

NLP Models for Mental Health-Related Text Analysis in Brazilian University Reddit Communities

Tiago C. B. Rocha¹, Leonardo P. Dallagnol¹,
Maicon C. Marques¹, Camila O. Silva¹, Júlia O. Vilaça², Jean R. Ponciano¹

¹Instituto de Ciências Matemáticas e de Computação – Universidade de São Paulo (USP)
São Carlos – SP – Brazil

²Faculdade Ciências Médicas de Minas Gerais, Belo Horizonte – MG – Brazil

{tiagocbr, dallagnolleonardo, maiconchavesmarques, caamilaoliv}@usp.br,
julia_vilaca@cienciasmedicasmg.edu.br, jeanponciano@icmc.usp.br

Abstract. *Online communities have become important spaces for mental health discussions among university students, generating large volumes of text that are difficult to analyze manually at scale. This work investigates automated methods for analyzing emotional and relational aspects of mental health-related text in three Brazilian university Reddit communities. We employed NLP models to evaluate two dimensions: the emotional tone expressed in the text and the presence or absence of supportive social interaction. The results indicate that automated methods are suitable for scalable identification of social support behaviors, while emotional tone analysis requires greater caution and is best applied to clearly polarized expressions rather than fine-grained neutral distinctions.*

1. Introduction

Mental health has become one of the main topics in public health, with a global increase in recent years [Fan et al. 2025]. Data from the World Health Organization indicate that Brazil has one of the highest prevalence rates of anxiety disorders in the world, affecting approximately 9.3% of the population [World Health Organization 2017]. Among the most susceptible groups are students, who face factors such as academic pressure, stress, and social isolation, often associated with anxiety and depression [Li et al. 2025]. In this context, studies have been conducted in Brazil to understand and support students' mental health [Oliveira et al. 2025]. At the same time, social platforms such as REDDIT have become relevant spaces for discussions about mental health, especially because they provide an anonymized environment that reduces barriers to addressing sensitive topics.

One way to obtain relevant information from social platforms is by classifying their textual content, such as user comments on posts. However, due to the large volume of data on these platforms, manual classification is often unfeasible. While both traditional Natural Language Processing (NLP) methods and Large Language Models (LLMs) have been employed for this task, most studies focus primarily on English-language datasets and generalized mental health communities, with limited attention to Portuguese-speaking populations or to university-centered online environments, with their linguistic and contextual specificities [Montejo-Ráez et al. 2024]. In addition, prior work frequently concentrates on either emotion classification or mental disorder detection [Inamdar et al. 2023, Ren et al. 2021, Sweetlin et al. 2025].

This scenario motivates our study, which investigates both emotional tone and supportive interaction in mental-health-related comments posted in Brazilian university communities on REDDIT. Specifically, we present an evaluation of LLMs and traditional NLP models for classifying REDDIT comments according to two complementary and equally important dimensions: *emotional valence* [Russell 1980] and *social support function* [House 1981]. In our context, emotional valence refers to the predominant emotional tone expressed in a comment, while social support concerns the presence of supportive or non-supportive interactional behavior in comments. By jointly analyzing these dimensions, we aim to provide mental health specialists and researchers with a more comprehensive understanding of how emotional distress and supportive dynamics manifest in online interactions, potentially aiding the identification of vulnerable individuals and the characterization of supportive communication patterns in digital environments.

2. Background

Reddit. REDDIT¹ is a social media platform organized into discussion communities, called subreddits, where users create posts and comments about a wide range of topics, including mental health. Due to the large amount of textual and emotionally expressive content available, REDDIT data has been widely used to construct datasets for depression and mental health analysis. Examples include the datasets created by Sampath and Durairaj [Sampath and Durairaj 2022] and by Herculano et al. [Herculano et al. 2024], which comprise REDDIT posts for depression detection, using content collected from the general public and without focusing on specific demographic or social groups.

Emotional Valence. The definition of emotional valence comes from the circumplex model of affect proposed by James A. Russell [Russell 1980], which organizes emotional states according to a pleasant–unpleasant dimension. In this context, textual content, for example, comments from social media platforms such as REDDIT, can be classified as positive, negative, or neutral according to their predominant emotional tone. Positive comments typically express emotions such as hope, satisfaction, encouragement, or joy, while negative comments are associated with sadness, frustration, fear, anger, or emotional distress. Neutral comments correspond to messages with no clear emotional polarity, usually presenting descriptive, objective, or ambiguous content.

Social Support Function. House’s social support theory [House 1981] allows for interpreting supportive behavior in interactions, which describes how social relationships can contribute to coping processes during stressful situations. Under this perspective, textual content could be evaluated according to the type of interpersonal assistance offered to the content creator. In this sense, online communication could be supportive (e.g., understanding, reassurance, advice, or motivational intent), or non-supportive (e.g., indifference, criticism without constructive intent, lack of engagement, or absence of helpful interaction). Finally, social support is not limited to direct assistance or explicit advice. It can also emerge, for example, through expressions of belonging and emotional validation.

3. Related work

Research on mental health analysis in online communities has increasingly employed NLP methods to investigate psychological and emotional patterns expressed in social me-

¹<https://www.reddit.com/>

dia texts. Several studies have explored the use of machine learning and transformer-based models to detect indicators of stress, depression, and psychological vulnerability in platforms such as REDDIT. Inamdar et al. [Inamdar et al. 2023] investigated stress detection using multiple NLP representations and demonstrated the scalability of automated approaches for mental health monitoring. Similarly, Ren and colleagues [Ren et al. 2021] proposed an emotion-aware neural architecture for depression detection on REDDIT, highlighting the relevance of emotional cues in identifying psychological distress. More recently, Sweetlin et al. [Sweetlin et al. 2025] introduced a hybrid BERT-LSTM model capable of detecting multiple mental health conditions while also employing generative AI techniques to rewrite harmful content into supportive messages, reinforcing the potential use of such architectures in online mental health environments.

Emotion analysis has also become an important dimension in computational mental health research, particularly in studies seeking to understand how affective expressions relate to psychological conditions. For example, the study by Guo et al. [Guo et al. 2021] modeled emotional transitions in online discussions to characterize mental disorders, arguing that emotional representations can provide robust indicators of mental health states beyond purely lexical features. In parallel, Atapattu et al. [Atapattu et al. 2022] introduced the EmoMent corpus, an emotion-annotated dataset focused on mental health discussions, demonstrating the complexity and ambiguity involved in labeling emotional expressions in psychologically sensitive contexts.

Social support and interpersonal interaction have also been investigated in online communities. Prior studies have shown that social media platforms can function as spaces for emotional validation, coping assistance, and community-based support [Ruihua et al. 2025]. In this context, Low et al. [Low et al. 2020] analyzed mental health support groups on REDDIT during the COVID-19 pandemic and demonstrated how NLP techniques can reveal patterns of vulnerability and supportive behavior at scale. Newer approaches have explored the use of large language models for mental health-related social discussions, including detecting supportive interactions and providing contextual recommendations [Aggarwal et al. 2024]. These studies highlight the relevance of relational dimensions in online mental health environments and support the use of automated methods to identify supportive social behaviors in large conversational datasets.

Despite advances, important gaps remain in the literature, particularly regarding studies that jointly investigate emotional valence and supportive interactions in Brazilian Portuguese mental health-related communication among university students. Our work contributes to the discussion by investigating the performance of both traditional NLP techniques and LLMs in analyzing emotional and social support dimensions in mental-health-related comments posted in Brazilian university communities on REDDIT.

4. Data and Methods

The methodological pipeline adopted in this study consists of four main stages, as shown in Figure 1: (i) data collection and preprocessing; (ii) manual annotation by an annotator with a background in psychology, following two analytical criteria: emotional valence and social support function; (iii) automated classification with traditional NLP models and LLMs; and (iv) performance evaluation. Details are provided below.

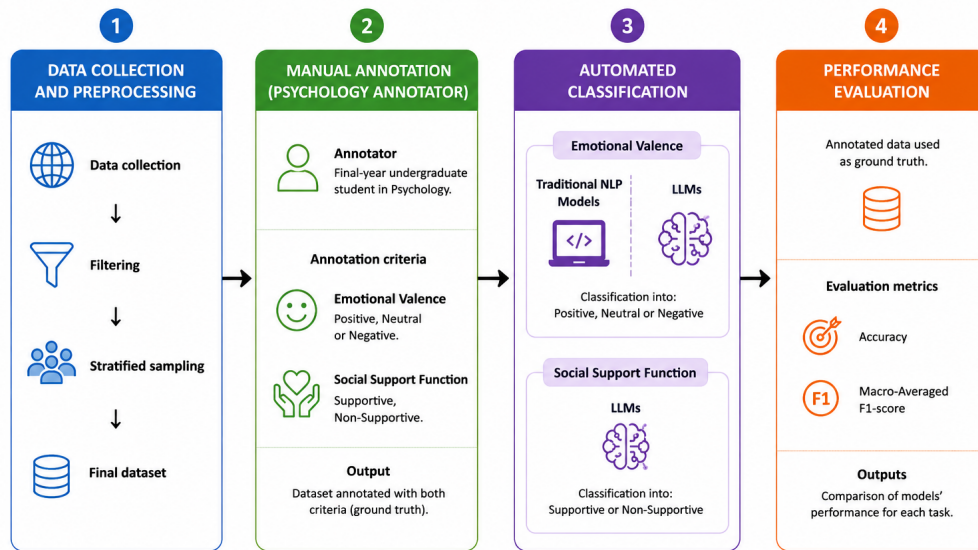


Figure 1. Methodological pipeline adopted in this study.

4.1. Dataset

Initially, the data were collected from BrStudentMH [Marques et al. 2026]², a public dataset containing posts and comments related to mental health within the Brazilian university context. This dataset includes content from the subreddits `r/usp` (65 thousand members), `r/faculdadeBR` (84 thousand members), and `r/askacademico` (13 thousand members). In addition, the posts are organized into keywords corresponding to mental health-related topics, which are used to categorize each post. In total, the dataset comprises 943 posts published between January 1, 2024, and April 20, 2025, as well as 15,680 comments published between January 1, 2024, and September 15, 2025. Posts and comments are written in Brazilian Portuguese.

A filtering step was then applied to the initial comment set to ensure analytical consistency. Comments marked as deleted or removed by the platform (370 instances), replies to other comments (7,981), and comments made by the original post author (2,821) were excluded. As overlap among these categories was possible, the effective number of removals was 8,176 comments, resulting in a filtered dataset of 7,504 comments.

Next, a stratified sampling of approximately 10% of this filtered comment dataset was conducted to create a manageable subset for manual annotation, as annotating the entire dataset would be impractical. The sampling procedure was implemented in Python using the PANDAS library. Each comment was associated with the keyword of its corresponding post, and the proportion of comments per keyword was computed. Comments were then randomly sampled while preserving these proportions. The final dataset used in this study corresponds to this stratified comment sample, comprising 745 comments: 66 from `r/askacademico`, 298 from `r/USP`, and 381 from `r/faculdadeBR`.

4.2. Expert Annotation

Expert annotation was performed by a final-year undergraduate psychology student, following the two independent analytical dimensions discussed in Section 2: emotional va-

²<https://github.com/MaiconChavesMarques/BrStudentMH-dataset>

lence and social support function. The annotation process was used to construct a reliable ground truth dataset for evaluating automated classification models across both dimensions, enabling a comparison between human judgments and machine-based predictions.

The resulting annotated dataset exhibited the following distribution of classes. For emotional valence, 104 comments (13.96%) were labeled as neutral, 298 (40.0%) as negative, and 343 (46.04%) as positive. For social support, 217 comments (29.13%) were labeled as not providing support, while 528 (70.87%) were labeled as providing support.

4.3. Models Employed

To evaluate the feasibility of automated classification, different LLMs were employed, including GROK 4.2, GPT-5.3, GEMINI 3, and DEEPSEEK V3. Each model was queried using an identical instruction-based prompt, designed to produce both emotional valence and social support function classifications in a structured output format. The prompt explicitly defined both tasks, provided theoretical grounding for each dimension, and included label definitions, decision rules, and output formatting constraints to ensure consistency across models.

In addition to LLMs, traditional NLP models were also evaluated for the emotional valence task. This set of models includes both transformer-based and lexicon-based approaches commonly used for sentiment analysis. The transformer-based models include four pretrained sentiment analysis models from Hugging Face: `cardiffnlp/twitter-xlm-roberta-base-sentiment` [Barbieri et al. 2022], `tabularisai/multilingual-sentiment-analysis`³, `ggrazzioli/cls_sentimento_sebrae`⁴, `lipaoMai/bert-sentiment-model-portuguese`⁵, while the lexicon-based approach is represented by the `LeIA` model [Almeida 2018], a Portuguese adaptation of the `VADER` sentiment analysis framework [Hutto and Gilbert 2014].

The use of traditional NLP models was restricted to the emotional valence task due to their alignment with the sentiment analysis paradigm, for which they are specifically designed and widely benchmarked. In contrast, social support classification involves more complex interpersonal and contextual interpretation, including the identification of intent, empathy, and relational function, which are not directly captured by conventional sentiment-oriented models. Within the scope of this study, the selected traditional NLP models do not address this task, which would typically require task-specific supervised training on annotated data.

4.4. Evaluation Metrics

Model performance was evaluated using accuracy and macro-averaged F1-score (F1-macro), using the manually annotated dataset as ground truth. Accuracy measures the overall proportion of correctly classified instances across all classes, defined as $Accuracy = \frac{TP+TN}{TP+TN+FP+FN}$, where TP, FP, TN, and FN denote true positives, false positives, true negatives, and false negatives, respectively.

The F1-score is based on two complementary measures: precision (P) and recall (R). Precision quantifies the proportion of correctly predicted positive instances among

³<https://huggingface.co/tabularisai/multilingual-sentiment-analysis>

⁴https://huggingface.co/ggrazzioli/cls_sentimento_sebrae

⁵<https://huggingface.co/lipaoMai/bert-sentiment-model-portuguese>

all instances predicted as positive, i.e., how reliable the positive predictions are. Recall measures the proportion of correctly identified positive instances among all actual positive cases, i.e., how well the model captures the true instances of a class. The F1-score combines these two aspects into a single metric by computing their harmonic mean, defined as $F1 = 2 \cdot \frac{P \cdot R}{P + R}$, where $P = \frac{TP}{TP + FP}$ and $R = \frac{TP}{TP + FN}$. This combination penalizes extreme imbalances between precision and recall, providing a more robust measure of performance than either metric alone. The macro-averaged F1-score is computed as $F1_{macro} = \frac{1}{K} \sum_{i=1}^K F1_i$, where K denotes the number of classes, giving equal weight to each class regardless of their frequency.

This combination of metrics was selected to provide a comprehensive evaluation of model performance. While accuracy captures overall correctness, it may be biased in the presence of class imbalance. In contrast, the macro-averaged F1-score ensures that performance across all classes is equally represented, making it appropriate for both emotional valence classification and social support detection tasks, where class distributions are not uniform.

5. Experimental Results

5.1. Emotional Valence Classification

Emotional valence classification proved to be challenging, as depicted in Table 1. The best performance was achieved by the GEMINI model, with 57.6% accuracy and 55.3% F1-macro. This result slightly outperforms traditional models such as LeIA (51.8% accuracy and 48.6% F1-macro), indicating incremental gains from the use of LLMs. Among non-LLM approaches, LeIA achieved the best overall performance, suggesting that lexicon-based methods remain competitive in this task.

The overall low performance on this task can be attributed to the highly contextual and ambiguous nature of emotional expressions in texts [Barrett 2017] and the resulting semantic proximity between classes, particularly between neutral and polarized categories. A tendency of LLMs to overpredict the neutral class was also observed, suggesting a conservative strategy in the presence of semantic uncertainty. This is illustrated in Figure 2a, which shows the confusion matrix for GEMINI. On the other hand, lexicon-based models exhibit a more polarized prediction pattern, as illustrated in Figure 2b for LeIA, likely driven by their reliance on explicit sentiment-bearing lexical markers.

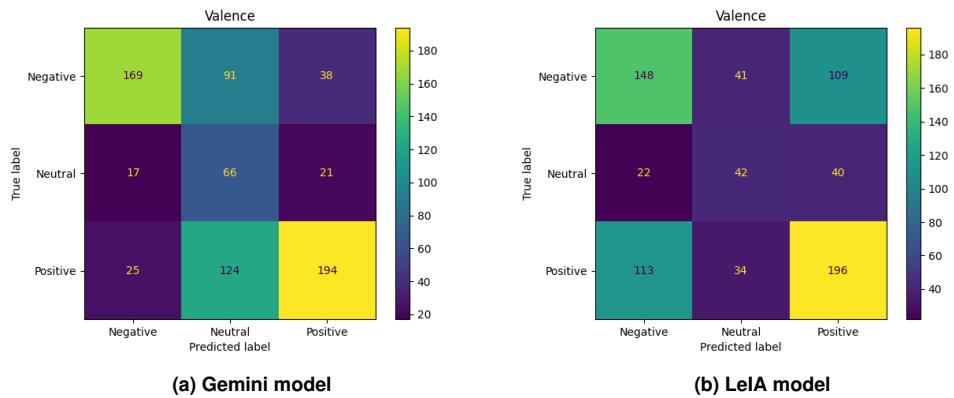
Despite the overall lower performance, the models achieved better results in the positive and negative classes, with most errors concentrated in the neutral class (Table 1). This pattern suggests that model outputs may be more reliably used in a selective manner, focusing on instances with clearer emotional polarity, while treating neutral or ambiguous cases with greater caution or requiring additional validation.

5.2. Social Support Classification

The LLMs' performance for the social support classification task is presented in Table 2. LLMs demonstrated a high capacity to determine whether a comment supports the post's author. The best-performing model was GROK, achieving 85.9% accuracy and 82.2% F1-macro. The results indicate that LLMs can capture relational and contextual aspects of communication, such as empathy and intent to help, suggesting their suitability for reducing reliance on manual labeling, especially in large datasets.

Table 1. Performance of models in emotional valence classification

Type	Model	Acc (%)	F1-macro (%)	F1 Neg. (%)	F1 Neu. (%)	F1 Pos. (%)
LLM	GPT	54.9	53.8	65.7	33.9	61.8
	Gemini	57.6	55.3	66.4	34.3	65.1
	Grok	53.6	52.1	64.9	32.4	59.1
	DeepSeek	46.6	47.2	60.2	33.1	48.2
NLP	LeIA	51.8	48.6	50.9	38.0	57.0
	CardiffNLP	41.5	39.2	60.9	26.0	30.6
	ggrazzioli	48.7	45.4	55.7	26.8	53.7
	lipaoMai	29.7	25.2	2.0	26.6	47.1
	tabularisai	44.3	40.1	56.5	26.6	37.2

**Figure 2. Confusion matrices for emotional valence classification.**

Such performance can be partly explained by the structure of the task itself. Unlike emotional valence classification, which involves three classes with often ambiguous boundaries (particularly between neutral and positive or neutral and negative classes), social support classification is formulated as a binary problem with more clearly defined and conceptually distinct categories.

6. Large-Scale Automated Classification

To further assess the applicability of automated classification at scale, the best-performing models for each task were applied to the full filtered dataset (7,504 comments): GROK for social support classification and GEMINI for emotional valence classification. Table 3 presents the overall distribution of predicted labels. The results show a relatively balanced distribution of emotional valence, with a slight predominance of negative comments (37.59%), followed by positive (35.47%) and neutral (26.93%). In contrast, social support is markedly more prevalent, with 71.82% of comments classified as supportive.

To contextualize the results across online communities, the r/USP subreddit was analyzed separately due to the relevance of the University of São Paulo (USP) in the Brazilian higher education context [Times Higher Education 2025], and compared with the other two subreddits (Table 3). The USP subset exhibits fewer negative comments (34.84%) compared to the others (39.45%), alongside a slightly higher proportion of positive and neutral instances. This may suggest a relatively less negative or more balanced emotional tone within the USP community. In terms of social support, both contexts show high levels of supportive interactions, with only minor variation (70.81% in USP vs. 72.49% in other subreddits). This indicates that supportive behavior is a consistent characteristic across communities.

Table 2. Performance of models in social function classification

Model	Acc (%)	F1-macro (%)	F1 NoSupp. (%)	F1 Supp. (%)
GPT	82.6	80.4	73.8	86.9
Gemini	85.1	82.2	74.9	89.4
Grok	85.9	82.2	74.1	90.3
DeepSeek	83.9	81.8	75.7	88.0

Table 3. Distribution and comparison of predicted labels across datasets

Metric	Class	Full Dataset		Comparison (%)	
		Count (n)	Percentage (%)	USP	Others
Valence	Negative	2821	37.59	34.84	39.45
	Positive	2662	35.47	37.13	34.36
	Neutral	2021	26.93	28.03	26.19
Support	Support	5389	71.82	70.81	72.49
	NoSupport	2115	28.18	29.19	27.51

However, these results should be interpreted with caution for emotional valence classification. As discussed in Section 5.1, performance on this task was lower, especially due to difficulties in distinguishing neutral from polarized expressions. In this sense, large-scale emotional valence analysis is better understood as indicative of general patterns rather than as a precise labeling of individual comments.

7. Limitations

The traditional NLP models used in this work for emotional valence classification were originally designed for sentiment analysis tasks. Although sentiment analysis and emotional valence are related concepts, they are not strictly equivalent: sentiment analysis is commonly used as polarity classification (positive, negative, neutral) in classical NLP approaches [Pang and Lee 2008], whereas emotional valence involves a broader interpretation of subjective emotional expression in user-generated text, which is inherently ambiguous and highly contextual [Nandwani and Verma 2021, Barrett 2017]. As a result, a potential mismatch exists between the original training objectives of these models and the task they were submitted to. This may help explain the low performance observed for traditional NLP approaches in the emotional valence classification task.

Another limitation of this study is the reliance on a single expert annotator to establish the ground truth. Given the subjective nature of emotional valence and social support labeling, relying on a single annotator may have introduced individual biases into the dataset, potentially affecting model evaluation.

8. Conclusion and Future Work

In this work, the use of NLP models for analyzing mental health-related comments from Brazilian university REDDIT communities was evaluated. The results indicate that LLMs are particularly effective for social support classification, achieving strong performance in identifying supportive interactions. In contrast, emotional valence classification proved to be considerably more challenging, with both LLMs and traditional NLP models showing limited performance, particularly in distinguishing neutral from polarized expressions.

Future work includes investigating alternative strategies for emotional valence classification, such as prompt refinement, ensemble approaches, and confidence-based filtering, as well as expanding the analysis to other online communities and platforms. Additionally, future studies will also involve multiple annotators and report inter-annotator

agreement measures to strengthen the reliability of the ground-truth labels. We also intend to incorporate statistical significance analyses when comparing LLMs, allowing for stronger conclusions regarding performance differences among the models.

Acknowledgments

This research was supported by the Brazilian agencies Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq, grant 144405/2025-3), Coordenação de Aperfeiçoamento de Pessoal de Nível Superior (CAPES, Finance Code 001), São Paulo Research Foundation (FAPESP, grant No. 2024/13328-9), and University of São Paulo (PUB-USP and PRPI Ordinance No. 1032, Process 22.1.09345.01.2).

Artificial Intelligence Usage

CHATGPT and GRAMMARLY were used for grammar revision, and CHATGPT was used to create Fig. 1. The authors assume full responsibility for the content of the work.

References

- Aggarwal, V., Thukral, S., Patel, K., and Chatterjee, A. (2024). Leveraging llms for mental health: Detection and recommendations from social discussions. In *2024 IEEE/WIC Int. Conf. on Web Intel. and Intellig. Agent Tech. (WI-IAT)*, pages 350–354. IEEE.
- Almeida, R. J. A. (2018). Leia - léxico para inferência adaptada. <https://github.com/rafjaa/LeIA>.
- Atapattu, T., Herath, M., Elvitigala, C., de Zoysa, P., Gunawardana, K., Thilakaratne, M., De Zoysa, K., and Falkner, K. (2022). Emoment: An emotion annotated mental health corpus from two south asian countries. In *Proceedings of the 29th International Conference on Computational Linguistics*, pages 6991–7001.
- Barbieri, F., Espinosa Anke, L., and Camacho-Collados, J. (2022). XLM-T: Multilingual language models in Twitter for sentiment analysis and beyond. In *Proceedings of the Thirteenth Language Resources and Evaluation Conference*, pages 258–266, Marseille, France. European Language Resources Association.
- Barrett, L. F. (2017). *How emotions are made: The secret life of the brain*. Pan Macmillan.
- Fan, Y., Fan, A., Yang, Z., and Fan, D. (2025). Global burden of mental disorders in 204 countries and territories, 1990–2021: results from the global burden of disease study 2021. *BMC Psychiatry*, 25(1):486.
- Guo, X., Sun, Y., and Vosoughi, S. (2021). Emotion-based modeling of mental disorders on social media. In *Int. Conf. on Web Intel. and Intellig. Agent Tech.*, pages 8–16.
- Herculano, A., de Paula, T.-H., Fernandes, D., and Rego, A. (2024). DepreRedditBR: Um conjunto de dados textuais com postagens depressivas no idioma português brasileiro. In *VI Dataset Showcase Workshop*, pages 77–90, Porto Alegre, RS, Brasil. SBC.
- House, J. (1981). *Work Stress and Social Support*. Addison-Wesley Series in Health Education. Addison-Wesley Publishing Company.
- Hutto, C. J. and Gilbert, E. (2014). VADER: A parsimonious rule-based model for sentiment analysis of social media text. In *Int. AAAI Conf. on Web and Social Media*, volume 8, pages 216–225.

- Inamdar, S., Chapekar, R., Gite, S., and Pradhan, B. (2023). Machine learning driven mental stress detection on reddit posts using natural language processing. *Human-Centric Intelligent Systems*, 3(2):80–91.
- Li, Q., Li, J., and Fan, Y. (2025). Addressing mental health in university students: a call for action. *Frontiers in Public Health*, 13:1614999.
- Low, D. M., Rumker, L., Talkar, T., Torous, J., Cecchi, G., and Ghosh, S. S. (2020). Natural language processing reveals vulnerable mental health support groups and heightened health anxiety on reddit during covid-19: Observational study. *Journal of medical Internet research*, 22(10):e22635.
- Marques, M. C., Silva, C. O., Dallagnol, L. P., Rocha, T. C. B., and Ponciano, J. R. (2026). BrStudentMH: A Reddit dataset for mental health analysis in the Brazilian university context. In *VIII Dataset Showcase Workshop*. accepted.
- Montejo-Ráez, A., Molina-González, M. D., Jiménez-Zafra, S. M., Ángel García-Cumbreras, M., and García-López, L. J. (2024). A survey on detecting mental disorders with natural language processing: Literature review, trends and challenges. *Computer Science Review*, 53:100654.
- Nandwani, P. and Verma, R. (2021). A review on sentiment analysis and emotion detection from text. *Social Network Analysis and Mining*, 11(1):81.
- Oliveira, R. A., Vêras, R. M., and Araújo, T. M. d. (2025). Assistência estudantil e promoção da saúde mental nas universidades federais brasileiras: uma revisão integrativa. *Avaliação: Revista da Avaliação da Educação Superior (Campinas)*, 30:e025030.
- Pang, B. and Lee, L. (2008). Opinion mining and sentiment analysis. *Foundations and Trends in Information Retrieval*, 2(1–2):1–135.
- Ren, L., Lin, H., Xu, B., Zhang, S., Yang, L., and Sun, S. (2021). Depression detection on reddit with an emotion-based attention network: algorithm development and validation. *JMIR medical informatics*, 9(7):e28754.
- Ruihua, L., Hassan, N. C., Qiuxia, Z., Sha, O., and Jingyi, D. (2025). A systematic review on the impact of social support on college students’ wellbeing and mental health. *Plos one*, 20(7):e0325212.
- Russell, J. A. (1980). A circumplex model of affect. *Journal of Personality and Social Psychology*, 39(6):1161–1178.
- Sampath, K. and Durairaj, T. (2022). Data set creation and empirical analysis for detecting signs of depression from social media postings. In *International Conference on Computational Intelligence in Data Science*, pages 136–151. Springer.
- Sweetlin, J. D., Vaishnavi, S., Keshavan, A., and Palanisamy, R. (2025). Detection of mental illness from reddit data using deep learning and mitigation using generative ai. In *2025 Int. Conf. on Computing, Intelligence, and Application (CIACON)*, pages 1–8.
- Times Higher Education (2025). Latin america university rankings 2026. <https://www.timeshighereducation.com/world-university-rankings/2026/latin-america-university-rankings>. Acesso em: 11 dez. 2025.
- World Health Organization (2017). *Depression and Other Common Mental Disorders: Global Health Estimates*. World Health Organization, Geneva.