

MindSense AI: An Intelligent Mental Wellness Companion Using Data Science, Machine Learning, NLP, and Generative AI

Ms. Prachi Surve¹, Shalaka Wagh², Vaishnavi Bhakad³, Shravani Patil⁴

¹Assistant Professor, D. Y. Patil Deemed to be University, School of Humanities & Sciences, Belapur, Maharashtra, India.

²Student, Department of MCA, D. Y. Patil Deemed to be University, School of Humanities & Sciences, Belapur, Maharashtra, India.

³Student, Department of MCA, D. Y. Patil Deemed to be University, School of Humanities & Sciences, Belapur, Maharashtra, India.

⁴Student, Department of MCA, D. Y. Patil Deemed to be University, School of Humanities & Sciences, Belapur, Maharashtra, India.

Abstract—Mental health difficulties such as stress, loneliness, and academic pressure are common, yet many people hesitate to discuss these feelings openly with those around them. MindSense AI addresses this gap as an AI-driven mental wellness companion that draws on natural language processing, machine learning, and generative AI to interpret the emotional tone and context of a conversation and to respond in a supportive, human-like manner. The proposed system brings together several components – text processing, emotion and sentiment detection, intent recognition, generative response formulation, and general well-being recommendations – into a single conversational pipeline. MindSense AI is conceived purely as a non-clinical support tool; it is not designed to diagnose mental health conditions or to substitute for licensed professionals. Data privacy, operational safety, and ethical AI practice are treated as core design requirements rather than afterthoughts, and this paper outlines how these principles are woven into each stage of the architecture, from data handling to response generation.

Key Words: Mental Wellness, MindSense AI, Natural Language Processing, Machine Learning, Generative AI, Sentiment Analysis, Emotion Detection, Conversational AI, Responsible AI

1. INTRODUCTION

Mental wellness shapes nearly every dimension of daily life – how people study, work, maintain relationships, and experience their surroundings. Yet conversations about stress, anxiety, or low mood are still often avoided, whether from stigma, lack of time, or simply not knowing whom to approach. Advances in natural language processing (NLP), machine learning, and artificial intelligence now make it possible for software to interpret human language with a level of nuance that was out of reach only a decade ago. The introduction of the transformer architecture reshaped how models capture relationships within text [1], and BERT later demonstrated

that bidirectional context could substantially improve performance across a wide range of language-understanding tasks [2].

MindSense AI builds on these developments to offer a conversational companion organised around a simple five-stage idea – Listen, Understand, Analyse, Respond, and Guide. In practice, this means the system reads what a user types, works out the sentiment and likely emotional state behind it, considers the intent and context of the conversation, and then produces a supportive, non-judgemental reply along with general well-being suggestions where appropriate. The system is explicitly framed as a non-clinical companion: it does not attempt to diagnose or treat any condition, and it is built with privacy, safety, transparency, and human oversight as guiding design principles rather than optional add-ons.

2. LITERATURE REVIEW

Several strands of prior work inform the design of MindSense AI. Vaswani et al. [1] introduced the transformer architecture, which relies on self-attention rather than recurrence to model relationships between words and has since become the backbone of most modern language models. Building on this direction, Devlin et al. [2] proposed BERT, a bidirectional model that reads context from both sides of a word simultaneously and achieved strong results across a range of language-understanding benchmarks.

Beyond core language modelling, researchers have also examined how conversational systems perform in mental-health settings. Fitzpatrick et al. [3] evaluated Woebot, a text-based conversational agent, and found that such agents can offer accessible, low-barrier support for mental wellness. To support fine-grained emotion recognition, Demszky et al. [4] released the GoEmotions dataset, comprising roughly 58,000 Reddit comments labelled across 27 distinct emotion categories plus a neutral class –

a resource that remains one of the more detailed emotion-labelled corpora available for this kind of work.

At a broader level, Balan et al. [5] reviewed conversational agents aimed at young people's mental health and reported both encouraging benefits and notable limitations, including concerns about engagement over time and the depth of support such tools can realistically provide. Tabassi and NIST [6] published an AI Risk Management Framework that foregrounds privacy, fairness, transparency, safety, and reliability as core criteria for responsible AI systems – principles this paper adopts directly. Finally, Li et al. [7] conducted a systematic review of AI-based conversational agents for mental health and well-being, concluding that while the evidence base shows promise, further long-term research and rigorous evaluation are still needed before such tools can be deployed with confidence.

3. RESEARCH GAP

Much of the existing literature treats emotion detection, chatbot design, and language modelling as separate research problems, each studied largely in isolation. Studies that focus on emotion classification rarely address how a system should respond once an emotion is detected, and conversational-agent research does not always engage deeply with privacy or safety considerations. MindSense AI is positioned to close this gap by bringing these strands together into a single, non-clinical framework that explicitly treats privacy, safety, and responsible AI as first-class design requirements rather than an afterthought bolted onto a working prototype.

4. PROBLEM STATEMENT

Everyday conversational text carries a great deal of emotional information, but extracting that information reliably is far from straightforward. Informal writing styles, slang, emojis, abbreviations, short or fragmented messages, sarcasm, spelling variation, and limited surrounding context all make automated emotion and intent recognition difficult. Rule-based chatbots, meanwhile, tend to fall back on repetitive, context-independent replies that can feel impersonal or even dismissive in a sensitive conversation. What is needed, therefore, is an integrated NLP and AI framework capable of analysing conversational text and generating supportive, context-aware responses, while consistently respecting the privacy, safety, and ethical boundaries appropriate to a mental-wellness application.

5. OBJECTIVES

- Develop an AI-based mental wellness companion for everyday conversational support.

- Preprocess and analyse conversational text for downstream NLP tasks.
- Detect sentiment and broad emotional categories expressed by the user.
- Interpret conversational context and user intent.
- Compare classical machine learning models against transformer-based approaches.
- Generate supportive, context-aware responses using generative AI.
- Offer general, non-clinical well-being suggestions.
- Visualise emotional and sentiment trends over time.
- Uphold privacy, safety, transparency, and responsible AI practice throughout.

6. PROPOSED METHODOLOGY

The system follows a nine-stage pipeline:

Data Collection → Preprocessing → Feature Extraction → Emotion & Sentiment Analysis → Context & Intent Detection → AI Response Generation → Recommendations → Visualisation → Evaluation

- Data Collection: labelled emotion datasets serve as the foundation for training and evaluation.
- Preprocessing: raw text is cleaned while deliberately retaining emojis and hashtags that carry emotional signal.
- Feature Extraction: TF-IDF vectors and transformer-based embeddings represent text numerically.
- Analysis: sentiment polarity and broad emotion categories are identified from the extracted features.
- Context & Intent: the current message is interpreted alongside recent conversation history.
- Response Generation: generative AI produces a safe, supportive reply grounded in the detected emotion and intent.
- Recommendation: general well-being suggestions are offered where appropriate.
- Visualisation: non-clinical emotional trends are presented back to the user.
- Evaluation: accuracy, response quality, safety, and usability are all assessed.

7. DATASET DESCRIPTION

GoEmotions offers a strong reference point for emotion-classification work of this kind. Demszky et al. [4] report approximately 58,000 English-language Reddit comments annotated across 27 emotion categories plus a neutral label (Fig. 1). For a student-level prototype such as MindSense AI, it makes sense to collapse these 27 fine-grained labels into a smaller set of broader categories – for instance, joy, sadness, anger, fear, and neutral – to keep

classification tractable. That said, any such mapping should be documented explicitly and evaluated on its own merits, since grouping emotions inevitably changes the nature, and difficulty, of the underlying classification task.

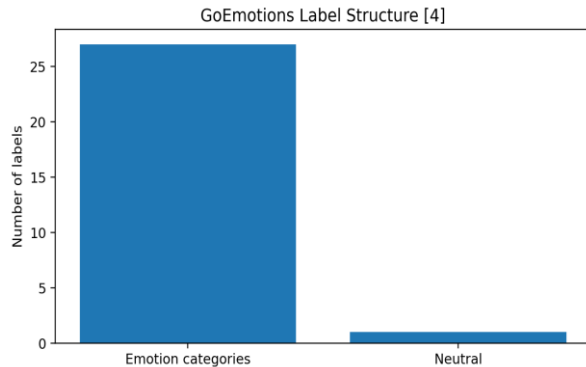


Fig -1: Label structure of the GoEmotions dataset reported by Demszky et al. [4]

8. DATA PREPROCESSING

- Normalise text into a consistent format (lowercasing, standard encoding).
- Strip out URLs and other non-informative noise.
- Tokenise the cleaned text.
- Remove stop words that add little semantic value.
- Apply stemming or lemmatisation to reduce words to a common root form.
- Handle emojis, hashtags, slang, and abbreviations carefully, since these often carry emotional weight rather than noise.
- Discard duplicate or empty entries.
- Preserve conversational context across turns wherever it is needed for downstream analysis.

9. FEATURE EXTRACTION AND NLP PROCESSING

- TF-IDF: a simple, interpretable baseline representation of text.
- Unigrams and bigrams: capture individual words and short two-word phrases respectively.
- BERT / Transformers: capture contextual meaning that bag-of-words methods miss [2].
- Sentiment Analysis: classifies text as positive, negative, or neutral.
- Emotion Analysis: identifies broader emotional patterns beyond simple polarity.
- Context and Intent Detection: helps the system understand not just what was said, but why.

10. MACHINE LEARNING AND GENERATIVE AI MODELS

Table 1 summarises the classification and generation techniques considered for MindSense AI, ranging from fast

classical baselines to transformer-based and generative models.

Table -1: Models used for classification and response generation

Model Technique /	Purpose	Role
Naive Bayes	Text classification	Fast baseline
Logistic Regression	Text classification	Interpretable baseline
SVM	Text classification	Strong TF-IDF-based model
Random Forest	Classification	Comparison model
BERT / Transformers	Context understanding	Advanced model
Generative AI	Response generation	Supportive conversation

11. SYSTEM ARCHITECTURE

User Input → Preprocessing → NLP Feature Extraction → Emotion & Sentiment Analysis → Context & Intent Detection → Safety Guidelines → Generative AI → Supportive Response → Recommendations → Visualisation Each stage builds on the output of the one before it. A dedicated safety-guidelines checkpoint sits just ahead of the generative AI component, screening candidate context before a response is generated, so that only appropriate, supportive replies reach the user.

12. RECOMMENDATION MODULE

The recommendation module offers general, non-clinical suggestions based on the broad emotional patterns detected during a conversation. Depending on context, this might include taking a short break, maintaining a regular daily routine, engaging in an enjoyable activity, reaching out to someone the user trusts, or, where appropriate, considering professional support. These suggestions are always framed as general guidance rather than medical advice, diagnosis, or treatment, and the system is careful never to present itself as a substitute for a qualified professional.

13. DATA VISUALIZATION MODULE

Visualisation helps users notice non-clinical patterns in their own conversations over time – which emotions come up most often, how sentiment shifts from day to day, and so on. Tools such as Matplotlib and Plotly are well suited to generating these views. Every visualisation is clearly

labelled as being intended for awareness and self-reflection, not as a clinical assessment of any kind.

14. EVALUATION FRAMEWORK

No experimental accuracy figures are reported in this paper, since the study describes a proposed system rather than a completed implementation with test results. Once built, the system should be evaluated on a held-out test set using the following standard metrics:

Accuracy = Correct predictions/Total predictions ... (1)

Precision = TP / (TP + FP) ... (2)

Recall = TP / (TP + FN) ... (3)

F1-score = $2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall})$... (4)

A confusion matrix should also be produced to give a class-by-class breakdown of correct and incorrect predictions. Beyond these quantitative measures, response quality should be reviewed by human raters against a clear rubric covering relevance, clarity, supportive tone, factual appropriateness, safety, and consistency with the system's non-clinical purpose.

15. EVIDENCE FROM RELATED RESEARCH

To situate MindSense AI within the wider evidence base, it is worth revisiting the numbers reported by Li et al. [7] in their systematic review. Their search identified 7,834 candidate records, of which 533 proceeded to full-text review; from these, 35 studies met the eligibility criteria, and 15 of those were randomized controlled trials (Fig. 2 and Fig. 3). These figures come directly from the published review and describe the state of the field generally – they are not results produced by the MindSense AI prototype itself.

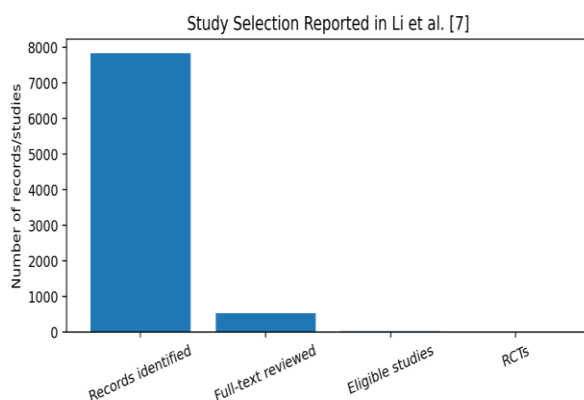


Fig -2: Study-selection counts reported by Li et al. [7]

Composition of 35 Eligible Studies in Li et al. [7]

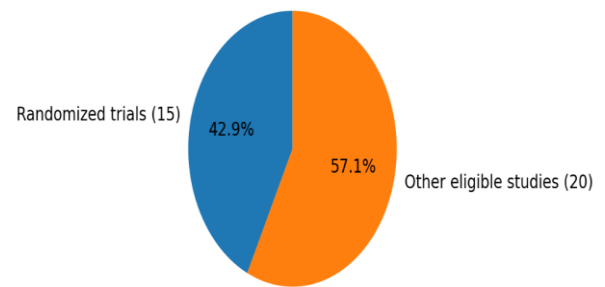


Fig -3: Of the 35 eligible studies in Li et al. [7], 15 were randomized trials and the remaining 20 were not

16. EXPECTED RESULTS AND DISCUSSION

Once implemented, MindSense AI is expected to function as a modular system: classical machine learning models such as Naive Bayes, logistic regression, and SVM can serve as fast, interpretable baselines, while transformer-based models are expected to improve contextual understanding at the cost of additional computational overhead. Evaluation will combine standard classification metrics (accuracy, precision, recall, F1-score) with qualitative assessment of response quality, safety, privacy handling, and overall usability. Actual results – quantitative and qualitative – will be reported once the system has been implemented and tested against real conversational data.

17. PRIVACY AND SECURITY

MindSense AI is designed to collect only the information strictly necessary for it to function, avoiding unnecessary personal details wherever possible. Conversational data is protected through secure storage, access controls, and encrypted communication channels. The system is also designed to explain clearly how user data is processed and stored, and to minimise the retention of sensitive information beyond what is actually required.

18. RESPONSIBLE AI, SAFETY, AND ETHICS

As a supportive, non-clinical system, MindSense AI is not designed to diagnose conditions or to replace professional mental health care. Its design follows established responsible-AI guidelines covering privacy, fairness, transparency, reliability, and accountability [6]. The team has also considered several practical challenges that any deployment must account for – ambiguous or sarcastic language, cultural variation in how emotion is expressed, dataset bias, model errors, and the risk of inappropriate or unsafe responses – and treats each of these as an active design concern rather than a hypothetical edge case.

19. FUTURE SCOPE

Future work could extend MindSense AI in several directions: adopting more advanced transformer architectures and large language models, adding multilingual support (for example, English-Hindi and English-Marathi conversations), improving context and intent detection, incorporating sarcasm detection, exploring voice and multimodal input, and building in explainable-AI features and privacy-preserving personalisation. Stronger safety mechanisms and longer-term user studies would also be valuable before any real-world deployment is considered.

20. CONCLUSION

MindSense AI brings together NLP, machine learning, data science, and generative AI into a single, non-clinical mental wellness companion capable of analysing conversational text, recognising broad emotions and sentiment, understanding context, and responding in a supportive way. The framework also incorporates visualisation, privacy, safety, and responsible-AI considerations as core parts of the design rather than optional extras. Its main contribution lies in this integrated, non-clinical approach; future research will focus on improving accuracy, extending multilingual support, and strengthening explainability, privacy, and safety. Ultimately, the system is intended to support users alongside, not in place of, qualified mental health professionals.

ACKNOWLEDGEMENT

The authors would like to thank their mentors and institution for their guidance and support throughout this work.

REFERENCES

- [1] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin, "Attention Is All You Need," *Advances in Neural Information Processing Systems (NeurIPS)*, 2017.
- [2] J. Devlin, M. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," *Proc. NAACL-HLT*, 2019.
- [3] K. K. Fitzpatrick, A. Darcy, and M. Vierhile, "Delivering Cognitive Behavior Therapy to Young Adults With Symptoms of Depression and Anxiety Using a Fully Automated Conversational Agent (Woebot): A Randomized Controlled Trial," *JMIR Mental Health*, 2017.
- [4] D. Demszky, D. Movshovitz-Attias, J. Ko, A. Cowen, G. Nemade, and S. Ravi, "GoEmotions: A Dataset of Fine-Grained Emotions," *Proc. ACL*, 2020.
- [5] S. Balan, A. Rathod, et al., "Automated Conversational Agents for Mental Health: A Systematic Review," *npj Digital Medicine*, 2024.
- [6] E. Tabassi, "Artificial Intelligence Risk Management Framework (AI RMF 1.0)," National Institute of Standards and Technology (NIST), 2023.
- [7] H. Li, R. Zhang, et al., "Artificial Intelligence-Based Conversational Agents for Mental Health and Well-Being: A Systematic Review," *npj Digital Medicine*, 2023.