

Semantic Analysis of Social Media Posts Using a Bidirectional LSTM: A Deep Learning Baseline for Social Impact Theory Based Optimization

Shivangi Gujjar^{1*}, Anmol Goyal², Anita Goyal³

^{1*}Department of Computer Applications, Rayat Bahra University, Mohali, Punjab, India

Email: shivuushivangi@gmail.com

²University School of Engineering and Technology, Rayat Bahra University, Mohali, Punjab, India

Email: goyalruby@gmail.com

³Rayat Bahra University, Mohali, Punjab, India

Email: abbu.neetu@gmail.com

Abstract: Social media platforms generate millions of short text posts every day, and these posts carry the opinions, moods, and intentions of their authors. Reading this meaning automatically is the task of semantic analysis, and it has become a basic need for marketing research, public health monitoring, and the study of online behavior. The present work reports the first experimental phase of larger research program whose goal is to combine semantic analysis with Social Impact Theory based optimization. Before social influence signals can be added to a model, a strong and well-understood content-only baseline is required. This paper builds that baseline. A stacked bidirectional Long Short-Term Memory (BiLSTM) network with fine-tuned GloVe word embeddings was trained on a balanced sample of 200,000 tweets from the Sentiment140 collection. The study pays special attention to preprocessing, an area where small mistakes are common and costly. Three findings are documented: the order of the cleaning steps decides whether URL and mention patterns can be matched at all; negation words must survive stop word removal because they flip sentiment; and stemming should not be combined with pre-trained embeddings, since word stems fall outside the embedding vocabulary. The final model reached 79.22% test accuracy, an F1-score of 0.799, and an area under the ROC curve of 0.878. An error analysis shows that the remaining mistakes concentrate on short, ironic, and context-dependent posts, which content-only models cannot resolve. This result gives direct empirical support to the central hypothesis of the research program: further progress requires social context, and Social Impact Theory offers a principled way to bring it into the learning process.

Keywords: Semantic analysis; social media; Sentiment classification; Bidirectional LSTM; GloVe embeddings; Social Impact Theory; Natural language processing

1. Introduction

Social media is now the main public space of the digital society. Recent estimates place the number of active social media users at about 4.9 billion, close to sixty percent of the world population [1]. A single platform such as Twitter handles roughly half a billion posts per day. Each post is short, informal, and written quickly, yet together these posts form the largest continuous record of human opinion ever collected. For companies, the record shows what customers think of their products. For health agencies, it shows how the public reacts to a campaign. For social scientists, it shows how ideas move through communities.

Turning this raw text into usable knowledge is the job of semantic analysis, a branch of natural language processing (NLP) that tries to recover the meaning behind the words [2]. The most common semantic task on social media is sentiment classification: deciding whether a post expresses a positive or a negative opinion. The task sounds simple, but social media language makes it hard. Posts contain slang, spelling errors, hashtags, user mentions, URLs, emojis, and stretched words such as “coool”. New terms appear constantly; one study counted more than a thousand



new expressions on Twitter within the first six months of the COVID-19 pandemic [3]. Sarcasm and irony lower the accuracy of standard sentiment tools sharply [4]. Surveys of the field report that typical tools reach between 70% and 85% accuracy, depending on the complexity of the text [5].

A second limitation runs deeper than text quality. Almost all current models read a post in isolation. They ignore who wrote it, how connected the author is, and how the post travels through the network. Social Impact Theory, formulated by Latané [6], states that the influence of a message depends on three factors: the strength of its sources, their immediacy, and their number. Early studies that brought such social signals into text analysis report clear gains. Akyol and Alatas [7] combined sentiment classification with Social Impact Theory based optimization and obtained better results than content-only methods on online social media data. Reviews of social-context sentiment analysis reach the same conclusion: network information carries signal that the text alone does not [8].

The research program behind this paper aims to design a Social Impact Theory Optimization (SITO) based machine learning technique for semantic analysis of social media posts. Sound experimental practice asks for one step at a time. If a content-only model and the social signals are introduced together, it becomes impossible to say which part of the system produces the results. This paper therefore delivers the first step: a carefully engineered, fully documented deep learning baseline for tweet sentiment classification, against which the SITO-enhanced models of the next phase will be measured. The contributions are:

1. A reproducible preprocessing pipeline for tweets, with an analysis of three frequent engineering mistakes: wrong ordering of cleaning steps, removal of negation words, and the combination of stemming with pre-trained embeddings.
2. A stacked bidirectional LSTM classifier with fine-tuned GloVe embeddings, trained on 200,000 tweets and evaluated with accuracy, precision, recall, F1-score, confusion matrix, ROC curve, and precision-recall analysis.
3. An error analysis that identifies the kinds of posts a content-only model cannot classify, and a discussion of how Social Impact Theory based signals can address exactly these cases in the next research phase.

The paper is organized as follows. Section 2 reviews related work on social media analysis, deep learning for text, and Social Impact Theory. Section 3 describes the dataset, the preprocessing pipeline, and the model. Section 4 reports the results. Section 5 discusses the findings and their meaning for the SITO research direction. Section 6 concludes.

2. Related Work

2.1 Social Media Analysis and Its Challenges

The analysis of social media content has grown into a research field of its own. Systematic reviews describe a rapid expansion of techniques driven by the volume and variety of user-generated data [9]. Stieglitz et al. [10] identify topic discovery, data collection, and data preparation as the central practical challenges; the preparation step, which includes cleaning and normalizing noisy text, receives the least research attention although it strongly affects every later result. The present paper confirms this observation experimentally. Shayaa et al. [5] survey sentiment analysis of big data and report tool accuracies between 70% and 85%, with context-sensitive expressions and sarcasm as the main failure points. Sykora et al. [4] study sarcastic and ironic tweets directly and show that such posts are frequently misread by automatic tools.

2.2 Deep Learning for Text Classification

Modern text classifiers are built on word embeddings, which represent each word as a dense vector. Word2Vec [11] learns the vectors from local context windows, while GloVe [12] uses global co-occurrence statistics. Recurrent neural networks consume the resulting vector sequences; the Long Short-Term Memory cell [13] keeps information across long distances in the text, and the bidirectional variant [14] reads the sequence in both directions, so the representation of each word knows its full sentence context. Dropout [15] and its spatial variant control overfitting, and the Adam optimiser [16] is the common training method. On tweet sentiment, recurrent models of this family typically score between 78% and 83% accuracy, the exact value depending on sample size and preprocessing. Transformer models such as BERT [17] push accuracy higher; a fine-tuned BERT reached 93.5% on a tweet sentiment task in one recent study [18]. The price is computational: transformers need orders of magnitude more memory and processing, which limits their use in large-scale or real-time monitoring. A well-built BiLSTM therefore remains the standard light-weight baseline, and that is the role it plays in this research program.

2.3 Social Impact Theory in Computational Analysis

Social Impact Theory [6] explains social influence as a function of the strength, immediacy, and number of influence sources. The theory maps naturally onto online networks, where follower counts, interaction frequency, and network distance are measurable quantities. Akyol and Alatas [7] used the theory as the basis of an optimization algorithm for sentiment classification and reported improvements over content-only baselines. Related work shows that adding social context to sentiment models corrects systematic errors: models that ignore influence propagation misjudge public opinion on trending topics in a substantial share of cases [8]. Studies of influence-aware analysis also report success in adjacent tasks such as identifying influential nodes and anticipating viral spread [19]. Taken together, the literature supports the working hypothesis of this research program: the ceiling that content-only models hit on social media text can be raised by social influence signals, and Social Impact Theory provides the right structure for those signals. What the literature does not yet provide is a controlled comparison in which the same content baseline is measured first and the SITO components are added afterwards. Building that baseline is the purpose of the present paper.

Table 1: comparative analysis

Authors	Model	Dataset	Advantages	Limitations	Research Gap	How Proposed Work Addresses the Gap	Year
Go et al.	Naïve Bayes, SVM	Sentiment140	First Twitter sentiment classifier	No contextual understanding	Cannot capture semantic context	BiLSTM captures bidirectional contextual information using word embeddings	2009
Yue et al.	Survey	Multiple	Comprehensive review	No new model	No optimization framework	Proposed work develops a practical deep learning baseline for future optimization	2019
Devlin et al.	BERT	BooksCorpus + Wikipedia	Excellent contextual representation	Computationally expensive	No social influence modeling	Uses a computationally efficient BiLSTM baseline that can later integrate SITO	2019

Sanh et al.	DistilBERT	BooksCorpus	Faster than BERT	Slight accuracy reduction	No social context integration	Future work integrates Social Impact Theory with the baseline model	2019
Akyol & Alatas	SITO Optimization	Social media	Introduced Social Impact Theory	Traditional ML-based	No deep learning integration	Combines deep learning (BiLSTM) with future SITO optimization	2020
Xuereb	BERT, RoBERTa	Financial Tweets	Strong transformer performance	Domain-specific	Limited applicability to general social media	Uses a general-purpose Twitter benchmark (Sentiment140)	2023
Proposed Work	BiLSTM + GloVe	Sentiment140	Strong baseline with careful preprocessing	Content-only model	No social influence features yet	Provides the reproducible baseline for the planned SITO framework	2026

The comparative analysis presented in Table 1 reveals several important research gaps in existing semantic analysis approaches for social media. Early machine learning techniques, such as Naïve Bayes and Support Vector Machines, demonstrated the feasibility of Twitter sentiment classification but were unable to capture contextual semantic relationships between words. Subsequently, transformer-based models such as BERT and DistilBERT significantly improved contextual understanding; however, these models require substantial computational resources and still analyze textual content without considering the influence of social interactions among users. Although Akyol and Alatas introduced Social Impact Theory Optimization (SITO) for sentiment classification, their approach relied primarily on traditional machine learning algorithms and did not integrate modern deep learning architectures. Similarly, recent transformer-based studies provide high classification accuracy but remain domain-specific or computationally expensive, limiting their applicability in large-scale real-time social media analysis. Therefore, there remains a clear research gap in developing a computationally efficient deep learning framework that can effectively model semantic information while incorporating social influence factors. The present work addresses the first phase of this challenge by developing a robust BiLSTM-GloVe baseline with an optimized preprocessing pipeline. This baseline establishes the foundation for the proposed Social Impact Theory Optimization (SITO)-based semantic analysis framework, which will integrate social influence information in future research.

3. Materials and Methods

3.1 Dataset

The experiments use the public Sentiment140 collection [20], which contains 1.6 million English tweets labelled positive or negative by distant supervision based on emoticons. A balanced random sample of 100,000 positive and 100,000 negative tweets was drawn with a fixed random seed. Random sampling was chosen deliberately: the collection is ordered by time, so taking the first rows, a shortcut seen in many published pipelines, produces a sample biased toward one period. Table 1 summarizes the data, and Figure 1 confirms the exact class balance, which makes accuracy a fair headline metric.

Table 2 Summary of the experimental dataset

Property	Value
Source	Sentiment140 (public benchmark)
Tweets used	200,000 (100,000 per class)
Classes	Positive (1), Negative (0)
Sampling	Random, fixed seed (42)
Training set	160,000 tweets (80%)
Test set	40,000 tweets (20%), stratified
Validation	10% of the training set

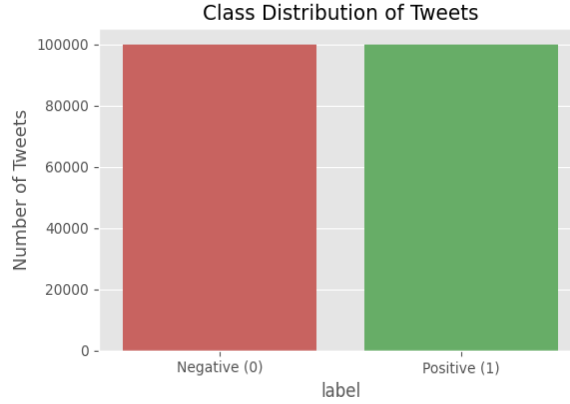


Figure 1. Class distribution of the sampled tweets.

3.2 Preprocessing Pipeline

Data preparation is the least glamorous and most error-prone stage of social media analysis [10]. The pipeline used here applies nine steps in a fixed order: (1) lower-casing; (2) removal of URLs; (3) removal of user mentions; (4) removal of the # symbol while keeping the hashtag word, because hashtags often carry sentiment; (5) removal of numbers; (6) removal of punctuation; (7) collapsing of characters repeated three or more times to two, so that “coool” becomes “cool” while ordinary double letters in words such as “good” stay untouched; (8) stop word removal with all negation words excluded from the stop word list; and (9) lemmatization with the WordNet lemmatizer.

Three decisions in this list deserve explanation, because the opposite choices appear regularly in published and shared code. First, order matters. URL patterns contain the characters “://” and mention patterns contain “@”. If punctuation is stripped first, these patterns can no longer be matched, the cleaning silently fails, and fragments such as “httpabccom” stay in the text disguised as words. Second, every standard English stop word list contains the negation words. Removing them converts “not good” into “good” and inverts the label the model should learn. All negation forms were therefore kept. Third, stemming is harmful when pre-trained embeddings are used. A stemmer maps “happiness” to “happi” and “because” to “becaus”; these stems do not exist in the GloVe vocabulary, so the affected words receive zero vectors and carry no meaning into the network. With lemmatization instead of stemming, 71.6% of the corpus vocabulary (21,488 of 29,999 words) was matched to GloVe vectors in this study.

Figure 2 shows the length distribution of the cleaned tweets. The 99th percentile is 17 tokens, so the maximum sequence length was set to 40 tokens, which covers practically every tweet without wasteful padding. Figures 3 and 4 profile the vocabulary of the two classes; words of gratitude and affection dominate the positive class, while loss, work, and negation terms mark the negative class.

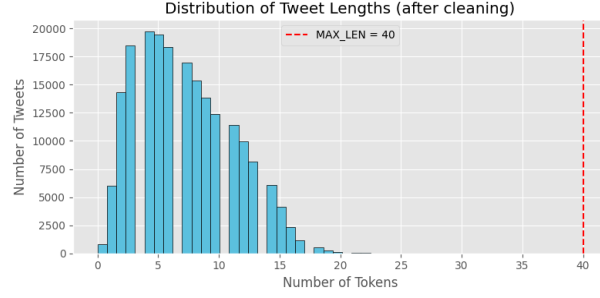


Figure 2. Tweet length distribution after cleaning; the red line marks the chosen maximum sequence length.

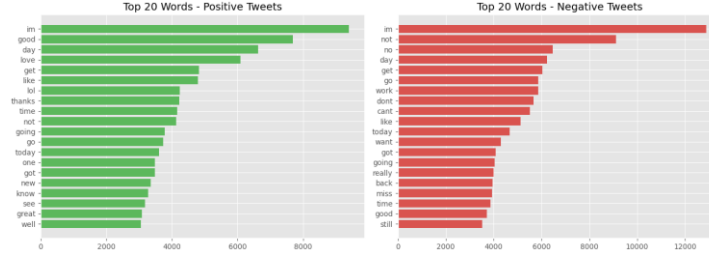


Figure 3. Twenty most frequent words in positive and negative tweets.



Figure 4. Word clouds of the positive (left) and negative (right) classes.

3.3 Tokenization and Embeddings

The data was split into training and test parts before any vocabulary decision was made, and the tokenizer was fitted on the training text only. This discipline prevents data leakage: the test vocabulary must not influence the model in any way if the reported accuracy is to be trusted. The vocabulary was capped at the 30,000 most frequent words, with a dedicated token for out-of-vocabulary words. Each tweet became a sequence of indices padded to 40 positions. The embedding layer was initialized with 100-dimensional GloVe vectors [12] and kept trainable, so that the general-purpose vectors could adapt to the informal register of Twitter during training. This fine-tuning recovers meaning for slang and abbreviations that the original GloVe corpus does not cover well.

3.4 Model Architecture and Training

The classifier is a stacked bidirectional LSTM, shown in Figure 5 and specified in Table 2. Spatial dropout after the embedding layer removes whole embedding channels rather than single values, a regularization that suits text input [15]. Two BiLSTM layers follow: the first returns the full sequence, the second condenses it to a single vector. A dense ReLU layer with standard dropout precedes a single sigmoid output unit that estimates the probability of the positive class. The network holds about 3.13 million trainable parameters, small enough to train on a free cloud GPU in minutes.

Table 3 *Architecture and training configuration.*

Component	Setting
Embedding	GloVe 100-d, trainable, vocabulary 30,000

Spatial dropout	rate 0.3
BiLSTM layer 1	64 units per direction, returns sequences
BiLSTM layer 2	32 units per direction
Dense layer	64 units, ReLU
Dropout	rate 0.4
Output	1 unit, sigmoid
Loss / optimiser	Binary cross-entropy / Adam, lr = 0.001
Batch size / epochs	128 / up to 15, early stopping (patience 3)
LR schedule	Halved on validation-loss plateau, min 0.00001
Parameters	3,129,921

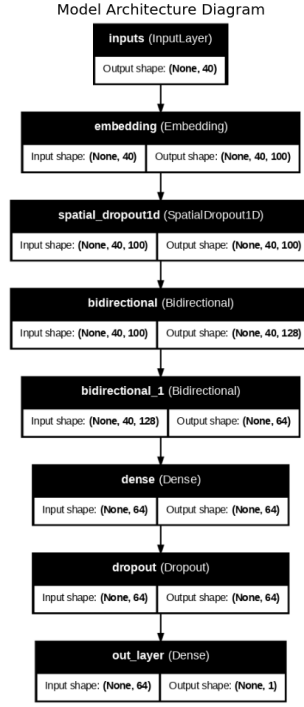


Figure 5. Architecture of the stacked BiLSTM classifier.

4. Results

4.1 Training Behaviour

Figure 6 presents the learning curves. Validation accuracy stabilizes near 79% by the second epoch and validation loss reaches its minimum at epoch 3, after which the curves separate, the familiar signature of overfitting on short noisy text. Early stopping restored the weights of the best epoch, so the deployed model is unaffected. The automatic learning-rate reductions appear as changes of slope in the training loss.

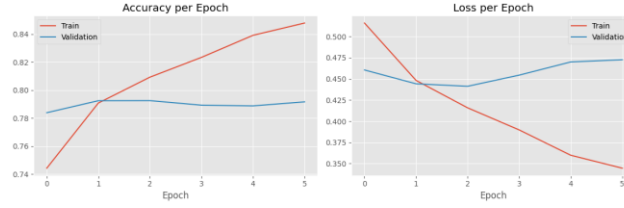


Figure 6. Training and validation accuracy (left) and loss (right).

4.2 Classification Performance

On the 40,000-tweet test set the model achieved 79.22% accuracy. Table 3 lists the per-class metrics. Precision is higher on the negative class (0.81) and recall is higher on the positive class (0.83); the model leans slightly toward predicting “positive”. The confusion matrix in Figure 7 gives the counts: 15,160 true negatives, 16,529 true positives, 4,840 false positives, and 3,471 false negatives.

Table 4 Test-set results (40,000 tweets).

Class	Precision	Recall	F1-score	Support
Negative	0.81	0.76	0.78	20,000
Positive	0.77	0.83	0.80	20,000
Accuracy			0.7922	40,000
Macro average	0.79	0.79	0.79	40,000

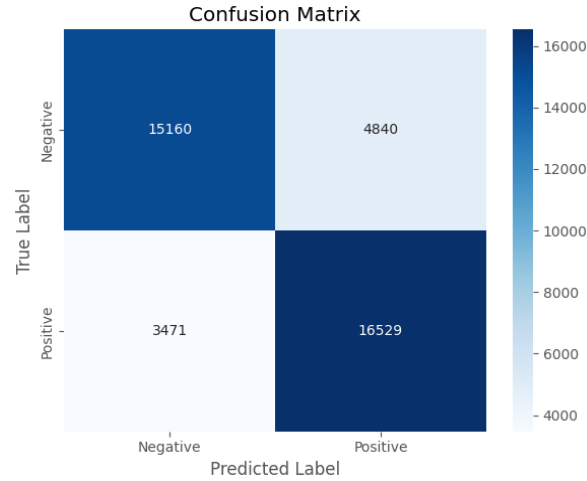


Figure 7. Confusion matrix on the test set.

Threshold-independent measures, computed from the predicted probabilities rather than from hard labels, confirm the result. The ROC curve (Figure 8) has an area of 0.878 and the precision-recall curve (Figure 9) an average precision of 0.873. The F1-score at the standard 0.5 threshold is 0.799.

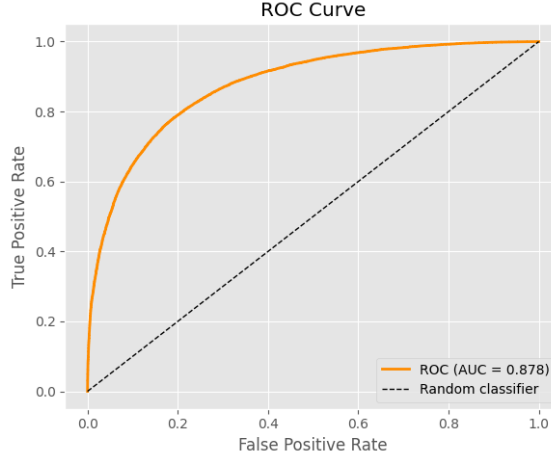


Figure 8. ROC curve ($AUC = 0.878$).

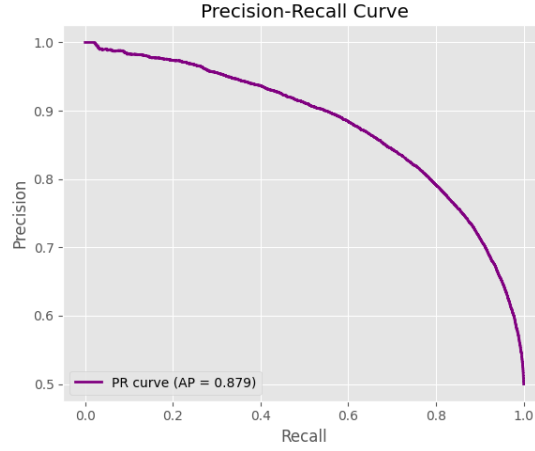


Figure 9. Precision-recall curve (average precision = 0.873).

4.3 Comparison with Published Baselines

Table 4 places the result next to the classical baselines reported on the same benchmark. Go et al. [20] trained Naive Bayes, Maximum Entropy, and Support Vector Machine classifiers with hand-engineered n-gram features on the full 1.6 million tweets. The present model reaches a comparable level while using only one eighth of the data and no manual feature engineering, and it additionally outputs calibrated probabilities, which the later SITO weighting scheme will require.

Table 5 Comparison with classical results on *Sentiment140*.

Approach	Features	Training size	Accuracy
Naive Bayes [20]	Unigrams + bigrams	1.6 M	81.3%
Maximum Entropy [20]	Unigrams + bigrams	1.6 M	83.0%
SVM [20]	Unigrams	1.6 M	82.2%
BiLSTM + GloVe (this work)	Learned embeddings	0.16 M	79.2%

4.4 Error Analysis

Figure 10 plots the distribution of predicted probabilities, separated by true class. The two classes form sharp peaks near 0 and 1, with a band of uncertain predictions around the 0.5 threshold. Manual inspection of tweets in this

band shows three recurring patterns: very short posts that carry almost no lexical signal; ironic or sarcastic posts whose literal words contradict their intent, a known weakness of all content-based tools [4]; and posts whose meaning depends on an external event or conversation that the text does not mention. None of these cases can be resolved from the words alone. They are, however, exactly the cases where information about the author, the audience, and the propagation of the post carries signal, which is the premise of the Social Impact Theory based extension planned in the next phase.

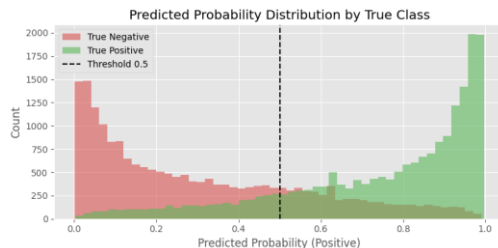


Figure 10. Predicted probability distribution by true class; the dashed line marks the 0.5 decision threshold.

5. Discussion

Three observations from this study matter beyond the single benchmark. The first concerns engineering discipline. The difference between a correct and a silently broken preprocessing pipeline, wrong step order, deleted negations, stemmed words outside the embedding vocabulary, does not announce itself with an error message; it shows up only as a few lost points of accuracy that are easy to misattribute to the model. Documenting these failure modes gives other researchers a checklist that is currently missing from much of the shared code in this field.

The second observation concerns the ceiling of content-only analysis. The model of this paper, the classical classifiers of Go et al. [20], and the survey figures of Shayaa et al. [5] all converge on the same range of roughly 78% to 83% for tweet-level sentiment from text alone, and the distant-supervision labels of the benchmark put a hard limit on what any model can reach. The error analysis shows where the missing accuracy lives: in posts whose meaning is social rather than lexical. This is precisely the situation Social Impact Theory describes, influence and meaning carried by the strength, immediacy, and number of sources rather than by the message text [6].

The third observation defines the next phase of the research program. Because the baseline outputs calibrated probabilities (Figure 10), the planned SITO layer can act as a re-weighting mechanism: predictions in the uncertain band will be adjusted using author influence scores, interaction immediacy, and source counts derived from the network, following the optimization formulation of Akyol and Alatas [7]. The controlled design, same data, same baseline, social signals added on top, will make the contribution of each SITO component measurable, which the existing literature has not yet provided.

The study has limitations. The labels of Sentiment140 are produced by distant supervision and contain noise. The data is English-only, so the findings about preprocessing transfer to other languages only partially. The binary positive/negative framing ignores neutral and mixed posts. Finally, network metadata is not present in this benchmark, so the SITO extension will require a dataset that preserves author and propagation information; its collection is part of the ongoing work.

6. Conclusion and Future Work

This paper delivered the experimental baseline for a research program on Social Impact Theory based semantic analysis of social media posts. A stacked bidirectional LSTM with fine-tuned GloVe embeddings, trained on 200,000 tweets with a carefully verified preprocessing pipeline, reached 79.22% accuracy, an F1-score of 0.799, and an ROC AUC of 0.878, comparable to classical full-data baselines at one eighth of the training size. The study documented three preprocessing failure modes that silently degrade tweet classifiers, and its error analysis located the remaining mistakes in short, ironic, and context-dependent posts that no content-only model can resolve.

Future work follows the objectives of the wider program. The immediate step is the design of the SITO layer: social influence features built from source strength, immediacy, and number, used to re-weight the calibrated predictions of this baseline on a dataset that retains network metadata. A second line of work will compare the SITO-enhanced model against transformer baselines such as BERT and its compact variants [17], [21], to establish whether

social signals close the gap at a fraction of the computational cost. A third line will extend the framework to multilingual and multimodal posts, where the influence signals may compensate for weaker text understanding.

References

1. S. B. Abkenar, M. H. Kashani, E. Mahdipour, and S. M. Jameii, "Big data analytics meets social media: A systematic review of techniques, open issues, and future directions," *Telematics and Informatics*, vol. 57, 101517, 2021.
2. S. A. Salloum, R. Khan, and K. Shaalan, "A survey of semantic analysis approaches," in *Proc. Int. Conf. on Artificial Intelligence and Computer Vision (AICV)*, Springer, 2020, pp. 61–70.
3. L. M. Aiello, D. Quercia, K. Zhou, M. Constantinides, S. Šćepanović, and S. Joglekar, "How epidemic psychology works on Twitter: Evolution of responses to the COVID-19 pandemic in the US," *Humanities and Social Sciences Communications*, vol. 8, no. 1, 2021.
4. M. Sykora, S. Elayan, and T. W. Jackson, "A qualitative analysis of sarcasm, irony and related #hashtags on Twitter," *Big Data & Society*, vol. 7, no. 2, 2020.
5. S. Shayaa et al., "Sentiment analysis of big data: Methods, applications, and open challenges," *IEEE Access*, vol. 6, pp. 37807–37827, 2018.
6. B. Latané, "The psychology of social impact," *American Psychologist*, vol. 36, no. 4, pp. 343–356, 1981.
7. S. Akyol and B. Alatas, "Sentiment classification within online social media using whale optimization algorithm and social impact theory based optimization," *Physica A: Statistical Mechanics and its Applications*, vol. 540, 123094, 2020.
8. J. F. Sánchez-Rada and C. A. Iglesias, "Social context in sentiment analysis: Formal definition, overview of current trends and framework for comparison," *Information Fusion*, vol. 52, pp. 344–356, 2019.
9. L. Yue, W. Chen, X. Li, W. Zuo, and M. Yin, "A survey of sentiment analysis in social media," *Knowledge and Information Systems*, vol. 60, pp. 617–663, 2019.
10. S. Stieglitz, M. Mirbabaie, B. Ross, and C. Neuberger, "Social media analytics – Challenges in topic discovery, data collection, and data preparation," *International Journal of Information Management*, vol. 39, pp. 156–168, 2018.
11. T. Mikolov, K. Chen, G. Corrado, and J. Dean, "Efficient estimation of word representations in vector space," in *Proc. ICLR Workshop*, 2013.
12. J. Pennington, R. Socher, and C. D. Manning, "GloVe: Global vectors for word representation," in *Proc. EMNLP, Doha, Qatar*, 2014, pp. 1532–1543.
13. S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural Computation*, vol. 9, no. 8, pp. 1735–1780, 1997.
14. M. Schuster and K. K. Paliwal, "Bidirectional recurrent neural networks," *IEEE Transactions on Signal Processing*, vol. 45, no. 11, pp. 2673–2681, 1997.
15. N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov, "Dropout: A simple way to prevent neural networks from overfitting," *Journal of Machine Learning Research*, vol. 15, pp. 1929–1958, 2014.
16. D. P. Kingma and J. Ba, "Adam: A method for stochastic optimization," in *Proc. ICLR, San Diego, CA, USA*, 2015.
17. J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," in *Proc. NAACL-HLT, Minneapolis, MN, USA*, 2019, pp. 4171–4186.
18. O. Xuereb, "Short-term stock market price prediction based on social media and news sentiment," M.S. thesis, University of Malta, 2023.
19. T. Khan, A. Michalas, and A. Akhunzada, "Fake news outbreak 2021: Can we stop the viral spread?," *Journal of Network and Computer Applications*, vol. 190, 103112, 2021.
20. A. Go, R. Bhayani, and L. Huang, "Twitter sentiment classification using distant supervision," *CS224N Project Report*, Stanford University, 2009.
21. A. Assiri, A. Gumaeci, F. Mehmood, T. Abbas, and S. Ullah, "DeBERTa-GRU: Sentiment analysis for large language model," *Computers, Materials & Continua*, vol. 79, no. 3, pp. 4925–4948, 2024.
22. P. He, X. Liu, J. Gao, and W. Chen, "DeBERTa: Decoding-enhanced BERT with disentangled attention," in *Proc. Int. Conf. on Learning Representations (ICLR)*, Vienna, Austria, 2021.
23. Z. Lan, M. Chen, S. Goodman, K. Gimpel, P. Sharma, and R. Soricut, "ALBERT: A lite BERT for self-supervised learning of language representations," in *Proc. Int. Conf. on Learning Representations (ICLR)*, Addis Ababa, Ethiopia, 2020.
24. K. Clark, M. Luong, Q. Le, and C. Manning, "ELECTRA: Pre-training text encoders as discriminators rather than generators," in *Proc. Int. Conf. on Learning Representations (ICLR)*, Addis Ababa, Ethiopia, 2020.
25. A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. Gomez, Ł. Kaiser, and I. Polosukhin, "Attention is all you need," in *Advances in Neural Information Processing Systems (NeurIPS)*, Long Beach, CA, USA, 2017.
26. S. Minaee, N. Kalchbrenner, E. Cambria, N. Nikzad, M. Chenaghlu, and J. Gao, "Deep learning–based text classification: A comprehensive review," *ACM Computing Surveys*, vol. 54, no. 3, pp. 1–40, 2021.
27. T. Brown, B. Mann, N. Ryder, M. Subbiah, J. Kaplan, P. Dhariwal, A. Neelakantan, et al., "Language models are few-shot learners," *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 33, pp. 1877–1901, 2020.
28. T. Otter, J. Medina, and J. Kalita, "A survey of the usages of deep learning for natural language processing," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 32, no. 2, pp. 604–624, 2021.