yourhapanextdoor/video/7460103481286724886?q=ai&t=1749214302632>.

[T22] THATSILLY.GUY. **its genuinely so harmful guys. I could go on and on about how this app oftentime....** TikTok, [s.l.], 03 jun. 2025. Disponível em: <https://www.tiktok.com/@thatsilly.guy/video/7511620888001826070?q=ai&t=1749214302632>.

[T23] MYALA.SHAY. **anyways 10/10 recommend this 🥺💕 #ai #chatgpt #therapy #therapytok #aitok.** TikTok, [s.l.], 06 out. 2024. Disponível em: <https://www.tiktok.com/@myala.shay/video/7422770960983411998?q=ai&t=1749214302632>.

[T24] GIANLUCA.MAURO. **You could have a therapy session with AI. I'm going to give you a quick demo, an....** TikTok, [s.l.], 07 nov. 2023. Disponível em: <https://www.tiktok.com/@gianluca.mauro/video/7298797317413457184?q=ai&t=1749214302632>.

[T25] IMAGINEMICHELE. **If you've been thinking about therapy—or just some self-reflection—did you know** TikTok, [s.l.], 26 mar. 2025. Disponível em: <https://www.tiktok.com/@imaginemichele/video/7486253211020578091?q=ai&t=1749214302632>.

[T26] THERAPYTOTHEPOINT. **Will AI be able to replace therapists? #therapistsontiktok #therapysessions #psy....** TikTok, [s.l.], 09 out. 2024. Disponível em:

[https://www.tiktok.com/@therapytothepoint/video/7423799649733053739?q=ai therapy&t=1749214302632](https://www.tiktok.com/@therapytothepoint/video/7423799649733053739?q=ai%20therapy&t=1749214302632).

[T27] LEOCANDRADE. 🧠 **O ChatGPT pode ser sua primeira conversa de autocuidado. “Me ajudou mais que ...** TikTok, [s.l.], 23 mai. 2025. Disponível em: [https://www.tiktok.com/@leocandrade/video/7507740744338263301?q=ai therapy&t=1749214302632](https://www.tiktok.com/@leocandrade/video/7507740744338263301?q=ai%20therapy&t=1749214302632).

[T28] SABRINA_RAMONOV. **POV “my AI therapist saved my relationship” - I hear/read lots of anecdotal stor....** TikTok, [s.l.], 27 mai. 2025. Disponível em: [https://www.tiktok.com/@sabrina_ramonov/video/7509177933618089246?q=ai therapy&t=1749214302632](https://www.tiktok.com/@sabrina_ramonov/video/7509177933618089246?q=ai%20therapy&t=1749214302632).

[T29] BRITTASTVN. **save for when you feel lost / healing alone 📷 prompt @1:00 💡 Bonus tip: ask it....** TikTok, [s.l.], 17 abr. 2025. Disponível em: [https://www.tiktok.com/@brittastvn/video/7494427573456325918?q=ai therapy mental health chatbot&t=1752013562366](https://www.tiktok.com/@brittastvn/video/7494427573456325918?q=ai%20therapy%20mental%20health%20chatbot&t=1752013562366).

[T30] DR.AUDRA.HORNEY. **chat GPT therapy: yay or nay? #mensmentalhealth #therapyformen #ai #chatgpt #the....** TikTok, [s.l.], 13 fev. 2025. Disponível em: [https://www.tiktok.com/@dr.audra.horney/video/7471019802253872414?q=ai therapy mental health chatbot&t=1752013562366](https://www.tiktok.com/@dr.audra.horney/video/7471019802253872414?q=ai%20therapy%20mental%20health%20chatbot&t=1752013562366).

[T31] WORRYBOUTYOSELFLOVE. **This was stupid incredible. Just go do it please. #chatgpt #ai #higherself #guid....** TikTok, [s.l.], 12 fev. 2025. Disponível em: [https://www.tiktok.com/@worryboutyoselflove/video/7470614599239322926?q=ai therapy&t=1749214302632](https://www.tiktok.com/@worryboutyoselflove/video/7470614599239322926?q=ai%20therapy&t=1749214302632).

[T32] ALBERTA.NYC. **There’s a reason why every company tells you not to put your business data into A....** TikTok, [s.l.], 28 ago. 2024. Disponível em: [https://www.tiktok.com/@alberta.nyc/video/7408353395431623982?q=ai therapy&t=1749214302632](https://www.tiktok.com/@alberta.nyc/video/7408353395431623982?q=ai%20therapy&t=1749214302632).

[T33] HOPEWITHHOLLY. **How I use AI to help me with my mental health. #mentalhealth #mentalhealthawaren....** TikTok, [s.l.], 05 mai. 2024. Disponível em: [https://www.tiktok.com/@hopewithholly/video/7365554785619086635?q=ai therapy&t=1749214302632](https://www.tiktok.com/@hopewithholly/video/7365554785619086635?q=ai%20therapy&t=1749214302632).

[T34] TANNIBLELECTER. **Battling borderline with AI. When I’m struggling with attachment I go to AI. The....** TikTok, [s.l.], 17 mai. 2025. Disponível em: [https://www.tiktok.com/@tanniblelecter/video/7505485530033933599?q=ai therapy&t=1749214302632](https://www.tiktok.com/@tanniblelecter/video/7505485530033933599?q=ai%20therapy&t=1749214302632).

[T35] DIGITALLAURAANDERSON. **This is better than 10 years of therapy omg. #ai #chatgpt #chatgptprompt #chatgp....** TikTok, [s.l.], 04 jul. 2025. Disponível em: [https://www.tiktok.com/@digitallauraanderson/video/7523368459233758477?q=ai therapy&t=1749214302632](https://www.tiktok.com/@digitallauraanderson/video/7523368459233758477?q=ai%20therapy&t=1749214302632).

[T36] LCSWKATE. **A chat bot is not capable of empathy or logic. They are designed to emulate huma....** TikTok, [s.l.], 18 abr. 2025. Disponível em:

[https://www.tiktok.com/@lcswkate/video/7494733147272154411?q=ai therapy&t=1749214302632](https://www.tiktok.com/@lcswkate/video/7494733147272154411?q=ai%20therapy&t=1749214302632).

[T37] CALLMEBELLY. **boy problems are being solved** 🧠 #chatgpt #openai #ai #therapist. TikTok, [s.l.], 22 mai. 2024. Disponível em: [https://www.tiktok.com/@callmebelly/video/7371917338431589675?q=ai therapy&t=1749214302632](https://www.tiktok.com/@callmebelly/video/7371917338431589675?q=ai%20therapy&t=1749214302632).

[T38] ITSMIPIAA. **#chatgpt #lifehack #lifecoach #freetherapy #ai #aitherapy #fyp #manifestationcoa....** TikTok, [s.l.], 25 abr. 2025. Disponível em: [https://www.tiktok.com/@itsmiapiao/video/7497279910151474474?q=ai therapy&t=1749214302632](https://www.tiktok.com/@itsmiapiao/video/7497279910151474474?q=ai%20therapy&t=1749214302632).

[T39] HEALTHYGAMER.GG. **Can AI Replace Therapy?** 🤖. TikTok, [s.l.], 19 mai. 2023. Disponível em: [https://www.tiktok.com/@healthygamer.gg/video/7234903504744271109?q=ai therapy&t=1749214302632](https://www.tiktok.com/@healthygamer.gg/video/7234903504744271109?q=ai%20therapy&t=1749214302632).

[T40] MAALTOKS. **why relying on chatgpt as your therapist can stunt your inner world #techtok #ai.** TikTok, [s.l.], 20 jun. 2025. Disponível em: [https://www.tiktok.com/@maaltoks/video/7518185303304129822?q=ai therapy&t=1749214302632](https://www.tiktok.com/@maaltoks/video/7518185303304129822?q=ai%20therapy&t=1749214302632).

[T41] ARTISTJANAE. **she just gets it idk #therapytiktok #chatgpt #ai #contentcreator #teamwork.** TikTok, [s.l.], 04 dez. 2024. Disponível em: [https://www.tiktok.com/@artistjanae/video/7444680107731569962?q=ai therapy&t=1752700820020](https://www.tiktok.com/@artistjanae/video/7444680107731569962?q=ai%20therapy&t=1752700820020).

APÊNDICES

Apêndice A – Planilha eletrônica usada durante o estudo

A planilha completa com os dados utilizados no estudo pode ser acessada livremente para consulta.

ASSIS, Isabelle. Estruturação da MLR - TCC "Uso Espontâneo de LLMs para Fins Terapêuticos e Apoio Socioemocional". 2025. Disponível em: https://docs.google.com/spreadsheets/d/1NLgKXopf2ki8izhpmpvJ9_7GR-aHrTE0DvKj2o4y0o8/edit?usp=sharing. Acesso em: 29 jul. 2025.

Apêndice B – Script Python utilizado para extração de dados do Reddit

```
import praw

reddit = praw.Reddit(
    client_id="xxx",
    client_secret="xxx",
    user_agent="xxx"
)

results = []

# subreddit = reddit.subreddit("chatgpt")
# for post in subreddit.search("therapy", sort="top", limit=50):

# subreddit = reddit.subreddit("mentalhealth+therapy+selfhelp")
# for post in subreddit.search("chatgpt", sort="top", limit=50):

for post in reddit.subreddit('all').search('ai therapy', sort='hot', limit=50):
    post.comments.replace_more(limit=0)
    top_comments = [comment.body.strip().replace("\n", ' ') for comment in
post.comments[:10]]
```



```

while len(top_comments) < 10:
    top_comments.append("")

    post_data = f"""
    -----
    Title: {post.title.strip()}
    Link: https://www.reddit.com{post.permalink}
    Date: {post.created_utc}
    Body: {post.selftext}
    Search term: AI
    Subreddit: {post.subreddit.display_name}
    Score / Upvote Ratio: {post.score} / {post.upvote_ratio}
    Number of comments: {post.num_comments}

    Top Comments:
    1. {top_comments[0]}
    2. {top_comments[1]}
    3. {top_comments[2]}
    4. {top_comments[3]}
    5. {top_comments[4]}
    6. {top_comments[5]}
    7. {top_comments[6]}
    8. {top_comments[7]}
    9. {top_comments[8]}
    10. {top_comments[9]}
    """

    results.append(post_data)
    print("Feito para post " + post.title.strip())

with open("Reddit_Posts.txt", "w", encoding="utf-8") as f:
    f.writelines(results)

print("✅ Pronto!")

```

Apêndice C – Prompt de Aplicação de Critérios de Qualidade (Artigos)

“O tema da minha dissertação é “Uso Espontâneo de LLMs como ferramentas de apoio terapêutico e socioemocional”, com o objetivo de investigar e documentar como usuários utilizam LLMs para fins terapêuticos/sociais e mapear impactos relatados.

Com base nisso, avalie (com a escala 0, 0.5 e 1) esse artigo com base nas seguintes métricas:

Tema relacionado

Metodologia clara

Modelo ou proposta bem definido

Discussão sobre resultados encontrados

Contém análise crítica sobre o fenômeno

Mencionou desafios, limitações ou ameaças

Dê a nota para cada métrica e justifique.”

Apêndice D – Prompts de Coleta de Evidências (Artigos)

“O tema da minha dissertação é “Uso espontâneo de LLMs como ferramentas para fins terapêuticos e apoio socioemocional”, com o objetivo de investigar e documentar como usuários utilizam LLMs para fins terapêuticos/sociais e mapear impactos relatados.

Com base nisso, avalie (com a escala 0 para Não responde, 0.5 para Responde parcialmente e 1 para Responde claramente) esse artigo com base nas seguintes perguntas de pesquisa:

QP1: Quais são as motivações que levam usuários a adotar LLMs para fins terapêuticos e apoio socioemocional?

QP2: Como ocorre a dinâmica de interação para fins terapêuticos e apoio socioemocional?

QP3: Quais benefícios subjetivos são relatados para fins terapêuticos e apoio socioemocional?

QP4: Quais riscos emergem desse uso não supervisionado para fins terapêuticos e apoio socioemocional?

QP5: Qual é o perfil dos usuários que recorrem a LLMs para fins terapêuticos e apoio socioemocional?

(Prompt 1) Dê a nota para cada métrica e justifique."

(Prompt 2) Dê a nota para cada métrica e como justificativa retorne citações diretas do texto, sem traduzir e sem textos intermediários como justificativa."

Apêndice E – Prompt de Categorização e Classificação das Respostas às Questões de Pesquisa

"You are analyzing a spreadsheet containing {{tipo da publicação}} and their corresponding answers to research questions about the use of AI as a tool for social interaction and informal psychotherapy.

Your task is to categorize the responses to {{questão de pesquisa}} into meaningful categories based on recurring themes and topics.

Instructions:

Extract responses from the column corresponding to {{questão de pesquisa}}.

Identify key themes, patterns, and recurring topics in the responses.

Assign each response to one or more relevant categories.

Provide a summary of the categories, explaining the criteria used for classification.

If a response does not fit into existing categories, create a new category and justify it.

If there is no response add to a "No response" category, do not exclude any {{tipo da publicação}} for the output.

Expected Output:

A list of categories with descriptions.

Each response mapped to one or more categories.

A summary explaining the classification logic.

The result of this analysis in a spreadsheet, the column with the category must have the name of the category.

Ensure that the categorization is consistent, meaningful, and representative of the data. If necessary, refine the categories to improve clarity."

Apêndice F – Tabelas de fontes das categorias

Tabela F1 – Fontes de principais temas dos artigos

Categoria	Nº	Fontes
Análise de uso	14	[A2, A4, A8, A9, A10, A11, A12, A13, A15, A16, A17, A20, A21, A22]
Avaliação de modelo existente	12	[A1, A4, A10, A12, A13, A14, A15, A17, A18, A20, A21, A22]
Criação de modelo especializado	8	[A2, A3, A5, A6, A7, A8, A16, A19]

Fonte: A autora (2025).

Tabela F2 – Fontes de motivações, segundo a literatura acadêmica

Categoria	Nº	Fontes
Apoio Emocional e Desabafo	15	[A2, A3, A6, A7, A9, A10, A12, A13, A15, A16, A18, A19, A20, A21, A22]
Privacidade e Anonimato	12	[A4, A5, A9, A10, A12, A13, A14, A17, A18, A19, A20, A21]
Companhia e Redução da Solidão	12	[A2, A3, A6, A7, A10, A12, A14, A15, A16, A18, A20, A22]
Acessibilidade e Conveniência	10	[A4, A5, A9, A10, A12, A13, A14, A17, A18, A21]
Autoexploração	6	[A3, A7, A16, A17, A19, A22]
Abordagem de Vulnerabilidades	4	[A4, A9, A15, A21]
Sem Resposta	3	[A1, A8, A11]
Satisfação de Necessidades	2	[A20, A22]

Busca de Aconselhamento	2	[A3, A10]
-------------------------	---	-----------

Fonte: A autora (2025).

Tabela F3 – Fontes de padrões de interação, segundo a literatura acadêmica

Categoria	Nº	Fontes
Apoio Terapêutico e Emocional	14	[A1, A2, A5, A6, A9, A11, A12, A13, A14, A16, A17, A18, A19, A21]
Companheirismo e Conexão Social	6	[A3, A7, A8, A15, A20, A22]
Regulação Emocional e Autorreflexão	4	[A11, A13, A16, A17]
Desenvolvimento de Habilidades Sociais	3	[A4, A9, A11]
Informação e Educação	2	[A1, A10]

Fonte: A autora (2025).

Tabela F4 – Fontes de padrões de interação, segundo a literatura acadêmica

Categoria	Nº	Fontes
Terapia	15	[N2, N3, N4, N6, N7, N8, N11, N13, N15, N16, N18, N19, N21, N23, N24]
Companhia e Redução da Solidão	11	[N1, N4, N6, N12, N13, N14, N15, N16, N18, N19, N22]
Outros	8	[[N9], [N10], [N25], [N12], [N1], [N25], [N17], [N20]]
Relacionamento Romântico	6	[N4, N6, N12, N14, N15, N19]
Apoio Emocional e Desabafo	2	[N5, N22]
Conversas Casuais	2	[N10, N23]
Apoio à Neurodiversidade	2	[N17, N19]

Fonte: A autora (2025).

Tabela F5 – Fontes de casos de uso, segundo postagens no Reddit

Categoria	Nº	Fontes
Apoio Emocional e Desabafo	37	[R2, R3, R4, R5, R6, R7, R8, R9, R10, R11, R12, R14, R15, R17, R18, R23, R24, R25, R27, R28, R29, R30, R31, R32, R33, R34, R35, R37, R38, R39, R40, R41, R42, R44, R45, R46, R47]
Busca por Conselhos e Orientação	31	[R1, R2, R4, R5, R9, R10, R11, R12, R13, R17, R18, R19, R21, R23, R24, R25, R26, R27, R28, R29, R30, R31, R32, R33, R35, R37, R39, R40, R42, R45, R47]

Autoexploração	28	[R1, R2, R4, R5, R8, R9, R10, R11, R12, R13, R15, R16, R17, R18, R21, R23, R24, R25, R26, R27, R30, R32, R33, R37, R39, R42, R45, R47]
Simulação de Terapia	18	[R2, R4, R10, R12, R16, R17, R18, R19, R20, R23, R25, R28, R33, R36, R37, R40, R42, R45]
Companhia e Redução da Solidão	8	[R6, R7, R9, R31, R34, R35, R37, R40]
Nenhum	2	[R22, R43]
Necessidades de Grupos Específicos	1	[R13]

Fonte: A autora (2025).

Tabela F6 – Fontes de casos de uso, segundo vídeos no TikTok

Categoria	Nº	Fontes
Simulação de Terapia	15	[T1, T4, T6, T7, T10, T13, T14, T17, T20, T24, T25, T27, T30, T34, T37]
Nenhum	14	[T3, T5, T9, T11, T12, T15, T16, T18, T19, T26, T32, T36, T39, T40]
Exploração psicanalítica	11	[T1, T2, T4, T6, T8, T10, T25, T29, T31, T35, T38]
Apoio Emocional e Desabafo	3	[T21, T23, T41]
Busca por Conselhos e Orientação	3	[T28, T33, T38]
Companhia	2	[T22, T38]

Fonte: A autora (2025).

Tabela F7 – Fontes de abordagens de postagens no Reddit

Categoria	Nº	Fontes
Relato de experiência positiva	37	[R1, R2, R4, R5, R6, R7, R8, R9, R10, R11, R12, R13, R14, R15, R17, R18, R21, R23, R24, R25, R26, R27, R28, R30, R31, R32, R33, R34, R35, R37, R38, R39, R41, R42, R45, R46, R47]
Recomendação de uso	7	[R5, R6, R16, R17, R19, R26, R40]
Relato de experiência negativa	5	[R3, R20, R29, R36, R44]
Cuidado/Crítica	4	[R3, R22, R36, R43]

Fonte: A autora (2025).

Tabela F8 – Fontes de abordagens dos vídeos no TikTok

Categoria	Nº	Fontes
Aviso contra o uso	12	[T5, T12, T15, T16, T18, T19, T22, T26, T32,

		T36, T39, T40]
Dicas de uso (Prompt)	11	[T1, T2, T4, T6, T8, T13, T25, T27, T29, T31, T35]
Recomendação	6	[T14, T23, T33, T34, T38, T41]
Análise	5	[T3, T7, T20, T28, T30]
Demonstração	5	[T10, T17, T21, T24, T37]
Aviso de cuidado	4	[T6, T9, T11, T24]

Fonte: A autora (2025).

Tabela F9 – Fontes de benefícios, segundo a literatura acadêmica

Categoria	Nº	Fontes
Apoio Emocional	16	[A2, A3, A5, A6, A7, A9, A10, A11, A12, A13, A14, A16, A17, A18, A20, A21]
Ambiente Seguro e Sem Julgamento	10	[A4, A5, A8, A9, A10, A12, A13, A18, A20, A21]
Bem-estar e Emoções Positivas	10	[A2, A3, A6, A8, A11, A12, A14, A15, A17, A21]
Companheirismo e Intimidade	9	[A7, A9, A10, A15, A16, A18, A20, A21, A22]
Autoexploração	9	[A2, A3, A9, A11, A13, A16, A19, A21, A22]
Desenvolvimento de Habilidades	9	[A3, A4, A7, A9, A10, A11, A13, A15, A19]
Qualidade da Interação e Personalização	8	[A4, A5, A6, A7, A8, A13, A16, A17]
Empoderamento e Agência	7	[A4, A9, A10, A12, A19, A20, A21]
Sem Resposta	1	[A1]

Fonte: A autora (2025).

Tabela F10 – Fontes de benefícios, segundo a cobertura jornalística

Categoria	Nº	Fontes
Apoio Emocional	11	[N4, N5, N7, N11, N12, N14, N15, N17, N19, N21, N22]
Outros	10	[[N3, N21], [N6], [N13], [N25], [N1], [N9], [N17], [N9], [N6]]
Acessibilidade/Baixo Custo	8	[N2, N3, N7, N11, N21, N22, N24, N25]
Companheirismo	8	[N1, N4, N6, N12, N13, N14, N15, N19]
Ambiente Seguro e Sem Julgamento	7	[N3, N4, N7, N11, N21, N22, N24]
Disponibilidade	5	[N7, N11, N14, N21, N24]

Sentir-se Ouvido	5	[N5, N14, N15, N22, N25]
Ajuda com Habilidades Sociais	5	[N6, N12, N13, N17, N19]
Reduz Estigma	4	[N3, N17, N19, N24]
Reduz a solidão	2	[N1, N4]

Fonte: A autora (2025).

Tabela F11 – Fontes de benefícios, segundo postagens no Reddit

Categoria	Nº	Fontes
Apoio Emocional	33	[R2, R4, R6, R7, R8, R9, R10, R11, R12, R13, R14, R15, R17, R18, R23, R25, R27, R28, R30, R31, R32, R33, R34, R35, R37, R38, R39, R40, R41, R42, R44, R45, R47]
Autoexploração	28	[R1, R2, R4, R5, R8, R9, R10, R11, R12, R13, R15, R16, R17, R18, R21, R23, R24, R25, R26, R27, R30, R32, R33, R37, R39, R42, R45, R47]
Ambiente Seguro e Sem Julgamento	26	[R1, R4, R6, R8, R9, R10, R12, R14, R15, R17, R23, R24, R27, R30, R31, R33, R34, R35, R37, R38, R39, R41, R42, R45, R46, R47]
Benefícios Terapêuticos	25	[R1, R2, R4, R6, R9, R10, R11, R12, R15, R17, R18, R21, R23, R25, R26, R27, R28, R30, R31, R33, R37, R38, R40, R42, R45]
Espaço para Expressão Livre	22	[R1, R2, R4, R6, R8, R9, R10, R15, R17, R23, R25, R31, R33, R34, R35, R37, R38, R40, R42, R45, R46, R47]
Acessibilidade e Conveniência	18	[R4, R8, R9, R10, R14, R15, R17, R18, R19, R21, R22, R23, R28, R33, R37, R39, R40, R42]
Fonte de sabedoria	18	[R1, R2, R5, R10, R11, R12, R15, R17, R18, R21, R22, R23, R24, R25, R27, R30, R33, R42]
Companheirismo e Presença Social	14	[R6, R7, R9, R10, R14, R17, R31, R33, R34, R35, R37, R40, R42, R44]
Empoderamento e Confiança	9	[R2, R10, R14, R15, R23, R25, R26, R33, R42]
Nenhum	5	[R3, R20, R29, R36, R43]

Fonte: A autora (2025).

Tabela F12 – Fontes de benefícios, segundo vídeos no TikTok

Categoria	Nº	Fontes
Autoexploração	14	[T1, T2, T4, T6, T8, T10, T13, T25, T28, T29, T31, T35, T37, T38]
Nenhum	10	[T11, T12, T15, T19, T22, T26, T32, T36, T39, T40]

Disponibilidade	6	[T1, T16, T24, T27, T30, T33]
Acessibilidade/Baixo Custo	5	[T1, T16, T18, T27, T30]
Não específico	5	[T3, T14, T17, T20, T41]
Sentir-se ouvido	3	[T21, T23, T34]
Escassez de profissionais	1	[T7]
Anonimato	1	[T30]

Fonte: A autora (2025).

Tabela F13 – Fontes de riscos, segundo a literatura acadêmica

Categoria	Nº	Fontes
Conselhos inadequados	13	[A4, A6, A7, A8, A9, A12, A13, A14, A16, A18, A19, A20, A21]
Riscos emocionais	12	[A4, A6, A9, A11, A12, A13, A14, A18, A19, A20, A21, A22]
Dependência	11	[A9, A10, A11, A13, A14, A15, A17, A18, A20, A21, A22]
Linguagem estigmatizante	9	[A1, A4, A6, A9, A10, A17, A18, A19, A22]
Segurança de dados	9	[A1, A6, A7, A9, A10, A15, A16, A17, A18]
Imprecisão e desinformação	7	[A1, A6, A9, A10, A13, A17, A21]
Falhas técnicas e de segurança	6	[A8, A9, A13, A17, A20, A22]
Preocupações éticas	4	[A1, A14, A18, A19]
Manipulação	3	[A8, A15, A18]
Má gestão de crise	3	[A12, A14, A18]
Sem resposta	2	[A2, A5]

Fonte: A autora (2025).

Tabela F14 – Fontes de riscos, segundo a cobertura jornalística

Categoria	Nº	Fontes
Outros	22	[[N5], [N14], [N24], [N22], [N12], [N1], [N25], [N9], [N4], [N16], [N5], [N18], [N20], [N3], [N23], [N24], [N2], [N23], [N12], [N2], [N10], [N8]]
Segurança de Dados	13	[N1, N3, N4, N6, N8, N10, N11, N13, N14, N15, N21, N22, N25]
Conselhos inadequados	9	[N3, N6, N7, N8, N11, N15, N16, N18, N21]
Manipulação	7	[N2, N4, N8, N10, N15, N16, N23]

Falta de nuance/empatia	5	[N7, N11, N18, N24, N25]
Lacunas regulatórias	4	[N6, N8, N16, N18]
Exploração financeira	2	[N12, N15]
Falta de reciprocidade	2	[N1, N13]
Desencorajar a interação humana	2	[N1, N13]
Potencial para vício/apego	2	[N19, N22]

Fonte: A autora (2025).

Tabela F15 – Fontes de riscos, segundo postagens no Reddit

Categoria	Nº	Fontes
Nenhum	28	[R1, R2, R4, R6, R8, R9, R10, R12, R13, R14, R15, R16, R17, R18, R21, R23, R24, R25, R27, R30, R32, R33, R34, R37, R39, R41, R45, R47]
Limitações na Interação	8	[R3, R5, R20, R22, R35, R40, R44, R46]
Viés ao Usuário	6	[R5, R19, R26, R29, R36, R44]
Desinformação	6	[R3, R5, R19, R26, R29, R36]
Falta de Supervisão	6	[R5, R11, R22, R26, R28, R38]
Isolamento Social e Estigmatização	4	[R7, R22, R31, R43]
Riscos para a Saúde Mental	3	[R29, R40, R43]
Interpretação Errônea por Parte dos Usuários	3	[R5, R40, R43]
Dependência e Uso Excessivo	2	[R29, R43]
Privacidade e Segurança de Dados	2	[R19, R31]
Custo ambiental	1	[R42]

Fonte: A autora (2025).

Tabela F16 – Fontes de riscos, segundo vídeos no TikTok

Categoria	Nº	Fontes
Nenhum	16	[T1, T2, T3, T8, T10, T14, T21, T23, T27, T28, T31, T34, T35, T37, T38, T41]
Não é um terapeuta de verdade	15	[T4, T6, T7, T11, T12, T13, T16, T17, T19, T24, T25, T26, T29, T30, T33]
Favorece o lado do usuário	7	[T6, T9, T16, T18, T20, T24, T36]
Intenções espelhadas	6	[T5, T6, T9, T16, T18, T36]

Isolamento	5	[T15, T16, T22, T30, T39]
Preocupações com privacidade	3	[T5, T16, T32]
Dependência	3	[T5, T22, T24]
Uso limitado	1	[T17]
Substituição de relacionamentos reais	1	[T39]
Decadência neurológica	1	[T40]

Fonte: A autora (2025).

Tabela F17 – Fontes de perfil dos usuários, segundo a literatura acadêmica

Categoria	Nº	Fontes
Em Situações Vulneráveis	12	[A2, A5, A6, A9, A10, A12, A13, A14, A16, A18, A20, A22]
População em Geral	7	[A1, A3, A8, A13, A14, A19, A22]
Definidos por Faixa Etária	6	[A2, A6, A12, A15, A17, A18]
Em Contextos Educacionais ou Profissionais	5	[A4, A7, A8, A11, A17]
Com Condições Clínicas Específicas	3	[A4, A10, A21]

Fonte: A autora (2025).

Tabela F18 – Fontes de faixa etária dos usuários mencionados na cobertura jornalística

Categoria	Nº	Fontes
Crianças	2	[N6, N17]
Adolescentes	3	[N3, N16, N21]
Jovens Adultos	4	[N11, N14, N22, N24]
Adultos	11	[N2, N5, N6, N7, N8, N12, N15, N17, N19, N21, N25]
Idosos	3	[N1, N12, N19]
Sem resposta	7	[N4, N9, N10, N13, N18, N20, N23]

Fonte: A autora (2025).

Tabela F19 – Fontes de condições de saúde mental de usuários na cobertura jornalística

Categoria	Nº	Fontes
Ansiedade	8	[N3, N4, N6, N7, N11, N21, N22, N24]
Depressão	6	[N3, N7, N11, N15, N21, N24]
Solidão	5	[N4, N6, N13, N15, N25]

Outros	4	[[N17], [N8], [N4], [N8]]
Luto	3	[N2, N22, N25]
TDAH	2	[N17, N19]
TOC	2	[N17, N21]
Estresse	2	[N3, N5]
Trauma	2	[N2, N7]

Fonte: A autora (2025).

Tabela F20 – Fontes de país de origem de usuários mencionados na cobertura jornalística

Categoria	Nº	Fontes
EUA	10	[N1, N2, N3, N6, N7, N8, N12, N16, N17, N19]
Sem resposta	9	[N4, N5, N9, N10, N11, N13, N15, N20, N23]
Reino Unido	6	[N2, N8, N17, N18, N19, N21]
China	3	[N14, N22, N24]
Vietnã	2	[N8, N19]
Taiwan	1	[N24]
Itália	1	[N6]
Nova Zelândia	1	[N7]

Fonte: A autora (2025).

Tabela F21 – Fontes de observações específicos feitas em postagens no Reddit

Categoria	Nº	Fontes
Prefere terapia com IA a terapia com humanos	14	[R4, R10, R12, R14, R15, R16, R18, R23, R26, R31, R32, R33, R37, R47]
Tem diagnóstico de problema psicológico	14	[R1, R3, R4, R9, R10, R17, R18, R23, R26, R28, R29, R33, R36, R42]
Faz terapia com humano além do uso com IA	12	[R1, R4, R9, R11, R24, R25, R28, R30, R37, R42, R45, R46]
Relata experiência negativa com terapia com humanos	11	[R15, R16, R18, R21, R23, R26, R32, R33, R36, R42, R47]
Relata solidão ou isolamento	6	[R6, R9, R14, R31, R33, R34]
Tem problemas com vício	2	[R13, R36]

Fonte: A autora (2025).

Tabela F22 – Fontes de gênero percebido de usuários nos vídeos sobre o tema no TikTok

Categoria	Nº	Fontes
Feminino	25	[T1, T2, T4, T5, T9, T12, T15, T16, T17, T19, T21, T23, T25, T28, T29, T30, T31, T32, T33, T34, T35, T36, T38, T40, T41]
Masculino	15	[T3, T6, T7, T8, T10, T11, T13, T14, T18, T20, T24, T26, T27, T37, T39]
Não-binário	1	[T22]

Fonte: A autora (2025).

Tabela F23 – Fontes de menções de plataformas específicas nos vídeos sobre o tema no TikTok

Categoria	Nº	Fontes
ChatGPT	28	[T1, T2, T4, T5, T6, T8, T9, T10, T13, T17, T18, T20, T21, T23, T24, T25, T27, T28, T29, T31, T32, T33, T34, T35, T36, T37, T38, T41]
Character.ai	1	[T22]
MetaAI	1	[T19]
Willow	1	[T14]
Nenhum	10	[T3, T7, T11, T12, T15, T16, T26, T30, T39, T40]

Fonte: A autora (2025).

Apêndice G – Roteiro de survey para brasileiros maiores de 18 anos

“Legenda:

(): Escolher apenas uma opção

[]: Escolher uma ou mais opções

Aaa: Resposta eliminatória, por não cumprir os critérios de inclusão ou falhar no teste de atenção

Todas as perguntas foram programadas como obrigatórias, mas a grande maioria possui opções que indicam o desejo de não responder.

Seção 1: Informações Demográficas

Esta seção coleta dados sobre o perfil dos participantes, como idade, gênero e escolaridade, para contextualizar as respostas e permitir uma análise mais detalhada de como diferentes grupos se relacionam com os LLMs.

Qual é a sua idade?

- ☐ *Menos de 18 anos*
- ☐ 18–24 anos
- ☐ 25–34 anos
- ☐ 35–44 anos
- ☐ 45–54 anos
- ☐ 55+ anos
- ☐ Prefiro não responder

Com qual gênero você se identifica?

- ☐ Masculino
- ☐ Feminino
- ☐ Não-binário
- ☐ Prefiro não informar
- ☐ Outro: _____

Qual é o seu grau de escolaridade?

- ☐ Ensino fundamental incompleto
- ☐ Ensino fundamental completo
- ☐ Ensino médio incompleto
- ☐ Ensino médio completo
- ☐ Ensino superior incompleto
- ☐ Ensino superior completo
- ☐ Pós-graduação (lato sensu ou stricto sensu)
- ☐ Prefiro não responder

Você possui ou já teve algum diagnóstico ou condição relacionada à saúde mental?

- ☐ Sim, com diagnóstico clínico
- ☐ Sim, mas sem diagnóstico formal
- ☐ Não
- ☐ Prefiro não responder

Você faz ou já fez algum tipo de acompanhamento psicológico ou psiquiátrico com algum profissional da área?

- ☐ Sim, atualmente faço
- ☐ Já fiz no passado, mas não faço mais
- ☐ Nunca fiz
- ☐ Prefiro não responder

Qual é a sua área de atuação ou estudos?

- ☐ Tecnologia e Computação
- ☐ Ciências Sociais e Humanas
- ☐ Ciências Exatas e Naturais
- ☐ Saúde
- ☐ Educação
- ☐ Artes e Comunicação
- ☐ Negócios e Economia
- ☐ Outro: _____
- ☐ Prefiro não responder

Contextualização

O que são Modelos de Linguagem de Grande Escala (LLMs)?

Modelos de Linguagem de Grande Escala (LLMs) são sistemas de inteligência artificial (IA) capazes de entender e gerar texto em linguagem natural, como se estivessem conversando com uma pessoa. Eles são treinados com grandes quantidades de texto, o que lhes permite responder perguntas, gerar conteúdos e manter conversas de forma fluida e coerente. Ferramentas como o ChatGPT e o DeepSeek são exemplos populares de LLMs, amplamente usados para tarefas como escrever, ajudar com estudos, responder dúvidas e até oferecer apoio emocional. Durante todo o questionário, o termo LLM será usado para se referir a essas ferramentas.

Você já utilizou um Modelo de Linguagem de Larga Escala (LLM), como ChatGPT, DeepSeek, Gemini ou similar?

- ☐ Sim
- ☐ Não

Seção 2: Uso geral de LLMs

Aqui, exploramos a frequência e os motivos que levam os participantes a usar Modelos de Linguagem de Larga Escala (LLMs). A seção também busca entender em que contextos as ferramentas são mais utilizadas.

O que motivou o seu primeiro uso de LLMs?

- ☐ Conteúdo em redes sociais
- ☐ Recomendação de amigos
- ☐ Curiosidade ou interesse pessoal
- ☐ Notícias ou reportagens sobre o assunto
- ☐ Indicação em contexto de estudo ou trabalho
- ☐ Não lembro/Prefiro não responder
- ☐ Outro: _____

Com que frequência você usa LLMs?

- ☐ Diariamente
- ☐ Algumas vezes por semana
- ☐ Algumas vezes por mês
- ☐ Raramente
- ☐ *Nunca usei*
- ☐ Prefiro não responder

Quais LLMs você já utilizou?

- ☐ ChatGPT
- ☐ DeepSeek
- ☐ Gemini (ex-Bard)
- ☐ Character.AI
- ☐ Claude
- ☐ Grok
- ☐ Não lembro/Prefiro não responder
- ☐ Outros: _____

Em que contexto você costuma usar LLMs?

- ☐ Trabalho

- ☐ Estudos
- ☐ Assistência com tarefas diversas (não relacionadas a trabalho/estudos)
- ☐ Curiosidade / entretenimento
- ☐ Conversas sobre questões pessoais
- ☐ Prefiro não responder
- ☐ Outros: _____

Você já conversou com um LLM sobre questões emocionais, como sentimentos, problemas pessoais, solidão, ansiedade etc.?

- ☐ Sim
- ☐ Não
- ☐ *Prefiro não responder*

Seção 3: Uso emocional de LLMs

Nesta seção, investigamos o uso de LLMs para apoio emocional, como desabafo, busca de companhia e alívio de estresse. Queremos entender como essas ferramentas são vistas como recursos para o bem-estar psicológico e social dos usuários.

Com que frequência você usa LLMs para fins pessoais/emocionais?

- ☐ Diariamente
- ☐ Algumas vezes por semana
- ☐ Algumas vezes por mês
- ☐ Raramente
- ☐ *Nunca usei*
- ☐ Prefiro não responder

Para qual(ais) dos seguintes fins você já utilizou LLMs?

- ☐ Desabafo sobre problemas pessoais
- ☐ Conselhos sobre problemas pessoais
- ☐ Opinião neutra sobre dilemas ou conflitos pessoais
- ☐ Companhia/interação social
- ☐ Geração de ideias ou pensamentos criativos
- ☐ Alívio ou distração de estresse ou ansiedade

- ☐ Simulação de um relacionamento afetivo/fraternal
- ☐ Prefiro não responder
- ☐ Outros: _____

Em que situações você recorreu ao LLM?

- ☐ Quando me senti sozinho(a)
- ☐ Quando tive dificuldades emocionais ou psicológicas
- ☐ Quando senti necessidade de apoio emocional
- ☐ Quando me senti estressado(a) ou ansioso(a)
- ☐ Prefiro não responder
- ☐ Outros: _____

Esta é uma verificação de atenção. Para passar no teste, você deve selecionar a alternativa "Responsabilidade". Com base no texto que você leu acima, qual é o princípio a ser selecionado?

- ☐ *Privacidade*
- ☐ *Transparência*
- ☐ *Justiça*
- ☐ Responsabilidade

Você já sentiu que conversar com um LLM foi útil para a sua saúde emocional?

- ☐ Sim, muito útil
- ☐ Sim, um pouco útil
- ☐ Não percebi nenhuma diferença
- ☐ Não, não foi útil
- ☐ Não, foi prejudicial de alguma forma
- ☐ Prefiro não responder

Outros: _____

O que te motivou a procurar um LLM em um contexto pessoal, em vez de conversar com alguém?

- ☐ Disponibilidade
- ☐ Neutralidade das respostas
- ☐ Falta de pessoas de confiança (amigos, familiares)

- ☐ Privacidade e anonimato
- ☐ Evitar desconforto com pessoas conhecidas
- ☐ Rapidez na resposta
- ☐ Facilidade de acesso
- ☐ Falta de acesso a profissionais qualificados
- ☐ Prefiro não responder
- ☐ Outros: _____

Você já sentiu que a interação com LLMs substituiu interações com outras pessoas?

- ☐ Sim, com frequência
- ☐ Às vezes
- ☐ Não, nunca
- ☐ Não tenho certeza/Prefiro não responder
- ☐ Outros: _____

Você sentiu que o LLM conseguiu responder adequadamente às suas emoções?

- ☐ Sim, me senti compreendido (a)
- ☐ Mais ou menos
- ☐ Não
- ☐ Não tenho certeza/Prefiro não responder
- ☐ Outros: _____

Seção 3: Considerações Éticas e Sociais

Aqui, abordamos as preocupações éticas e impactos sociais que o uso de LLMs pode gerar, incluindo questões de privacidade, dependência emocional e possíveis riscos no uso dessas tecnologias para fins terapêuticos.

Você tem preocupações sobre usar LLMs para apoio emocional ou interação social?

- ☐ Sim, muitas preocupações
- ☐ Sim, algumas preocupações
- ☐ Não tenho preocupações
- ☐ Não sei dizer/Prefiro não responder
- ☐ Outros: _____

Quais das seguintes preocupações você acha mais relevantes ao usar LLMs?

- ☐ Não tenho preocupações
- ☐ Privacidade de dados
- ☐ Dependência emocional
- ☐ Falta de empatia real
- ☐ Efeitos negativos na saúde mental a longo prazo
- ☐ Uso inadequado das respostas fornecidas
- ☐ Respostas inadequadas a problemas humanos
- ☐ Enviesamento de respostas para agradar o usuário
- ☐ Outros: _____

Seção 4: Perguntas Abertas

Esta seção permite que os participantes compartilhem suas próprias percepções e experiências sobre o uso de LLMs. As respostas abertas proporcionam um espaço para insights mais profundos e subjetivos que não poderiam ser capturados por perguntas fechadas.

Você acha que usar uma IA para apoio emocional é algo positivo, negativo ou neutro? Explique.

Você acha que a acessibilidade dos LLMs pode trazer benefícios para pessoas que enfrentam dificuldades para encontrar apoio emocional, como em áreas com poucas opções de terapia?

Você gostaria de acrescentar algo sobre suas experiências com IAs como ferramentas de uso emocional?”

ANEXOS

Anexo A – Regras da comunidade r/therapy no Reddit

“Hello, r/therapy!

We have received several reports, comments, and messages regarding AI in our community. We have come to the conclusion to implement an AI policy for our community as outlined below. If you have any questions, comments, or concerns, please do not hesitate to contact us!

Best regards,

r/therapy Mod Team

Policy:

Discussion - We allow discussion of the ethics, impact, and results of the use of AI in therapy and as therapy.

Promotion - While discussion of AI and AI therapy is allowed, promotion of specific sites, tools, or of AI as a replacement for therapy is not. While AI can be a supplemental tool in mental health, it is not currently a safe, effective replacement for therapy.

Example:

Allowed: “I think AI could help the mental health community by doing [x]”

Not Allowed: “Real therapists are all narcissists. AI is the best way to get therapy.”

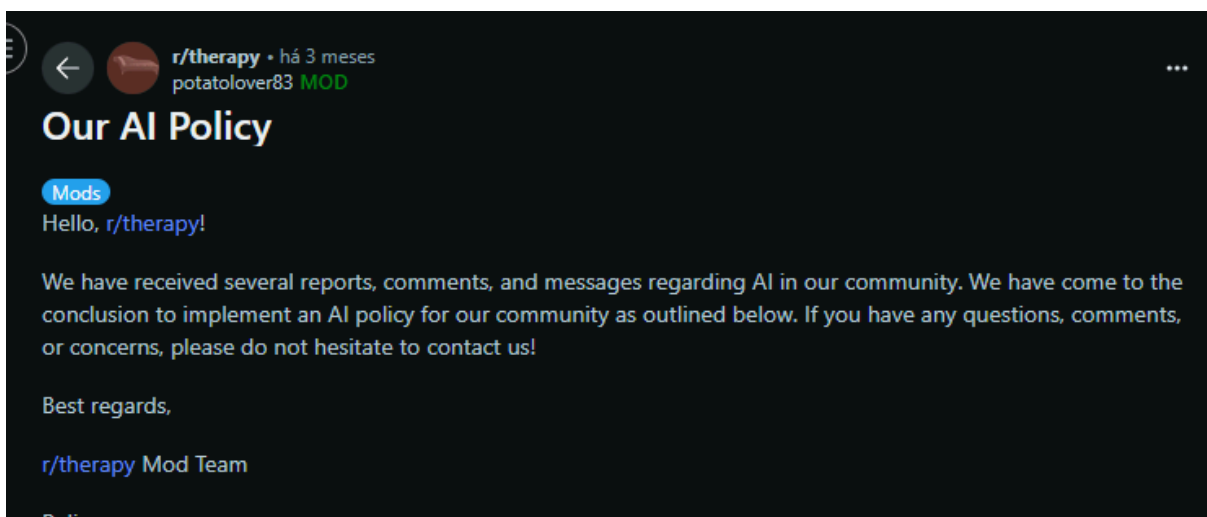
Use - The purpose of r/therapy is for authentic, human interactions. The use of generative AI to write posts or comments is prohibited. You are welcome to use AI to check facts (note: AI does get things wrong), come up with synonyms, and otherwise proofread your content but using AI to fully write your posts/comments is not allowed.

Example:

Allowed: Asking AI for a synonym, fact check, or to have a concept explained

Not Allowed: Pasting a question to AI and then replying with the AI's response.

(Note: these examples are not exhaustive and removal of posts and comments under the AI fall under moderator discretion)”



Fonte: REDDIT. Our AI policy. r/therapy, 2025. Disponível em: https://www.reddit.com/r/therapy/comments/1jxupif/our_ai_policy/. Acesso em: 22 jul. 2025.

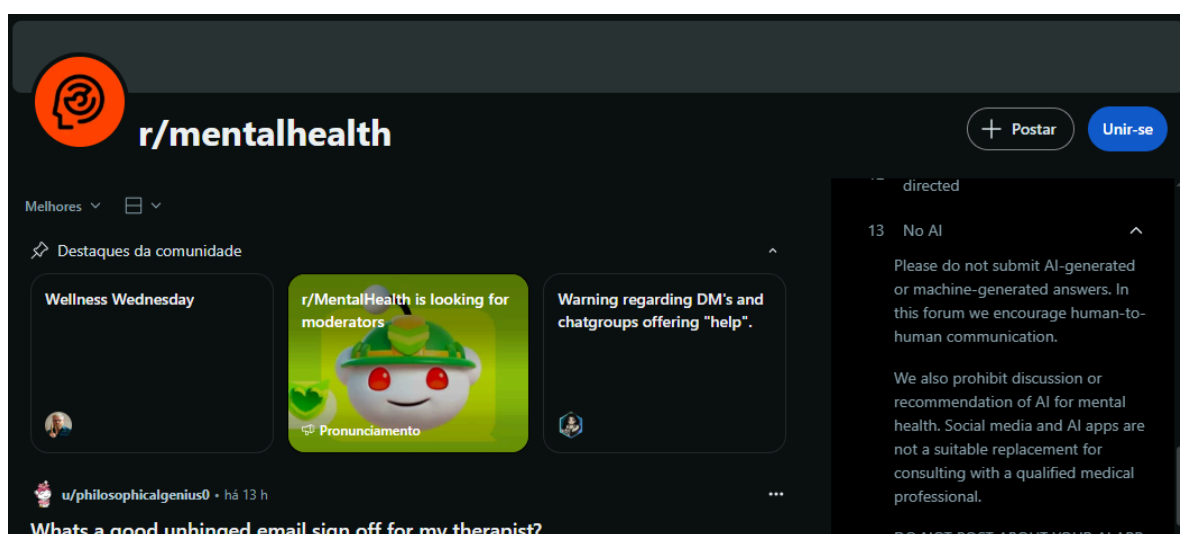
Anexo B – Regras da comunidade r/mentalhealth no Reddit

“13 No AI

Please do not submit AI-generated or machine-generated answers. In this forum we encourage human-to-human communication.

We also prohibit discussion or recommendation of AI for mental health. Social media and AI apps are not a suitable replacement for consulting with a qualified medical professional.

DO NOT POST ABOUT YOUR AI APP.”



Fonte: REDDIT. r/mentalhealth. Disponível em: <https://www.reddit.com/r/mentalhealth/>. Acesso em: 22 jul. 2025.