



Deposited via The University of Sheffield.

White Rose Research Online URL for this paper:

<https://eprints.whiterose.ac.uk/id/eprint/244264/>

Version: Published Version

Proceedings Paper:

Chen, Z., Preiss, J. and Bath, P.A. (2025) Detecting changes in mental health status via Reddit posts in response to global negative events. In: Angelova, G., Kunilovskaya, M., Escribe, M. and Mitkov, R., (eds.) Proceedings of the 15th International Conference on Recent Advances in Natural Language Processing - Natural Language Processing in the Generative AI Era. 15th International Conference on Recent Advances in Natural Language Processing (RANLP 2025), 08-10 Sep 2025, Varna, Bulgaria. Incoma Ltd. Shoumen, BULGARIA, pp. 227-233. ISBN: 9789544520984. ISSN: 2603-2813.

<https://doi.org/10.26615/978-954-452-098-4-027>

Reuse

This article is distributed under the terms of the Creative Commons Attribution (CC BY) licence. This licence allows you to distribute, remix, tweak, and build upon the work, even commercially, as long as you credit the authors for the original work. More information and the full terms of the licence here:

<https://creativecommons.org/licenses/>

Takedown

If you consider content in White Rose Research Online to be in breach of UK law, please notify us by emailing eprints@whiterose.ac.uk including the URL of the record and the reason for the withdrawal request.

Detecting Changes in Mental Health Status via Reddit Posts in Response to Global Negative Events

Zenan Chen, Judita Preiss, Peter A Bath

School of Information, Journalism and Communication, University of Sheffield

{zchen249, judita.preiss, p.a.bath}@sheffield.ac.uk

Abstract

Detecting population-level mental health responses to global negative events through social media language remains understudied, despite its potential for public health surveillance. While pretrained language models (PLMs) have shown promise in mental health detection, their effectiveness in capturing event-driven collective psychological shifts – especially across diverse crisis contexts – is unclear. We present a prototype evaluation of three PLMs for identifying population mental health dynamics triggered by real-world negative events. We introduce two novel datasets specifically designed for this task. Our findings suggest that DistilBERT is better suited to the noisier global negative events data, while MentalRoBERTa shows the validity of the method on the Covid-19 lockdown tidier data. SHAP interpretability analysis of 500 randomly sampled posts revealed that mental-health related vocabulary (anxiety, depression, worthlessness) emerged as the most influential linguistic markers for mental health classification.

1 Introduction

Global events characterized by war, violence, discrimination, and political uncertainty can significantly impact individuals worldwide (Moitra et al., 2023). Such disturbing and traumatic world news can affect mental health even among those not directly involved in these events (Thompson et al., 2019).

Mental health datasets predominantly consist of social media content (Mauriello et al., 2021). These datasets use membership in specific mental health communities as proxy indicators, for example, Shing et al. (2018) identify individuals with depression based on their participation in depression-focused subreddits.

This study attempts to investigate the population-level psychological patterns triggered by global

negative events using PLMs. To this end, this paper introduces an event selection framework – a semi-automated pipeline combining Wikipedia’s event catalog with Google Trends-driven significance ranking. This study resolves the temporal incomparability of search trends (due to undisclosed absolute volumes) through cross-event normalization. This enables relatively equitable comparison of events across disparate periods (e.g., 2010 earthquakes vs. 2022 wars), establishing an adaptable methodology to identify globally impactful events. This work also validates PLMs for detecting population-level mental health language changes following negative events, with interpretability analysis confirming mental health terms as primary classification markers.

2 Related Work

2.1 Detecting the Dynamics of Mental Health

Monitoring changes in individuals’ mood over time is essential for understanding and managing mental health conditions (Shalom and Aderka, 2020). The 2022 CLPsych Shared Task (Tsakalidis et al., 2022) focused on identifying moments of mood shifts – specifically transitions between positive and negative states – in longitudinal social media posts. More recently, the CLPsych 2025 Shared Task (Tseriotou et al., 2025) continued this line of work, emphasizing the importance of modeling mental health dynamics across user timelines.

Psychological research has established a negativity bias – the tendency for negative experiences and expressions to be more salient than positive ones (Baumeister et al., 2001). Given the success of RoBERTa (Liu et al., 2019), for example Chakraborty et al. (2025) successfully fine-tuned the model to classify adaptive versus maladaptive self-states, this is the model adopted in this work.

2.2 Detecting Mental Health in Responses to Global Negative Events

Prior work has demonstrated the potential of NLP for monitoring population-level mental health impacts during large-scale crises. Studies have tracked emotional shifts on social media (Lwin et al., 2020) and detected emerging psychological issues in helpline conversations (Raveau et al., 2023), providing clear evidence of crisis-induced mental health deterioration across populations. These approaches illustrate how PLMs and large-scale textual data can effectively monitor population-level mental health trends, but have primarily focused on single global crises.

However, to our knowledge, there is no work on the effect on population level mental health state of global negative events, such as natural disasters, wars, or social unrest. To bridge this gap, we explore the ability of PLMs to identify text written before and after global negative event.

3 Dataset Creation

For this investigation, the focus will be on social media posts, specifically Reddit due to its availability, surrounding global negative events.

3.1 Reddit Corpus Collection

Reddit is an open social media platform, providing anonymous space for users to discuss stigmatic topics and self-report personal issues.

Initial data collection Following previous works (Shing et al., 2018; Ji et al., 2021; Inna and Çagri, 2018), individuals likely to be suffering from depression are identified by extracting content from depression-related subreddits. These communities are known to attract users who either self-identify as experiencing depression or report having received a clinical diagnosis. Users were obtained in these communities (r/depression, r/depressed, and r/depression_help) between 2020 and 2024, from a publicly available Reddit dataset (stuck_in_the_matrix et al., 2025). Then, complete footprints (posts and comments from all subreddits) of these users were collected with the official API. Finally, a minimum activities threshold of 100 was set to filter out non-active users. This initial dataset contains data from 90,779 unique users contributing 54,525,664 submissions across 132,282 subreddits.

Restriction to English text To minimize noise stemming from the use of other languages, this study restricted analysis to English texts. Prior to language detection, non-linguistic content, including URLs, email addresses, username mentions, and other metadata elements were removed. Then FastText¹, an open-source library for text representation and classification, was employed to determine the most likely language of each post.

Exclusion of bot accounts Automated accounts (“bots”) were identified and excluded to ensure linguistic authenticity. The publicly available bot detection algorithm² analyzed username patterns (e.g., containing keywords like “bot”), comment patterns (e.g., extremely high posting frequency or repetitive structure), and content repetitiveness (e.g., identical or near-identical comments across posts). Accounts matching these patterns were flagged and excluded from further analysis.

3.2 Global Negative Events Time Period Determination

Since the focus is on using PLMs detecting a change in mental health state surrounding a negative global event, both these events and their occurrence time period must be identified.

Global events collection This study downloaded 1,094 global events and their descriptions from Wikipedia³, including major political elections, natural disasters, economic crises, sporting competitions, cultural ceremonies, technological breakthroughs, public health emergencies, military conflicts, and significant social movements that occurred worldwide during the study period.

Keywords extraction To identify the most significant events, Google Trends was used based on the assumption that events more frequently searched on Google at the worldwide level should be more influential globally. To automatically extract keywords from Wikipedia event descriptions for Google Trends data collection, Llama3.3-70B⁴

¹<https://fasttext.cc/>

²The bot detection algorithm is available at <https://github.com/scottenriquez/BotDetection-Algorithm>

³For example: <https://en.wikipedia.org/wiki/2020> from 2020 to 2024. Similar pages exist for 2021-2024, following the format <https://en.wikipedia.org/wiki/2021>, covering the period 2020-2024.

⁴<https://huggingface.co/meta-llama/Llama-3.3-70B-Instruct>

and Command R+⁵ models were selected based on their established capabilities in text generation tasks (Rodrigues et al., 2025; Liu et al., 2024).

After computing and comparing the gaps between the event occurrence dates sourced from Wikipedia and peak search weeks identified through their respective keywords, Llama3.3-70B was finally selected for its superior temporal accuracy. Llama3.3-70B successfully identified events with smaller time discrepancies in 9/20 test cases compared to Command R+’s 5/20. Additionally, 8/20 of Llama3.3-70B’s cases showed keyword relevance for Google Trends search but with longer time discrepancies, while only 3/20 cases were unsuitable for Google Trends data collection.

Google Trends data collection Since Google Trends provides only relative search interest rather than absolute volumes, and limits comparisons to a maximum of five keywords (Figure 1), a novel normalization approach to compare global events was developed. The Google Trends historical data for each event were downloaded via DataforSEO API⁶. For each event, its peak search week was identified and pairwise comparisons with three baseline keywords: “weather” (consistent daily interest), “covid-19” (major global event), and “black friday” (predictable seasonal pattern) was conducted.

Google Trends returns three distinct types of values. First, “null” indicates insufficient data for the specific event within the given time period. Second, “<1” indicates that the event was searched within the given period but at volumes too small to reach the minimum reporting threshold for any search data to be displayed. To make sure the “<1” has its numerical meaning, the number of “0.5” is assigned to the value “<1”. As it will be used in the combination formulas, it needs to have a numeric representation, and the number “0.5” distinguishes it from 0 and it is different to 1. Third, numerical values between 1 and 100 represent properly quantified relative search interest that can be directly used. By standardizing comparative ratios and averaging relative interest scores, a metric was established, enabling cross-event popularity comparison (Table 1).

While this methodology effectively identified globally significant events, it did not optimally identify specifically negative events as intended. For

⁵<https://huggingface.co/CohereLabs/c4ai-command-r-plus-4bit>

⁶<https://dataforseo.com/>

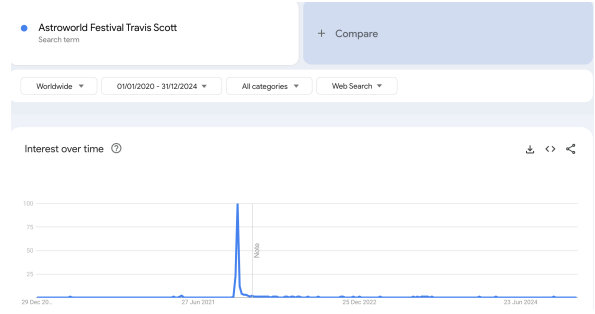


Figure 1: Search interest trends for “Astroworld Festival Travis Scott” worldwide from January 1, 2020, to December 31, 2024. Data source: Google Trends.

example, the keywords “open ai chatgpt” received a higher score than both “Astroworld Festival Travis Scott” and “Trump assassination attempt” despite the latter representing more clearly traumatic incidents. This limitation necessitated the manual identification of events which were negative to ensure appropriate dataset construction.

The new method defined negative global events as significant occurrences that cause harm, suffering, or disruption at regional or global scales, typically involving death, injury, displacement, environmental damage, economic loss, or infrastructure destruction. After manual selection, the dataset includes 233 negative events (e.g., 70 conflict-related events, 52 natural disaster-related events, and 35 health crisis etc.).

Time-window determination The analysis is centered on each event’s peak search week, extending the window two days before and after as the exact peak day cannot be extracted within 5-years search range. Due to the overlap between events, only the portion of non-overlapping (pre-event or post-event) data was retained, resulting in 36 pre-event periods (5.5%) and 38 post-event periods (6.3%).

Data sampling To mitigate potential biases from overrepresented time periods and to prevent individual users from dominating the dataset, a balanced sampling strategy was implemented. Posts were filtered by subreddit (mental health and personal experience-related subreddits, such as r/depression, r/AskReddit) and content length (minimum 200 characters), with labels assigned based on timestamps relative to event dates. The procedure then samples evenly across pre-event and post-event windows to ensure balanced representation, using the minimum available post count as the target

Event	vs. black friday	vs. weather	vs. covid-19	avg score
Astroworld Festival Travis Scott	<1:22 ⁷	<1:56	<1:14	0.0224
Trump assassination attempt	<1:2	1:57	<1:1	0.2529
open ai chatgpt	<1:<1	<1:72	<1:2	0.4213

Table 1: Example of the comparison between different events

threshold.

In the final stages, the process applies user-level caps to prevent individual users from dominating the dataset, calculating the median posts per user and randomly sampling down to this cap. The final balancing step equalizes pre-event and post-event post counts by randomly sampling down to the smaller group size, ultimately creating a balanced dataset that controls for event representation, user influence, and temporal period distribution.

3.3 Covid-19 Lockdown Data

Unlike other events in this study which may have more subtle or localized effects, Covid-19 lockdowns demonstrably impacted mental health worldwide. Since the impact of the events identified in the dataset in Section 3.2 is unclear, the methodology of using PLMs to detect the occurrence of a negative event is verified on a second dataset, known to have invoked mental health responses (Rani et al., 2024).

Baseline period January 10-18, 2020 was identified as the pre-lockdown window. This window was ideal as it: (1) preceded widespread international Covid-19 awareness, (2) did not overlap with other global negative events.

Initial Covid-19 lockdown period March 22-28, 2020, which coincided with implementation of widespread lockdown measures across multiple countries (Calfas et al., 2020; News, 2020) was selected and it also represented the peak of global search for “covid-19 lockdown” and “covid-19” according to Google Trends.

Data sampling Same pre- and post-event filtering strategy as in the global event data was used. However, for the Covid-19 lockdown data, the additional balancing steps for user representation is omitted to ensure there are sufficient data for analysis.

3.4 Final dataset

Posts were labeled based on their temporal relationship to global negative events and Covid-19

lockdown onset, with pre-event and pre-lockdown posts assigned as baseline mental health status (Label 0) and posts after-event and after-lockdown period labeled as worse mental health status (Label 1). This temporal labeling approach assumes that the global negative events and Covid-19 lockdown measures represent a population-level stressor that would manifest in changed mental health discourse patterns, providing a computational proxy for detecting event-driven mental health impacts at scale, rather than individual clinical diagnoses.

Global negative events data This sampled global negative events data comprised 25,185 users and 53,644 posts, with 26,822 (50%) posts labeled as “baseline mental health status” and 26,822 (50%) posts labeled as “worse mental health status”. The mean post content length is 468.78 characters.

Covid-19 lockdown data The Covid-19 lockdown data comprised 2,204 users and 6,345 submissions, with 3,418 posts (53.87%) labeled as “baseline mental health status” and 2,927 posts (46.13%) labeled as “worse mental health status”. The mean post content length is 680.36 characters.

4 Methods

A five-fold cross-validation with user-level splitting to prevent data leakage (Tsakalidis et al., 2022) was employed, and each training set was balanced by downsampling the majority class to a 50/50 distribution. Performance metrics (accuracy, precision, recall, F1) were averaged across all folds. Evaluation compared three PLMs: RoBERTa (Liu et al., 2019) and its mental health tuned variant Mental-RoBERTa (Ji et al., 2022), as well as the smaller model DistilBERT (Sanh et al., 2019).

Each model was fine-tuned with a learning rate of 1e-5, batch size of 8, maximum sequence length of 512 tokens, and early stopping to prevent overfitting. The training process utilized a cross-entropy loss function.

Model	avg_acc (%)	avg_prec	avg_rec	avg_f1
RoBERTa	0.5118	0.5014	0.7915	0.5418
MentalRoBERTa	0.5083	0.5096	0.9598	0.6627
DistilBERT	0.5157	0.5258	0.8668	0.6380

Table 2: Comparison of model performance on Covid-19 lockdown data

5 Results and Discussion

Performance on global negative event data DistilBERT attained numerically higher averaged F1 scores (0.6553) than RoBERTa (0.5795) and MentalRoBERTa (0.5558) in 5-fold cross-validation.

DistilBERT’s distilled knowledge from BERT has been stated to be effective in identifying linguistic markers of psychological distress across diverse mental health conditions, aligning with its efficacy in mental health conditions classification (Oh et al., 2023). While MentalRoBERTa and RoBERTa usually outperform DistilBERT on mental health tasks, the distilled model is more likely to tolerate noise in the data (Sanh et al., 2019).

The lower performance of the models overall may be due to the impact of global events across different communities: global events may affect users differently based on their backgrounds. For instance, “Aleppo Russian airstrikes Syrian” would likely have significantly greater impact on Syrian users or those with connections to the region than on the broader Reddit community.

Performance on Covid-19 lockdown data In the Covid-19 lockdown data featuring sharp temporal boundaries, MentalRoBERTa attained the highest averaged F1 score (0.6627), closely followed by DistilBERT (0.6380) (Table 2). This supports our earlier conclusion about model performance patterns.

The Covid-19 lockdown data’s clear pre/during-lockdown demarcation demonstrates the validity of the proposed method. The higher performance of MentalRoBERTa over DistilBERT on this, cleaner data also supports the conclusion regarding DistilBERT’s utility for the global task due to its noise tolerance.

SHAP analysis To validate the assumption of the event of Covid-19 lockdown producing detectable changes in the language use of people with mental health conditions, we examined the fine-tuned MentalRoBERTa model with SHapley Additive exPlanations (SHAP) explainability analysis (Figure 2). The model’s ability to correctly identify posts indi-

cating deteriorated mental health status (LABEL_1) through the use of linguistic markers such as “hopeless”, “damned” and “feel greedy wrong” demonstrates that language patterns do indeed shift in measurable ways following global crises. These results support our core hypothesis that mental health discourse exhibits detectable linguistic changes in response to external events.

To provide model interpretability insights, we applied SHAP analysis to randomly sampled posts (N=500) from Covid-19 lockdown data, extracting 72 unique high-importance linguistic words (threshold = 0.002). This analysis reveals specific words and phrases that most strongly influence the model’s mental health classification decisions, providing transparency into the model’s decision-making process.

SHAP analysis suggested the model learned to identify key linguistic markers associated with mental health distress (“anxiety”, “depression”, “worthless”), Covid-19 related stressors (“virus”, “quarantine”) and relationship concerns (“rejected”, “relationship”). The discovery of these patterns confirmed our hypothesis, that global negative events affect mental health, since the models appear to be based on depression markers.

6 Conclusion and Future Work

When exploring the effect of global negative events on population level mental health, we show the utility of DistilBERT in distinguishing posts pre- and post- such events. The data is shown to be too noisy for bigger, less noise tolerant, models, such as MentalRoBERTa. The approach is validated, and the conclusion supported, by a separate evaluation on a much cleaner dataset surrounding Covid-19 lockdown. The noise found in the global event dataset could be addressed in future work by combining the global approach with an individual-based approach. Future work will also incorporate manual content analysis to provide deeper qualitative insights into the mental health language use patterns identified by our approach.

Furthermore, our SHAP interpretability analysis

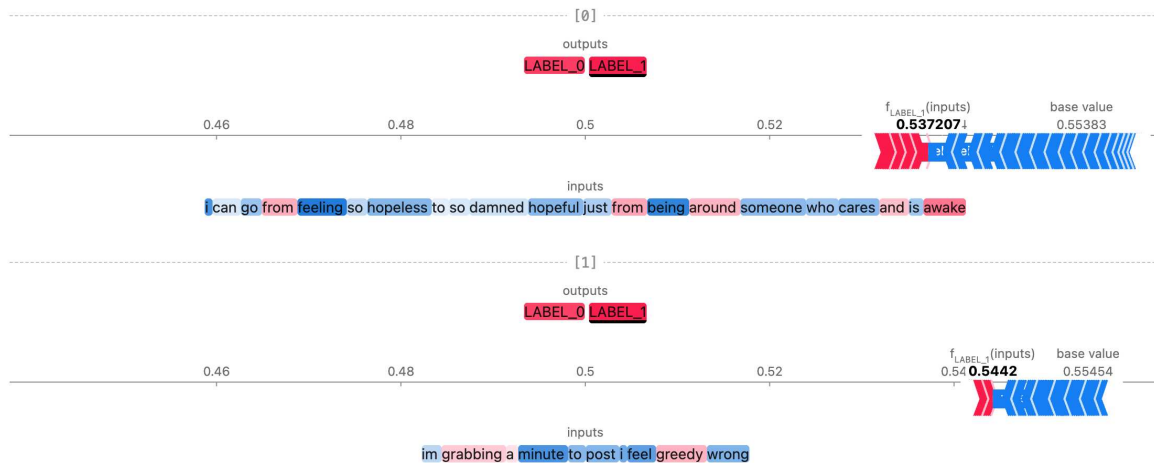


Figure 2: SHAP analysis of fine-tuned MentalRoBERTa on sample texts from SHAP

revealed that mental health-related vocabulary (anxiety, depression, worthlessness) emerged as the most influential linguistic markers for mental health classification. This finding validates our temporal labeling approach and demonstrates the model’s ability to identify mental health-related linguistic markers in post initial lockdown discourse.

7 Limitations

Our study faces some limitations. First, the psychological consequences of global negative events vary significantly across individuals based on geographic proximity, cultural context, and personal resilience levels. Our event selection methodology prioritized clean, non-overlapping time periods, which may have resulted in excluding some significant negative events that occurred during overlapping timeframes. This approach, while necessary for methodological clarity, potentially omits events with substantial psychological impact.

Second, labeling mental state changes based strictly on pre-/post-event time windows introduces confounding variables and temporal resolution mismatch. Personal life events (e.g., divorce, job loss) coinciding with global negative events may distort the perceived event-mental health linkage. Psychological responses to events may unfold over weeks (e.g., grief processing) or manifest abruptly (e.g., panic attacks), yet our fixed labeling window fails to accommodate such dynamics. Despite this limitation, the SHAP interpretability analysis provides evidence that this simplified labeling captured meaningful signal, as the model learned to identify clinically relevant vocabulary (anxiety,

depression, worthlessness) and pandemic-related stressors (virus, quarantine).

Third, users with pre-existing conditions (e.g., chronic depression) might exhibit stable negative language patterns, diluting event-specific linguistic shifts.

8 Ethics Statement

Ethical approval for the study was obtained from the University of Sheffield Information School Ethics Committee (ethical application reference 064309). The approved ethical application specifically covered data collection and model fine-tuning, but excluded manual analysis or in-depth qualitative analysis to protect both researcher well-being and Reddit user privacy. The automated approach to data collection and model fine-tuning mitigates potential ethical concerns associated with manual examination of sensitive mental health content.

Acknowledgments

This research was supported by funding from the Economic and Social Research Council (ESRC). We also thank Nicholas Musembi for his technical assistance in resolving code implementation issues.

References

- Roy F Baumeister, Ellen Bratslavsky, Catrin Finkenauer, and Kathleen D Vohs. 2001. *Bad Is Stronger than Good*. *Review of General Psychology*, 5(4):327–370.
- Jennifer Calfas, Margherita Stancati, and Chuin-Wei Yap. 2020. *California Orders Lockdown for State’s 40 Million Residents*.

- Suchandra Chakraborty, Sudeshna Jana, Manjira Sinha, and Tirthankar Dasgupta. 2025. [Self-State Evidence Extraction and Well-Being Prediction from Social Media Timelines](#). In *Proceedings of the Tenth Workshop on Computational Linguistics and Clinical Psychology*. Association for Computational Linguistics.
- Pirina Inna and Çöltekin Çağrı. 2018. [Identifying Depression on Reddit: The Effect of Training Data](#). In *Proceedings of the 2018 EMNLP Workshop SMM4H: The 3rd Social Media Mining for Health Applications Workshop and Shared Task*, pages 9–12.
- Shaoxiong Ji, Xue Li, Zi Huang, and Erik Cambria. 2021. [Suicidal Ideation and Mental Disorder Detection with Attentive Relation Networks](#). *Neural Computing and Applications*.
- Shaoxiong Ji, Tianlin Zhang, Luna Ansari, Jie Fu, Prayag Tiwari, and Erik Cambria. 2022. [MentalBERT: Publicly Available Pretrained Language Models for Mental Healthcare](#). In *Proceedings of the Language Resources and Evaluation Conference*.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. [RoBERTa: A Robustly Optimized BERT Pretraining Approach](#). In *ICLR 2020 Conference*.
- Yixin Liu, Alexander R. Fabbri, Jiawen Chen, Yilun Zhao, Simeng Han, Shafiq Joty, Pengfei Liu, Dragomir Radev, Chien-Sheng Wu, and Arman Cohen. 2024. [Benchmarking Generation and Evaluation Capabilities of Large Language Models for Instruction Controllable Summarization](#). In *NAACL 2024 Findings*.
- May Oo Lwin, Jiahui Lu, Anita Sheldenkar, Peter Johannes Schulz, Wonsun Shin, Raj Gupta, and Yinpeng Yang. 2020. [Global Sentiments Surrounding the COVID-19 Pandemic on Twitter: Analysis of Twitter Trends](#). *JMIR Public Health Surveill*, 6(22).
- Matthew Louis Mauriello, Thierry Lincoln, Grace Hon, Dorien Simon, Dan Jurafsky, and Pablo Paredes. 2021. [SAD: A Stress Annotated Dataset for Recognizing Everyday Stressors in SMS-like Conversational Systems](#). In *CHI EA '21: Extended Abstracts of the 2021 CHI Conference on Human Factors in Computing Systems*.
- Modhurima Moitra, Shanise Owens, Maji Hailemariam, Katherine S. Wilson, Augustina Mensa-Kwao, Gloria Gonese, Christine K. Kamamia, Belinda White, Dorraine M. Young, and Pamela Y. Collins. 2023. [Global Mental Health: Where We Are and Where We Are Going](#). *Current Psychiatry Reports*, 25(7):301–311.
- BBC News. 2020. [Covid-19: Milestones of the global pandemic](#).
- Jaedong Oh, Mirae Kim, Hyejin Park, and Hayoung Oh. 2023. [Are You Depressed? Analyze User Utterances to Detect Depressive Emotions Using DistilBERT](#). *Applied Science*, 13(10):6223.
- Saima Rani, Khandakar Ahmed, and Sudha Subramani. 2024. [From posts to knowledge: Annotating a pandemic-era reddit dataset to navigate mental health narratives](#). *Applied Science*.
- María P. Raveau, Julián I. Goñi, José F. Rodríguez, Isidora Paiva-Mack, Fernanda Barriga, María P. Hermosilla, Claudio Fuentes-Bravo, and Susana Eyheramendy. 2023. [Natural language processing analysis of the psychosocial stressors of mental health disorders during the pandemic](#). *npj Mental Health Res*, 2(17).
- José Frederico Rodrigues, Henrique Lopes Cardoso, and Carla Teixeira Lopes. 2025. [Evaluating Llama 3 for Text Simplification: A Study on Wikipedia Lead Sections](#). In *WWW '25: Companion Proceedings of the ACM on Web Conference 2025*.
- Victor Sanh, Lysandre Debut, Julien Chaumond, and Thomas Wolf. 2019. [Distilbert, a distilled version of bert: smaller, faster, cheaper and lighter](#). In *5th Workshop on Energy Efficient Machine Learning and Cognitive Computing - NeurIPS 2019*.
- Jonathan G. Shalom and Idan M. Aderka. 2020. [A meta-analysis of sudden gains in psychotherapy: Outcome and moderators](#). *Clinical Psychology Review*, 76.
- Han-Chin Shing, Suraj Nair, Ayah Zirikly, Meir Friedenberg, Hal Daumé III, and Philip Resnik. 2018. [Expert, Crowdsourced, and Machine Assessment of Suicide Risk via Online Postings](#). In *Proceedings of the Fifth Workshop on CLPsych*, pages 25–36.
- stuck_in_the_matrix, Watchful1, and RaiderBDev. 2025. [Reddit comments/submissions 2005-06 to 2024-12](#).
- Rebecca R Thompson, Nickolas M Jones, E Alison Holman, and Roxane Cohen Silver. 2019. [Media exposure to mass violence events can fuel a cycle of distress](#). *ScienceAdvances*, 5.
- Adam Tsakalidis, Jenny Chim, Iman Munire Bilal, Ayah Zirikly, Dana Atzil-Slonim, Federico Nanni, Philip Resnik, Manas Gaur, Kaushik Roy, Becky Inkster, Jeff Leintz, and Maria Liakata. 2022. [Overview of the CLPsych 2022 Shared Task: Capturing Moments of Change in Longitudinal User Posts](#). In *Proceedings of the Eighth Workshop on Computational Linguistics and Clinical Psychology*, pages 184–198. Association for Computational Linguistics.
- Talia Tseriotou, Jenny Chim, Ayal Klein, Aya Shamir, Guy Dvir, Iqra Ali, Cian Kennedy, Guneet Singh Kohli1, Anthony Hills, Ayah Zirikly, Dana Atzil-Slonim, and Maria Liakata. 2025. [Overview of the CLPsych 2025 Shared Task: Capturing Mental Health Dynamics from Social Media Timelines](#). In *Proceedings of the 10th Workshop on Computational Linguistics and Clinical Psychology (CLPsych 2025)*, pages 193–217. Association for Computational Linguistics.