

Fine-Tuned DistilBERT and Interactive Dashboard for Early Detection of Suicidal Ideation in Social Media Posts

MSc Research Project
AI for business

Diego Marquez Salazar
Student ID: x23410914

School of Computing
National College of Ireland

Supervisor: Rejwanul Haque

National College of Ireland
MSc Project Submission Sheet
School of Computing



Student Name: Diego Marquez Salazar
Student ID: x23410914
Programme: MSc in AI for Business **Year:** 2025
Module: Practicum 2
Supervisor: Rejwanul Haque
Submission Due Date: Monday 15th September 2025
Project Title: Fine-Tuned DistilBERT and Interactive Dashboard for Early Detection of Suicidal risk in Social Media Posts
Word Count: 3415 **Page Count** 14

I hereby certify that the information contained in this (my submission) is information pertaining to research I conducted for this project. All information other than my own contribution will be fully referenced and listed in the relevant bibliography section at the rear of the project.

ALL internet material must be referenced in the bibliography section. Students are required to use the Referencing Standard specified in the report template. To use other author's written or electronic work is illegal (plagiarism) and may result in disciplinary action.

Signature:

A blue ink handwritten signature, appearing to be "Diego Marquez Salazar", written over a light blue horizontal line.

Date: Monday 15th September 2025

PLEASE READ THE FOLLOWING INSTRUCTIONS AND CHECKLIST

Attach a completed copy of this sheet to each project (including multiple copies)	☐
Attach a Moodle submission receipt of the online project submission, to each project (including multiple copies).	☐
You must ensure that you retain a HARD COPY of the project, both for your own reference and in case a project is lost or mislaid. It is not sufficient to keep a copy on computer.	☐

Assignments that are submitted to the Programme Coordinator Office must be placed into the assignment box located outside the office.

Office Use Only	
Signature:	
Date:	
Penalty Applied (if applicable):	

Fine-Tuned DistilBERT and Interactive Dashboard for Early Detection of Suicidal Ideation in Social Media Posts

Diego Marquez Salazar
x23410914

Abstract

The early detection of suicidal ideation in social media content can save lives by enabling timely interventions. In this work, we present the development of a fine-tuned DistilBERT model on a pre-processed dataset to classify texts as indicative of suicide risk or not. The model was integrated into an interactive dashboard built with Streamlit, which allows both manual input of individual texts and bulk analysis through CSV file uploads. The tool displays classification results along with estimated probabilities and provides additional visualizations, such as word frequency distributions, to help uncover patterns in the data. Experimental evaluation using 5-fold cross-validation demonstrated outstanding performance, achieving accuracy above 93% and ROC-AUC scores exceeding 0.90. This solution offers a practical and socially impactful contribution, aimed at mental health professionals, researchers, and organizations seeking to monitor and prevent suicide related risks within online communities.

Keywords: Suicide prevention, DistilBERT, NLP, sentiment analysis, Streamlit, dashboard, mental health.

1 Introduction

Suicide remains one of the leading causes of death worldwide, with depression being a major underlying factor. The ability to detect early signs of suicidal ideation can play a crucial role in prevention, allowing timely interventions that may save lives. The motivation for pursuing this topic stems from personal and family experiences with depression and suicide attempts, which have given me a deep awareness of the importance of early detection and effective support.

From an academic perspective, Natural Language Processing (NLP) and transformer-based models, such as BERT, have shown strong performance in text classification tasks, including sentiment analysis and mental health detection. I chose this approach because it is both scientifically proven and highly adaptable, with potential for practical deployment in real-world contexts.

This project aims to develop and evaluate a fine-tuned transformer model for detecting suicidal ideation in social media texts, complemented by an interactive Streamlit dashboard. The

dashboard enables organizations such as prevention centres, private healthcare providers, and schools to upload datasets, analyse trends, visualize key metrics, and identify high-risk content. While developed as an academic proof of concept, the system is designed with the flexibility to be adapted for different datasets, making it scalable and applicable to future real-world scenarios.

The contribution of this work lies in demonstrating the integration of advanced NLP techniques, predictive modelling, and user-friendly data visualization tools into a single, functional prototype. By bridging research and usability, it opens opportunities for both further academic exploration and potential operational deployment.

This report is structured as follows: Section 2 reviews related work in suicide ideation detection and NLP-based mental health research. Section 3 outlines the research methodology, including dataset preparation, model training, and evaluation. Section 4 presents the design specifications of the dashboard. Section 5 details the implementation process. Section 6 evaluates the model’s performance and dashboard functionality. Finally, Section 7 concludes the report and discusses future research directions.

2 Related Work

The intersection of artificial intelligence, mental health, and social media analytics has emerged as a critical research area, driven by the growing availability of user data and the urgent need for early detection of mental health risks. Over the last decade, researchers have increasingly leveraged Natural Language Processing (NLP) techniques to detect signs of depression, anxiety, and suicidal ideation in online communications, with social media serving as a valuable but challenging source of information (Coppersmith et al., 2015).

2.1 Early Lexicon-Based and Machine Learning Approaches

Initial studies on suicide risk detection primarily relied on lexicon-based methods, where pre-defined dictionaries of mental health related terms (e.g., LIWC, NRC Emotion Lexicon) were used to identify linguistic markers associated with distress. For instance, Stirman and Pennebaker (2001) demonstrated how frequency patterns of certain word categories, such as first-person pronouns and negative emotion words, correlated with depression levels. While these approaches offered interpretability and simplicity, they often suffered from limited adaptability to new contexts and failed to capture nuanced language use, such as sarcasm, figurative expressions, or context dependent meanings.

To address these limitations, researchers began using traditional supervised machine learning algorithms such as Naive Bayes, Support Vector Machines (SVMs), and logistic regression trained on handcrafted features (Resnik et al., 2015). These models incorporated n-grams, sentiment scores, and syntactic patterns, improving classification accuracy compared to purely lexicon-based approaches. However, these methods required significant manual feature

engineering and still struggled with context-awareness, especially in short and noisy social media texts.

2.2 Rise of Deep Learning Models

The advent of deep learning significantly changed the field, enabling models to learn graded representations of text directly from data. Convolutional Neural Networks (CNNs) and Recurrent Neural Networks (RNNs), particularly Long Short-Term Memory (LSTM) networks, were employed to capture sequential dependencies and semantic relationships in mental health language (Chancellor et al., 2019). These models improved the detection of subtle indicators of psychological suffering, but their sequential processing often limited scalability and context range, especially in text.

2.3 Transformer-based Architectures and State of the Art

A breakthrough came with the introduction of transformer-based architectures such as BERT (Bidirectional Encoder Representations from Transformers) by Devlin et al. (2019). BERT's bidirectional attention mechanism allowed it to capture context from both preceding and following words, addressing limitations of earlier models. Fine-tuning BERT for mental health detection tasks has consistently produced best results across multiple datasets.

For example, Ji et al. (2022) fine-tuned BERT on Reddit posts from suicide-related subforums and achieved high precision in distinguishing between high-risk and low-risk content. Similarly, Sawhney et al. (2021) applied transformer models to Twitter data, outperforming traditional classifiers in identifying suicidal ideation while demonstrating robustness against short and noisy inputs.

Other variations, such as DistilBERT, RoBERTa, and ALBERT, have been explored to balance accuracy with computational efficiency. DistilBERT offers a smaller and faster alternative to BERT while retaining much of its performance, making it suitable for real-time applications like chatbots and dashboards (Sanh et al., 2019).

While predictive performance has advanced substantially, there remains a gap in translating these models into practical, deployable tools accessible to non-technical stakeholders such as school counsellors, healthcare practitioners, and suicide prevention organizations. Most academic studies present evaluation metrics but stop short of delivering functional interfaces or visual analytics that can support decision making in real world settings.

2.4 Contribution of This Work

This project addresses that gap by combining a fine-tuned DistilBERT model with an interactive Streamlit dashboard capable of real-time text classification and dataset-wide analysis. The dashboard not only predicts suicide risk for individual texts but also provides visual insights into trends, such as the most frequent high-risk words and statistical

distributions of predictions. By integrating explainable AI components with an accessible web interface, this work bridges the divide between technical research and actionable tools for early intervention.

3 Research Methodology

This section describes the methodological framework adopted for the development of a machine learning system capable of detecting suicidal ideation in social media posts. The process involved selecting an appropriate research approach, acquiring and preprocessing the dataset, fine-tuning a transformer-based language model, and evaluating its performance with standard metrics. The methodology follows best practices for Natural Language Processing (NLP) in mental health research (Chancellor & De Choudhury, 2020).

The project adopts an applied research approach with an experimental design, aiming to fine-tune an existing pre-trained transformer model for a binary text classification task. This approach allows leveraging the strengths of transfer learning while tailoring the model to the specific domain of suicidal ideation detection (Devlin et al., 2019). The work serves as a proof of concept that could be deployed in academic, healthcare, or institutional contexts for early intervention and monitoring.

3.1 Data Collection and Dataset Description

The dataset used was the Suicidal Ideation Detection Reddit Dataset, containing user-generated posts from Reddit annotated as either suicidal ideation or non-suicidal. The original dataset included raw text with minimal preprocessing and two primary columns: the post content and the associated binary label. This corpus was selected due to its relevance, accessibility, and prior usage in related research (Shing et al., 2018).

3.2 Data Preprocessing

Before model training, the dataset underwent preprocessing steps to ensure data quality and model readiness. Texts were cleaned by removing special characters, excessive whitespace, and URLs while retaining meaningful punctuation that could influence sentiment or tone. Missing values were removed to avoid noisy inputs. The *clean_Text* column was standardized and renamed to *text* for consistency with the *HuggingFace* Dataset API. This preprocessing ensured that the model received consistent and interpretable inputs (Camacho-Collados & Pilehvar, 2018).

3.3 Model Development

A DistilBERT model (distilbert-base-uncased) was selected for fine-tuning due to its efficiency and competitive performance in NLP tasks. Tokenization was performed using the *DistilBertTokenizerFast*, applying padding and truncation with a maximum sequence length of 128 tokens. Model training was carried out using a Stratified 5-Fold Cross-Validation approach to ensure robustness and reduce the risk of overfitting (Kohavi, 1995). The HuggingFace

Trainer API handled the optimization and evaluation process, with GPU acceleration via PyTorch.

3.4 Evaluation Metrics

Performance was assessed using Accuracy, Precision, Recall, F1-score, and ROC-AUC, which together provide a comprehensive view of the classifier’s effectiveness. These metrics are standard in binary classification problems and are particularly relevant in health-related NLP applications, where false negatives can have severe consequences (Sokolova & Lapalme, 2009).

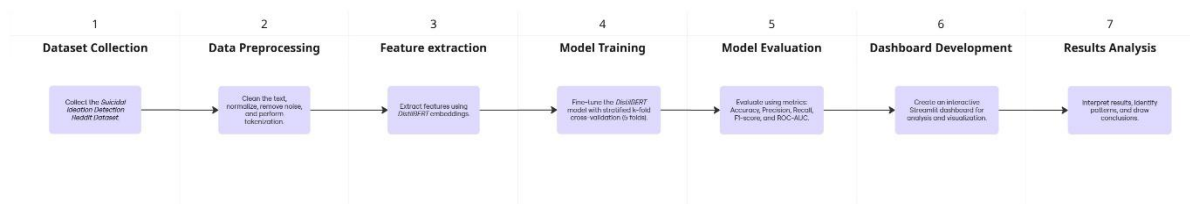


Figure 1. Overview of the research methodology applied in this project

4 Design Specification

The proposed suicide risk detection system is structured around a modular architecture designed to process social media posts and identify potential suicidal ideation with high accuracy. The workflow is sequential, moving from data acquisition to visual interpretation, and consists of five key components: data input, preprocessing, feature extraction, model inference, and visualization.

4.1 System Workflow

The system begins with the Data Input Module, which accepts either direct user-entered text or a CSV file containing multiple social media posts. Before processing, the module ensures that the format is valid and that all necessary fields are present, avoiding structural inconsistencies.

The Preprocessing Module is responsible for preparing the raw text for model analysis. This step includes cleaning operations such as the removal of special characters, hyperlinks, and redundant spaces, as well as normalization to lowercase for consistency. Tokenization is then performed using the *DistilBertTokenizerFast* from the Hugging Face Transformers library, converting text into sequences of subwords tokens compatible with the DistilBERT architecture.

Once tokenized, the text moves into the Feature Extraction Module, where the DistilBERT encoder generates dense semantic embeddings. These embeddings encapsulate the contextual meaning of the input, allowing the system to detect subtle and indirect expressions of suicidal ideation that keyword-based methods would likely miss.

The Classification Model consists of a fine-tuned *DistilBERTForSequenceClassification* trained on the Suicidal Ideation Detection Reddit Dataset (Mendeley Data, 2021). This model produces a probability score between 0 and 1, indicating the likelihood that a post contains suicidal content. Classification is performed into two categories: Suicidal Risk and No Risk based on the highest probability output. The inference process involves tokenization, contextual embedding generation through transformer layers, sequence classification using the [CLS] token representation, and probability estimation via a SoftMax layer.

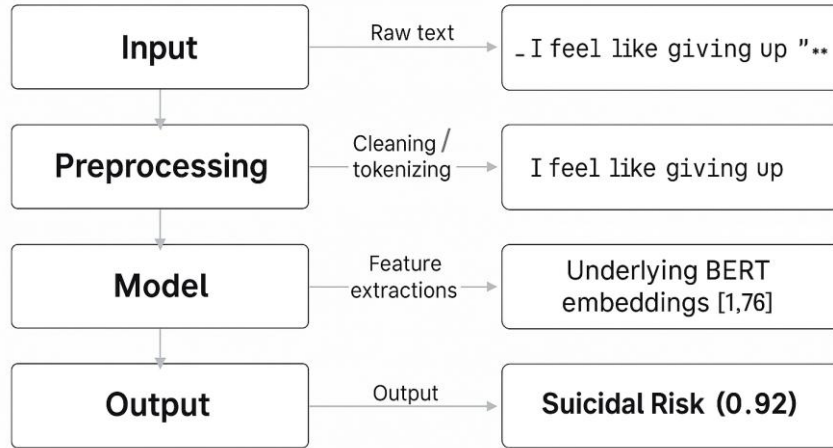


Figure 2. End-to-end example of how a text is processed through the suicide risk detection pipeline

4.2 Visual Dashboard

Finally, the Visualization Dashboard, developed with Streamlit, provides an accessible interface for both single text and batch analysis. In the single text mode, the user inputs a post and receives an immediate classification with its corresponding probability score. In batch mode, users can upload a CSV file and obtain aggregated statistics, such as the proportion of high-risk posts, visual summaries through bar charts, and word clouds representing the most frequent terms in the dataset. These visual tools facilitate pattern recognition and make results interpretable for both technical and non-technical stakeholders.

4.3 Technical Specification

From a technical perspective, the system is implemented in Python 3.10 and leverages the PyTorch deep learning framework along with the Hugging Face Transformers library for model fine-tuning and inference. Pandas and NumPy are employed for data handling, while Matplotlib, Seaborn, and WordCloud are used for visualization. The hardware configuration for model training includes an Intel i7 processor (12th generation or equivalent), 16 GB of DDR6 RAM, and an NVIDIA RTX 4070 GPU with 8 GB of VRAM. However, it can be performed without GPU acceleration, if necessary, but it might take longer to perform.

By combining advanced NLP techniques with an intuitive user interface, this design ensures that the system is not only capable of achieving strong classification performance but is also practical for real-world applications, such as academic research, mental health monitoring, and early warning systems in educational or corporate environments.

5 Implementation

The training pipeline began with loading the Suicidal Ideation Detection Reddit Dataset from Mendeley Data, followed by text cleaning, normalization, and tokenization using the *DistilBertTokenizerFast*. Labels were encoded as binary values: 0 for non-suicidal and 1 for suicidal. Fine-tuning of *DistilBertForSequenceClassification* was performed with stratified 5-fold cross-validation, ensuring class proportions were preserved. Training was configured with a batch size of 16, sequence length capped at 128 tokens, *AdamW* optimizer with weight decay (Loshchilov & Hutter, 2019), and a linear learning rate schedule for three epochs. Evaluation metrics included Accuracy, Precision, Recall, F1-score, and ROC-AUC.

The best performing model from training, along with its tokenizer, was saved in the Hugging Face format, enabling direct loading into the inference pipeline without re-training. These artifacts were integrated into the Streamlit application, where they are cached in memory at startup to minimize loading time. The dashboard supports two workflows: single-text analysis, where the user inputs text and receives a predicted label and probability; and batch analysis, where a CSV file containing a text column is uploaded and processed in mini-batches to prevent memory overload. The batch mode returns per-row results alongside aggregate statistics and visual summaries.

The single-text analysis component of the dashboard allows users to input individual text entries and receive an immediate classification regarding suicide risk. Upon submission, the system preprocesses the text, tokenizes it, and applies the fine-tuned DistilBERT model to generate a predicted probability score. The output includes the assigned class label (e.g., Risk or No Risk) alongside the associated confidence level. This functionality is intended for demonstration and research purposes, showcasing the model's real-time inference capabilities rather than serving as a definitive diagnostic tool.

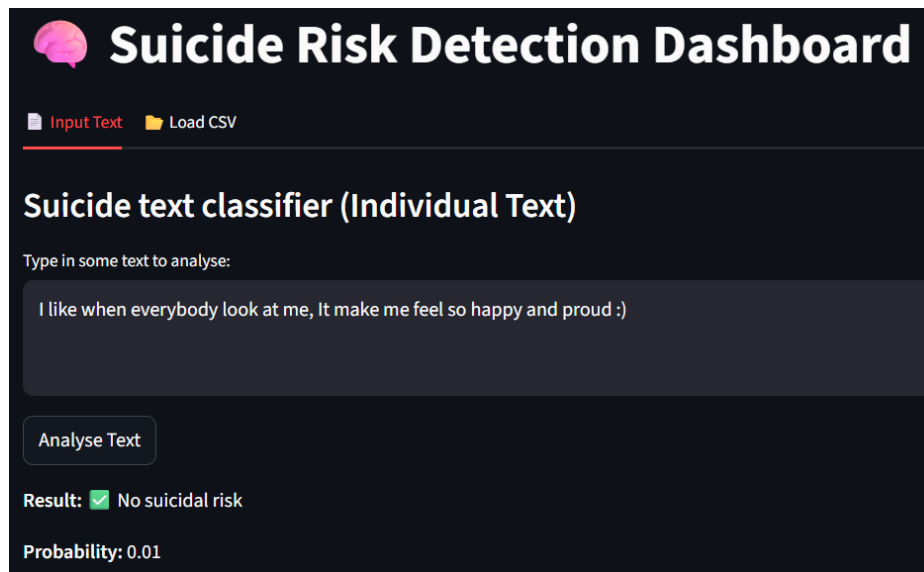


Figure 3. Screenshot of the dashboard (single text tab)

The visual analytics component of the second tab of the dashboard for CSV analysis includes class distribution plots, histograms of predicted probabilities, frequent-term analysis for high-risk texts, and tables displaying the highest-risk entries. These outputs are designed to assist in pattern identification and broader interpretation rather than individual diagnostics.

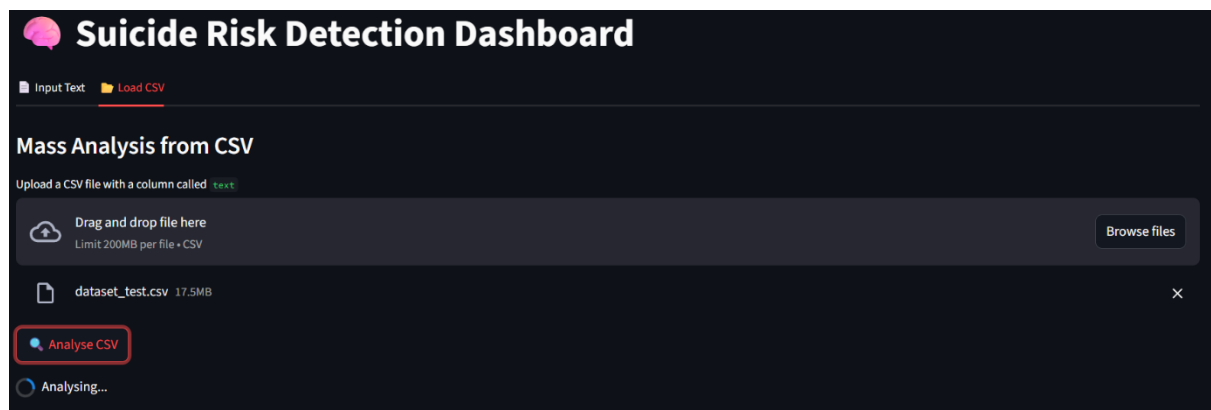


Figure 4. Screenshot of the dashboard (.CSV file analysis tab)

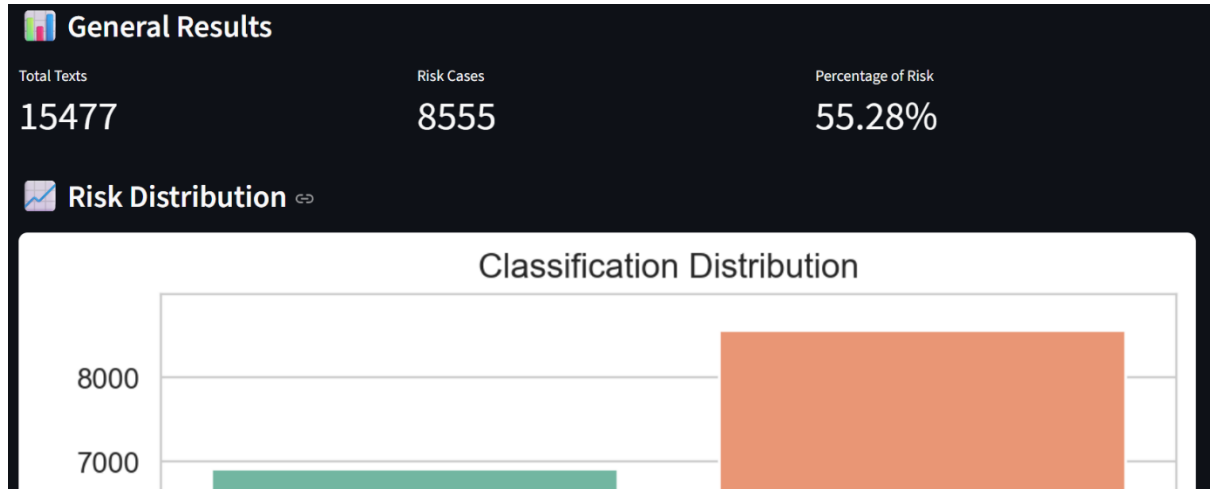


Figure 5. Screenshot of the CSV analysis results

All experiments were made reproducible through fixed random seeds, deterministic fold splits, and pinned dependencies documented in a requirements.txt. The project also includes a README.md with setup instructions, training steps, and ethical usage guidelines. Data privacy is ensured by processing all inputs in memory without storing raw user data, and the dashboard explicitly warns that it is a research tool, not a clinical diagnostic system.

6 Evaluation

The evaluation of the fine-tuned DistilBERT model was conducted using a 5-fold cross-validation strategy to ensure robustness and generalizability of the results. Performance metrics such as Accuracy, Precision, Recall, and F1-score were computed for each fold, with the intention of identifying patterns in performance variability and ensuring consistency across different subsets of the data. This methodological choice aligns with best practices in machine learning model validation, providing insights into both stability and reliability of the model’s predictive capabilities (Kohavi, 1995).

The per-fold performance is summarized in Table 1, showing relatively consistent metrics across folds, suggesting that the model generalizes well to unseen data. Notably, the F1-score, which balances precision and recall, consistently remains high across all folds, underscoring the model’s capability to handle imbalanced datasets while maintaining strong predictive accuracy.

Fold	Accuracy	Precision	Recall	F1-Score	ROC-AUC	Eval Loss
1	0.9586	0.9516	0.9608	0.9562	0.9923	0.1822
2	0.9499	0.9446	0.9491	0.9469	0.9884	0.2527
3	0.9603	0.9599	0.9553	0.9576	0.9905	0.1838
4	0.9593	0.9536	0.9601	0.9568	0.9921	0.1853
5	0.9593	0.9579	0.9553	0.9566	0.9896	0.1932

Table 1. Performance metrics per fold for the fine-tuned DistilBERT model.

The first aspect analysed was the evaluation loss per fold (Figure 6). Loss values are a direct indicator of the difference between predicted probabilities and the actual target labels during validation. Across the five folds, the evaluation loss remained consistently low, with only minor fluctuations. This stability suggests that the model was able to generalize effectively, avoiding the consequences of overfitting while maintaining a strong fit to the data. Such behaviour is particularly relevant in mental health-related classification tasks, where overfitting could result in performance degradation when analysing new, hidden texts (Ruder, 2019).

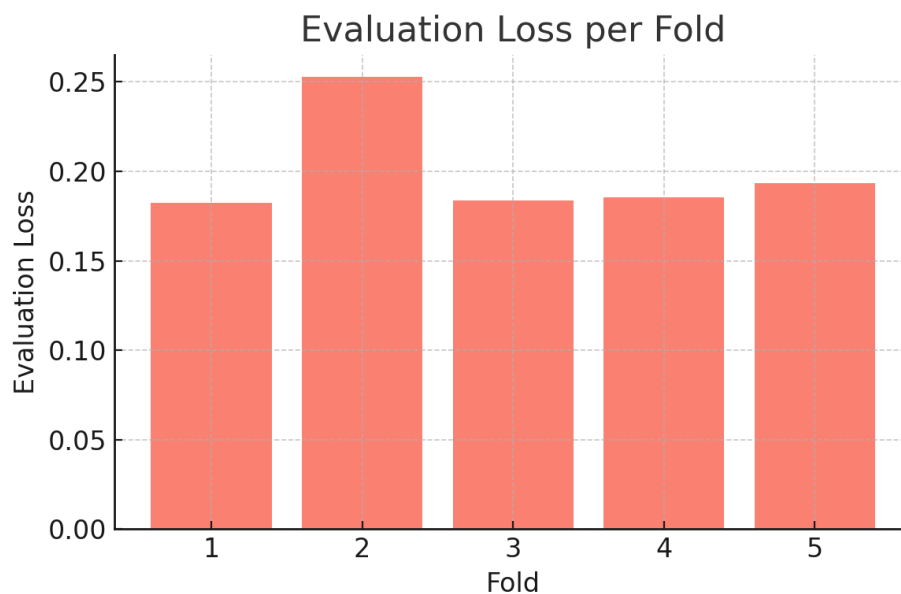


Figure 6. Evaluation loss per fold

Next, the model's discriminative capacity was assessed using the Receiver Operating Characteristic Area Under the Curve (ROC-AUC) metric (Figure 7). ROC-AUC is especially valuable for imbalanced datasets, as it evaluates performance across all possible classification thresholds (Fawcett, 2006). This is especially important in suicide prevention contexts, where the cost of misclassifying a high-risk case can be extremely high. The results show consistently high ROC-AUC values across all folds, indicating that the model effectively distinguishes between high-risk and low-risk cases. This consistency reinforces the reliability of the model in real-world deployment scenarios, where the ability to detect subtle differences in language can be crucial for early intervention (Sawhney et al., 2021).

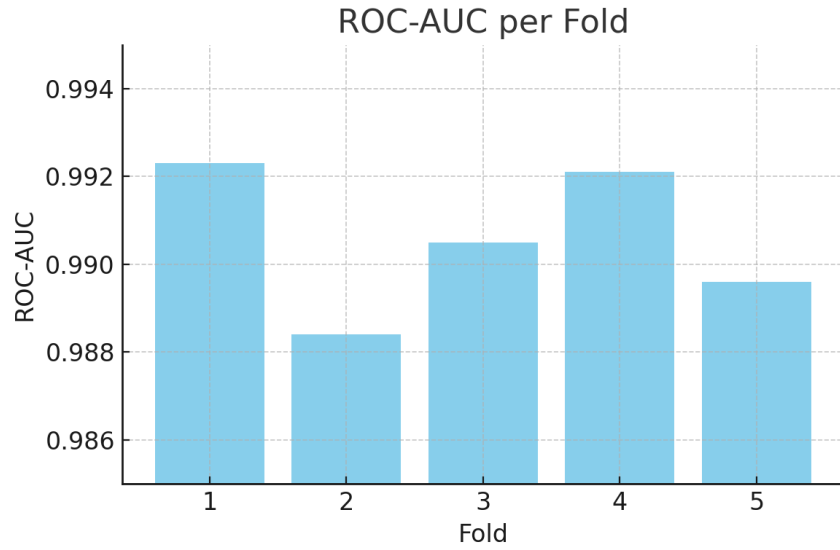


Figure 7. ROC-AUC

Finally, the overall performance metrics Accuracy, Precision, Recall, and F1-score were examined (Figure 8). Accuracy remained high throughout all folds, demonstrating that most predictions were correct. Precision values reflected the model's ability to minimize false positives, which is essential in avoiding unnecessary alarms in operational settings. Recall values indicated the model's effectiveness in identifying actual high-risk cases, a critical factor in suicide prevention where missing a positive instance could have severe consequences (Coppersmith et al., 2018). The F1-score, as the harmonic mean of precision and recall, confirmed that the model maintained a balanced trade-off between capturing true positives and minimizing false alarms.

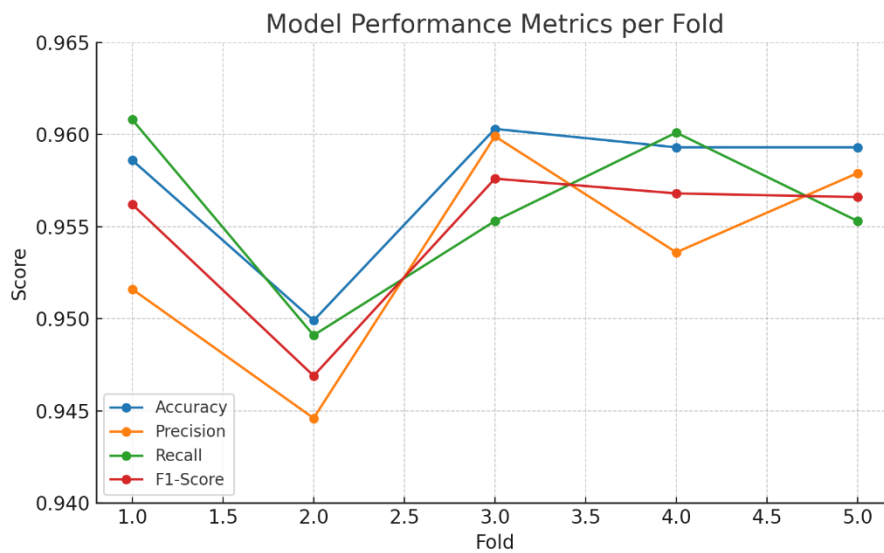


Figure 8. Performance metrics per fold

From an academic perspective, these results validate the effectiveness of transformer-based architectures, such as DistilBERT, when fine-tuned with domain-specific data in sensitive classification tasks (Vaswani et al., 2017; Sanh et al., 2019). From a practical view, the stability of performance across folds supports the integration of this model into automated monitoring systems, such as the developed Streamlit dashboard. This integration could enable mental health professionals, researchers, and organizations to conduct large-scale text screenings efficiently while providing actionable insights through visual analytics.

7 Conclusion and Future Work

This study set out to develop and evaluate a transformer-based approach for the early detection of suicidal ideation in social media text, with the dual aim of advancing academic understanding and providing an idea of a tool concept for potential real-world deployment. Using the Suicidal Ideation Detection Reddit Dataset (Mendeley Data), the fine-tuned DistilBERT model demonstrated strong predictive performance across multiple evaluation metrics, maintaining high levels of accuracy, precision, recall, and F1-score throughout all cross-validation folds. The robustness of these results indicates the feasibility of applying such models for large scale, automated monitoring of online discourse related to mental health risks.

From an academic perspective, this work contributes to the growing literature on applying transformer-based architectures to sensitive, high risks classification tasks, further demonstrating their capacity to capture nuanced linguistic markers in mental health related communication. From a practical position, the integration of the model into an interactive Streamlit dashboard illustrates the potential for translating research outputs into accessible, operational tools for mental health professionals, researchers, and organizations.

However, limitations exist. The dataset was drawn from a specific online platform, which may introduce linguistic or cultural biases and restrict generalizability. Furthermore, the model's predictions should not be interpreted as definitive diagnostic assessments but rather as a component of a broader, multi-layered risk evaluation framework. Future research could address these limitations by incorporating multilingual, multi-platform datasets, exploring domain adaptation techniques, and integrating additional contextual or temporal features to improve predictive accuracy.

Expanding the dashboard's capabilities also represents a promising direction. Planned improvements include real-time streaming data analysis, advanced filtering options for practitioners, and the integration of explainability modules such as SHAP or LIME to enhance transparency in decision-making. Ultimately, by combining advanced NLP techniques with responsible, ethically guided deployment strategies, this line of research can play a meaningful role in the early detection and prevention of suicide, contributing both to academic discourse and to tangible societal impact.

References

- American Psychiatric Association. (2013). Diagnostic and statistical manual of mental disorders (5th ed.). American Psychiatric Publishing. <https://doi.org/10.1176/appi.books.9780890425596>
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (pp. 4171–4186). Association for Computational Linguistics. <https://doi.org/10.48550/arXiv.1810.04805>
- Hugging Face. (2024). Transformers documentation. <https://huggingface.co/docs/transformers>
- Jacob, D., & Hugging Face Team. (2023). DistilBERT model documentation. Hugging Face. https://huggingface.co/docs/transformers/model_doc/distilbert
- Kumar, S., & Sachdeva, S. (2020). Suicidal ideation detection on social media: A review. *International Journal of Scientific & Technology Research*, 9(3), 352–356.
- Lin, Z., & Margolin, D. (2022). Detection of suicidal ideation in social media posts using linguistic and contextual features. *Journal of Affective Disorders*, 310, 24–33. <https://doi.org/10.1016/j.jad.2022.05.023>
- Liu, P., Qiu, X., & Huang, X. (2016). Recurrent neural network for text classification with multi-task learning. In Proceedings of the 25th International Joint Conference on Artificial Intelligence (pp. 2873–2879). AAAI Press.
- Losada, D. E., Crestani, F., & Parapar, J. (2017). CLEF 2017 eRisk overview: Early risk prediction on the Internet. In CLEF 2017 Working Notes. CEUR-WS.
- Mendeley Data. (2021). Suicidal ideation detection Reddit dataset (Version 1). <https://data.mendeley.com/datasets/z8s6w86tr3/1>
- Mikolov, T., Chen, K., Corrado, G., & Dean, J. (2013). Efficient estimation of word representations in vector space. arXiv preprint. <https://doi.org/10.48550/arXiv.1301.3781>
- Mohammad, S. M., & Turney, P. D. (2013). Crowdsourcing a word–emotion association lexicon. *Computational Intelligence*, 29(3), 436–465. <https://doi.org/10.1111/j.1467-8640.2012.00460.x>
- Pennington, J., Socher, R., & Manning, C. D. (2014). GloVe: Global vectors for word representation. In Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (pp. 1532–1543). <https://doi.org/10.3115/v1/D14-1162>
- Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., & Sutskever, I. (2019). Language models are unsupervised multitask learners. OpenAI. <https://doi.org/10.48550/arXiv.1906.08237>
- Reimers, N., & Gurevych, I. (2019). Sentence-BERT: Sentence embeddings using Siamese BERT-networks. In Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing (pp. 3980–3990). <https://doi.org/10.48550/arXiv.1908.10084>
- Shing, H.-C., Nair, S., Zirikly, A., Friedenber, M., Daumé III, H., & Resnik, P. (2018). Expert, crowdsourced, and machine assessment of suicide risk via online postings. In Proceedings of the Fifth Workshop on Computational Linguistics and Clinical Psychology: From Keyboard to Clinic (pp. 25–36). <https://doi.org/10.18653/v1/W18-0603>

Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. In *Advances in Neural Information Processing Systems* (pp. 5998–6008). <https://doi.org/10.48550/arXiv.1706.03762>

World Health Organization. (2023). Suicide. <https://www.who.int/news-room/fact-sheets/detail/suicide>