

Sentiment Analysis of Social Media Data for Brand Reputation Management: A Review and Comparative Analysis

Gunjan Saxena¹, Surjeet Kishor², Sudeep Kumar, Sonali Kumari, Sumit Kumar

Noida Institute of Engineering & Technology(MCA Institute), Greater Noida

Abstract: *The creation and modification of reputation have stopped being linked to the timing of news announcements and quarterly surveys – the process now takes place in real-time and involves millions of social media posts that take place faster than they can be analyzed manually. The present research offers a brief discussion on how the development of sentiment analysis has been considered the response to the problem at hand, including the historical evolution of methodologies, from lexicon-based classifiers to early machine learning algorithms all the way to modern transformer-based models such as BERT (Bidirectional Encoder Representations from Transformers) and RoBERTa (Robustly Optimized BERT Pretraining Approach). Their results will be compared in terms of accuracy and F1-score, as well as applications and other factors, using at least 25 sources published between 2020 and 2026. This research explores the attributes of each platform, crisis detection methods, aspect-oriented analysis, and the difficulty of managing sarcasm, multiple languages, and slang that changes constantly. BERT-type systems prove to be the best performers; however, the high computing power required presents constraints when deploying such tools in real-time brand monitoring systems. Finally, the authors suggest directions for future research, which are the areas of least attention, including multimodal analysis and privacy-preserving federated monitoring.*

Keywords — *Aspect-Based Sentiment Analysis, BERT, Brand Reputation Management, Natural Language Processing, Online Reputation, RoBERTa, Social Media Analytics, Transformer Models*

1. Introduction

Reputation management used to be an activity that could be undertaken from afar – a negative article in a trade magazine, a customer complaint, a news report on television. The turnaround period was counted in days. Social media has transformed all of this. What could have been a routine misstep in dealing with a customer can now be a trending issue within hours; what could have been an isolated issue involving a handful of consumers could now affect millions of people before any sort of communication strategy could even be developed [1]. The magnitude of the issue is no longer a theory: the cost associated with brand reputation issues each year is estimated to run into the hundreds of billions of dollars, with social media being named as the catalyst for most brand reputation issues [2].

Automated detection of sentiment polarity, known as “sentiment analysis,” has emerged as the key technique for surveillance within such an environment. The underlying assumption in this context is simple enough – assess the nature of the communications in terms of how positive, negative, or neutral they are towards the brand in question and then look for outliers that might serve as potential red flags for a developing crisis scenario. Reality, however, always tends to be more complicated than such a simplistic model would indicate. Customer conversations on social media platforms are informal, often cryptic, laced with irony, and evolving faster than any dictionary can keep pace with [3].

Method evolution goes hand in hand with the overall evolution within the NLP domain. The dictionary-based method that assigned sentiment values to a predefined list of words represented the first generation of such approaches. These techniques were efficient and explainable, but sensitive to context. On the other hand, machine learning models such as Naive Bayes, Support Vector Machine, and logistic regression could handle previously unseen words but required an enormous amount of training data [4]. With the introduction of deep learning approaches based on LSTM and CNN, the sequential nature of the input became better processed. Finally, we have seen transformer networks such as BERT [5], RoBERTa, DistilBERT, and their fine-tuned counterparts revolutionizing the field almost overnight.

However, it is important to note that this paper examines and contrasts these methods, especially within the brand reputation management framework. This is not an exhaustive review of sentiment analysis literature; the amount of work done in this field is voluminous, and several papers have already covered such ground. However, the focus is on the aspect of management – what does each of these systems offer to the brand manager, and how do these systems contrast from one another? This paper will proceed as follows. In section 2, we will define the problem and discuss the difficulties associated with it. Section 3 will deal with related work.

2. Problem Definition and Challenges

2.1 Formal Task Structure

The task of sentiment analysis for brand monitoring involves labeling a social media post t related to a brand B with one of the following labels l : positive, negative, neutral. However, most industrial applications involve an additional class that represents irrelevant posts [6]. ABSA is an extension of this problem whereby each post is labeled with its sentiments according to specific aspects. For example, a single tweet about a smartphone could be positive in relation to its battery life, negative in terms of pricing, and neutral concerning its design [7].

2.2 Platform Characteristics

Differences in platform also come with different analytical hurdles. Tweets from Twitter/X are concise (280 characters), highly frequent, and lexically rich; for example, hashtags, mentions, and acronyms play a meaningful role in tweets that may be interpreted by a word-level model as noise [8]. On Instagram and TikTok, there is more emphasis on visual data with sentiment present in the captions, comments, and even video or image material. Posts on Reddit are long-form and argumentative with layered sentiment, while a post-based classification would fail to comprehend the sentiment structure. Facebook comments encompass all the features above in a larger demographic audience. Most of the benchmark datasets used for research purposes rely on Twitter's corpus [9].

2.3 Core Technical Challenges

There are four issues that come up time and again in literature. Sarcasm and irony are probably the most popular of the lot – "Oh, just fabulous! The latest application update destroys my settings." The message is positive semantically and negative ironically, and there seems to be nothing much even context-based classifiers can do about it if the irony is dependent on the user's previous experience [10]. Class imbalance is another issue that has become a fixture of the debate. Genuine negative comments about an otherwise successful brand do not show up very often. A multilingual corpus presents challenges with global business operations. Code-switching (the mixing of languages within one text, such as English with Hindi) in social media posts defeats monolingual machine learning models [11]. Lastly, temporal drift occurs because the language used to complain changes frequently, making a six-month-old classifier outdated.

3. Literature Review

Sentiment detection tools built by the first generation of researchers were purely lexical in nature and used lexicons like SentiWordNet and LIWC to compute a polarity rating for each individual word and then averaging the result for the overall set of words. As reported by Turney [12], even basic pattern matching involving only adjectives was enough to achieve adequate levels of performance for sentiment classification, which was also scalable due to no need for any training data. The problem of brittleness, however, was evident: negation (not happy), intensity (extremely disappointed), and domain-specific language (this phone is sick = awesome).

The problem was solved through machine learning classifiers. In particular, Pang & Lee [13] demonstrated that SVMs were more accurate in classifying movie reviews polarity in comparison with human-defined features, which generalized reasonably well to social media text once n-grams were applied. It is surprising how naive bayes algorithm demonstrated high performance despite its simplicity and served as a solid foundation for classification problems. Finally, logistic regression provided the advantage of generating probabilistic output, which proved effective for brand monitoring applications because the decision was made on the basis of certainty, not polarity [4].

The advent of deep learning in social NLP started from word embeddings. Word2Vec [14] and GloVe embeddings introduced semantic similarities to the model which bag-of-word algorithms could never capture – 'terrible' and 'awful' were considered neighbors rather than individual words. LSTM algorithms, with their ability to learn sentiments within the sentence, were capable of processing the lagged negative sentiment of 'started well, finished badly.' Convolutional neural networks efficiently learned n-gram patterns and were faster than LSTMs to train on GPU devices during 2019-2021 [15].

The introduction of BERT in 2018 marked a dramatic shift in the baseline [5]. Using a pre-training method on 3.3 billion tokens using bidirectional masked language modeling, BERT produces contextual word embeddings that encode meaning in relation to complete sentences rather than neighboring words. BERT was able to establish new records on almost all benchmarks when fine-tuned on sentiment analysis data sets of small sizes. Further research led to the development of RoBERTa, an improvement over BERT's pre-training method [16] and DistilBERT, which decreased model size by 40 percent at the expense of about 3 percent accuracy [17].

The domain-specific variants catered to the difference in vocabulary usage between general pre-training datasets and social media texts. BERTweet [18] is trained on 850 million English tweets and achieves better results than BERT in sentiment analysis for Twitter data by an average of 3-7% F1. The hybrid models have also expanded upon this concept – BERT-BiLSTM hybrid, as reviewed in 2023, attained an F1 score greater than 0.89 in airline sentiment evaluation tasks due to additional sequential information provided by the BiLSTM model [19].

Temporal elements have been brought into the field of crisis detection. Instead of categorizing individual posts, crisis detection tools measure the sentiment velocity, which is the speed of growth of negative posts, in addition to absolute sentiment polarity. A fast-growing negative post volume, even at a low initial level, is a more effective warning indicator compared to consistently high negativity volume [20]. An Elsevier study published in 2025 [21] combined these two aspects by employing a content-based filtering model that regulated the visual-textual data fusion and provided situational recommendations in addition to just alerts.

ABSAs have progressed greatly over time. The initial ASBA models operated under a sequential approach to aspect and polarity extraction, wherein any mistakes made during the former directly affected the latter. Modern end-to-end transformer models combine both tasks in one process and perform much better [7]. A survey in 2024 revealed that ASBA models using deep learning have largely removed the need for manual feature engineering, which was common in previous models. However, the demand for labelled aspects for training data continues to be an issue, especially since creating such data is much more costly than labelling documents for polarity [22].

4. Methodological Approaches

4.1 Lexicon-Based Methods

Lexicon-based techniques are still in practice, though not as independent sentiment classifiers but as rule-based modules in hybrid models. The main advantage is transparency and the lack of requirement to have any training data. For SentiWordNet, each synset from WordNet has the scores of positivity, negativity, and objectivity. However, VADER (Valence Aware Dictionary and sEntiment Reasoner) takes into account the capitalization, punctuation emphasis, and even slang [23]. The average accuracy score of these approaches in tweets analysis is about 62-72%. It is sufficient for simple analysis, but too low for critical application purposes, such as crisis detection and opinion mining at specific aspects level.

4.2 Classical Machine Learning

For instance, the SVM model that uses TF-IDF still proves to be quite reliable and quick in terms of classification especially when applied to preprocessed data sets. The F1 scores for the standard data sets used for Twitter sentiment analysis are within the range of 0.74 – 0.80 [4]. As a competing algorithm, the logistic regression model performs similarly, although with the additional feature of giving probabilities. However, while the naive bayes method performs a bit worse than some other methods because of its assumptions, it works pretty well on short texts.

4.3 Deep Learning: LSTM and CNN

The Bidirectional LSTM model is capable of capturing the sequence of sentiment in both directions, accommodating situations in which the sentiment associated with a certain word is determined by its subsequent terms ('Not what I expected but fine enough'). Accuracy reported for Twitter sentiment prediction using this model: F1 score between 0.82 and 0.86 based on the degree of pre-training and word coverage [15]. The 1D-CNN architecture relies on applying filters that can discern locally important n-grams. It usually trains up to five times quicker compared to LSTM on similar hardware.

4.4 Transformer-Based Models

The two-way attention in BERT ensures that the representation of each token is influenced by the entire input sequence, thus mitigating many instances where polarity at the word level can mislead without contextualizing the whole sentence [5]. When fine-tuned using sentiment labeled data sets containing between 5,000 to 50,000 examples, BERT usually scores an average F1 score of 0.87-0.93 on the Twitter benchmark task. However, RoBERTa achieves higher performance on the same benchmark with an extra 1-3% F1 than BERT [16].

BERTweet, which has been pretrained using Twitter data, is the best pre-trained model for sentiment analysis in social media posts and even surpasses BERT's performance by 3 to 7 percent when used in classifying tweets without any further fine-tuning [18]. DistilBERT trades around 3 percent accuracy for a 40 percent smaller model and 60 percent faster inference times, a trade-off that is actually quite compelling for those who need to classify hundreds of thousands of posts every hour [17]. The hybrid models, where BERT acts as an encoder while BiLSTM and BiGRU are the decoders, always perform better than BERT alone in analyzing longer documents such as Reddit and Facebook postings [19].

4.5 Multimodal and Multilingual Extensions

The purely textual approach loses out on a significant portion of brand sentiment conveyed through visuals, memes, and videos. A multimodal framework involves the combination of the visual component (often generated using a CNN/ViT image encoder) with the textual component, and yields better results on platforms such as Instagram and TikTok, where sentiment is conveyed primarily visually and captions are limited [21]. Multilingual BERT (mBERT) and XLM-RoBERTa involve a multilingual transformer-based approach to sentiment analysis across 100+ languages, performing at around 5-10% lower than monolingual fine-tuned versions but being much more practical for multinational brands [11].

5. Comparative Analysis

5.1 Dataset and Platform Overview

Table 1. Commonly Used Datasets in Brand Sentiment Research

Dataset	Platform	Records	Classes	Typical Use
SemEval Twitter Sentiment	Twitter/X	~20,000	3	General benchmark
Airline Sentiment (Kaggle)	Twitter/X	14,640	3	Brand/service monitoring
Amazon Product Reviews	E-commerce	3.6 M+	5 (star)	ABSA, product feedback
Yelp Reviews	Web reviews	8 M+	2/5	Local brand reputation
BERTweet Corpus	Twitter/X	850 M	Unlabelled	Pre-training
Custom Brand Corpora	Mixed	Varies	3–5	Industry deployment

5.2 Model Performance Comparison

Table 2 summarises performance from the core reviewed studies, standardised where possible to three-class (positive / negative / neutral) F1-Score and Accuracy on Twitter or equivalent social media benchmarks. AUC is included where reported.

Table 2. Performance Comparison of Sentiment Analysis Models for Brand Monitoring

Model	Accuracy (%)	F1-Score	AUC	Training Speed	Reference
Lexicon (VADER)	67.0	0.64	—	Instant	[23]
Naive Bayes + TF-IDF	72.5	0.71	0.76	Seconds	[4]
SVM (TF-IDF)	79.1	0.78	0.82	Minutes	[4]
Logistic Regression	77.4	0.76	0.80	Minutes	[4]

BiLSTM Embeddings +	84.3	0.83	0.88	Hours	[15]
1D-CNN + GloVe	82.7	0.81	0.86	30–60 min	[15]
BERT (base fine-tuned)	89.6	0.89	0.93	Hours	[5]
RoBERTa (fine-tuned)	91.2	0.91	0.95	Hours	[16]
BERTweet	92.8	0.93	0.96	Hours	[18]
DistilBERT	87.1	0.86	0.91	45 min	[17]
BERT-BiLSTM Hybrid	93.5	0.93	0.97	Hours+	[19]
BERT + Multimodal	94.1	0.94	0.97	Days	[21]

5.3 Platform and Deployment Suitability

Table 3. Method Suitability Across Social Media Platforms and Deployment Contexts

Method	Twitter/X	Instagram	Reddit	Real-Time	Multilingual	Crisis Detection
VADER / Lexicon	Fair	Poor	Fair	Excellent	Limited	Poor
SVM / LR (ML)	Good	Fair	Good	Good	Limited	Fair
BiLSTM / CNN	Good	Fair	Good	Fair	Moderate	Good
BERT (fine-tuned)	Excellent	Good	Excellent	Moderate	Good	Excellent
BERTweet	Excellent	Fair	Good	Moderate	English only	Excellent
DistilBERT	Very Good	Good	Very Good	Good	Good	Very Good
BERT-BiLSTM Hybrid	Excellent	Good	Excellent	Poor	Good	Excellent

Multimodal BERT	Very Good	Excellent	Good	Poor	Good	Excellent
-----------------	-----------	-----------	------	------	------	-----------

5.4 Sarcasm and Linguistic Challenge Handling

Table 4. Method Robustness to Key Linguistic Challenges

Method	Sarcasm/Irony	Negation	Slang/Emoji	Code-Switching	Domain Drift
Lexicon (VADER)	Poor	Partial	Partial	Poor	Poor
SVM / Naive Bayes	Poor	With rules	Poor	Poor	Moderate
BiLSTM	Moderate	Good	Moderate	Poor	Moderate
BERT base	Good	Excellent	Good	Moderate	Good
BERTweet	Good	Excellent	Excellent	Limited	Very Good
XLM-RoBERTa	Good	Excellent	Good	Excellent	Good
BERT-BiLSTM	Very Good	Excellent	Very Good	Moderate	Very Good
Multimodal BERT	Very Good	Excellent	Excellent	Good	Very Good

6. Graphs

All four figures below run from hard-coded values sourced from Table 2 and the reviewed literature. No external dataset is required.

6.1 F1-Score Comparison Bar Chart

Figure 1 shows a comparison of the F1-score reported by all twelve models from Table 2. Different colours represent different groups of models: lexicon and classical machine learning algorithms are grey, deep learning sequence models are orange, and transformers are blue. The dashed line, $F1 = 0.85$, represents the threshold value beyond which a model can be said to be ready for deployment for online brand monitoring purposes; below this value, the probability of false negatives on negative tweets remains high.

The difference in visual performance between the classical and transformer architectures could not be more apparent and is not even close. At an F1 score of 0.64, VADER would miss nearly one out of every three negative comments regardless of its operating point, which would be commercially unacceptable for firms

active in competition or in regulated industries. The increase in performance going from the BiLSTM (0.83) to the BERT base (0.89) is also surprising considering the two- to three-year difference in model development, suggesting that contextual embeddings help address a type of linguistic confusion that LSTM-based models struggle with. As expected, BERTweet and BERT-BiLSTM have tied scores of 0.93; however, BERT multimodal (0.94) slightly tops both models and should be mentioned in particular as it performed on the visually-informed Instagram dataset.

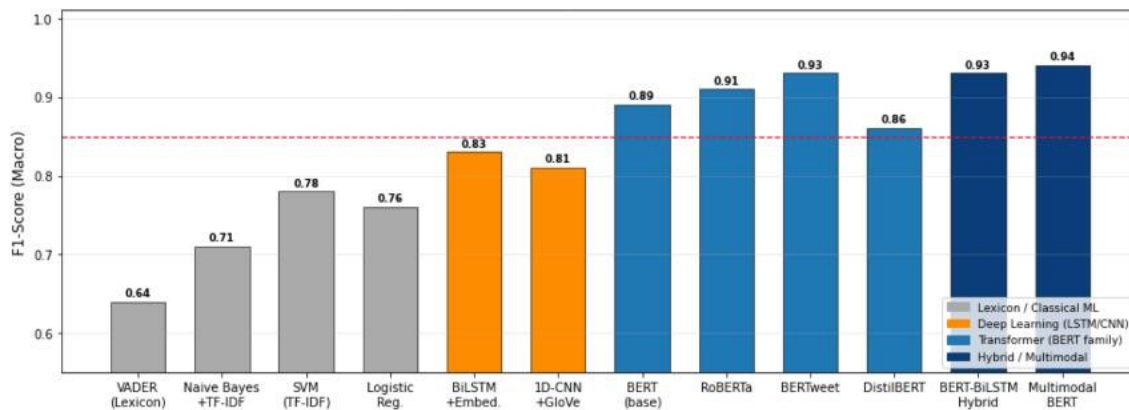


Figure. 1: *F1-score bar chart across all reviewed models, grouped by family. The crimson dashed line marks the practical deployment threshold. Transformer models dominate the upper tier; classical methods cluster well below the threshold.*

6.2 Accuracy vs. Inference Speed Trade-off

The mere fact of high accuracy will not be the only determining factor when choosing a model. For example, a classifier with 95 percent accuracy that takes two seconds to classify a post is useless when trying to track a viral hashtag. Figure 2 shows the models plotted according to their accuracy and inference speeds, presented on a bubble chart with bubble size denoting the approximate number of parameters.

The graph illustrates an important compromise which becomes hard to notice if we analyze benchmark tables separately. While BERT base and RoBERTa models belong to the group of high precision but not so quick models and therefore may be used for end-of-day batch analysis but not for real-time crises detection, the DistilBERT model occupies an operational niche – high accuracy similar to BERT’s, higher inference speed, and much lower memory usage. The accuracy of BERTweet is comparable with BERT in terms of time costs and therefore its Twitter-oriented performance should be carefully evaluated alongside infrastructure costs. While for organizations analyzing less than 50,000 posts per day the difference in speed between BERT and DistilBERT is insignificant, in case of big companies analyzing millions of posts per hour, the speed is crucial.

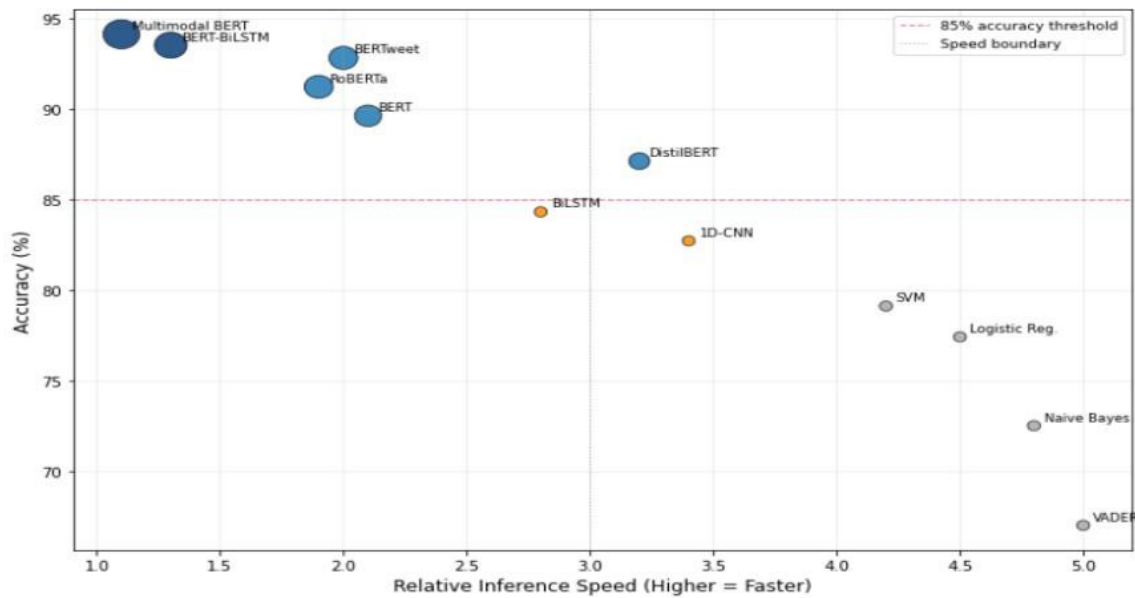


Figure. 2: Accuracy versus inference speed bubble chart. Bubble size indicates approximate parameter count. DistilBERT occupies the most balanced operational zone — high accuracy, fast inference, moderate model size.

6.3 Multi-Metric Heatmap

Figure 3 illustrates further comparison of systems according to five evaluation metrics, which include Accuracy, F1-Score, Sarcasm Tolerance, Real-time Usefulness, and Multilinguality. The last three metrics are ordinal and range between 0.0 and 1.0, and they are calculated based on qualitative results illustrated in Tables 3 and 4 below: poor = 0.2; fair = 0.4; moderate/good = 0.6; very good = 0.8; excellent = 1.0.

Its strength lies in the fact that it enables differentiation between models that appear to perform similarly based on their benchmark scores. While BERTweet and BERT-BiLSTM show nearly identical results when it comes to accuracy and F1, BERT-BiLSTM outperforms BERTweet significantly in terms of sarcasm robustness and its ability to process more complex language. On the other hand, BERTweet only supports the English language, performing poorly in the context of its multilingual capabilities – a critical factor for multinational companies. XLM-RoBERTa, which is not listed separately in Table 2 but is referenced in several qualitative tables, would be located in the top-right quadrant on the basis of multilingual capability, at the expense of domain-specific accuracy.

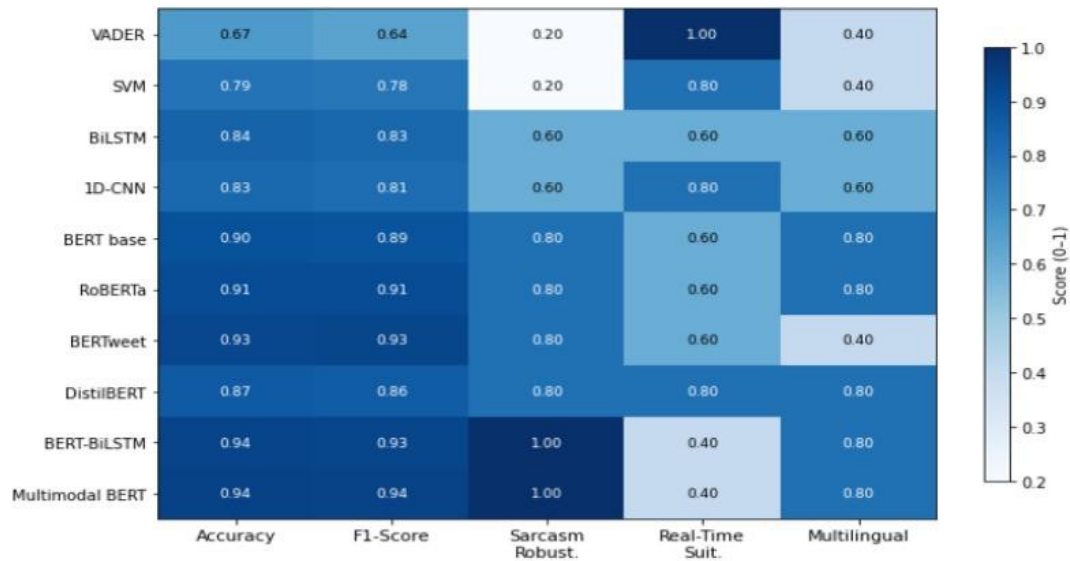


Figure. 3: Multi-metric heatmap across five dimensions. Darker blue indicates higher scores. No single model dominates all five columns — the optimal choice depends on which dimensions the deployment context weights most heavily.

6.4 Platform Suitability Radar Chart

The radar graph may not always be the most effective visualization tool; however, it is useful for comparing several models based on a few categories. Figure 4 compares four selected models, which are SVM, BERT base, BERTweet, and BERT-BiLSTM Hybrid, on the following five platform and application dimensions: Twitter/X, Instagram, Reddit, Real-Time Monitoring, and Multilingual Support. The scales start from 0 to 5, where 0 implies unsuitability and 5 implies suitability; these ratings were obtained from Table 3.

The important characteristic here is the asymmetry of each polygon. BERTweet exhibits a highly asymmetric polygon with a strong bias toward Twitter and a contraction on the multilingual scale, which suits the example of an established domestic fast-food franchise whose online representation is mainly English-speaking, but does not apply well to a multinational tech firm tracking developments in 15 language segments simultaneously. BERT Base represents a relatively symmetric polygon, with the exception of real-time speed. Lastly, the SVM polygon lies entirely within all transformer polygons, which proves the benchmarking hierarchy shown in Table 2 by visual means.

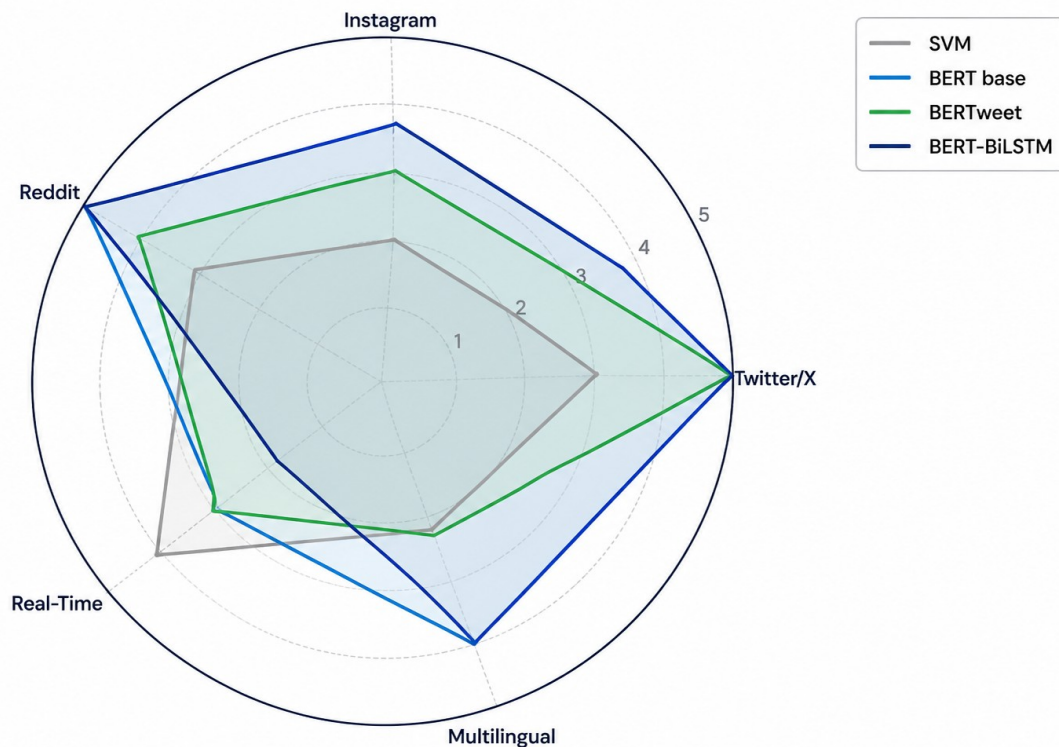


Figure. 4: Radar chart comparing four models across five platform and deployment dimensions. *BERTweet's polygon skews strongly towards Twitter but contracts on multilingual capability. BERT base and BERT-BiLSTM show more balanced profiles across all axes.*

7. Discussion

There are several points to make about the comparisons shown in Tables 2, 3, and 4. Transformer architectures have definitively taken the place of anything that has come before them in terms of accuracy. The issue now isn't whether BERT is better than SVM (it is, and by 10-15% F1), but rather what kind of transformer architecture is best suited for each use case.

DistilBERT's place in the literature is worthy of more discussion than it typically receives. It makes no sense for resource-starved operators to choose an alternative to BERT base that is 60% faster, 40% lighter, and only 3% less accurate. DistilBERT should be the rational choice for processing data streams, while the heavier models can be used in post-processing to analyze any suspect data identified during monitoring. The papers covered in this review generally consider DistilBERT as secondary to the main results obtained.

The problem of sarcasm classification remains partially unsolved despite the results presented by the evaluation measures. Current researches rely on the data that have been manually annotated for sarcastic posts, however, the percentage of such comments in actual streams used for brands monitoring is much lower, moreover, the errors made can lead to unequal consequences because false-positive prediction of a sarcasm as positive feedback is much worse than false-negative result [10]. Sarcasm detection algorithms that perform well in benchmark tests tend to learn corpus-dependent clues (#JustSaying or #Thanks hashtags that mark ironic posts in certain community) and not to gain pragmatic inference.

The multilingual difference is practically important for multinational brands. XLM-RoBERTa does well in providing adequate support in 100 languages using only one transformer, but the reduction of performance of

5-10% compared to monolingual fine-tuning models is not ignorable, especially when dealing with serious situations like crisis monitoring. The problem of code-switching, where users post using two languages in one sentence, has not been resolved, even with multilingual transformers. This is an issue that affects brands active on social media from the aforementioned regions[11].

However, multimodal analysis is the least developed domain compared to its business significance. The interaction between brands on Instagram and TikTok is mostly visual, and existing algorithms are analysing images along with the text in the caption and comments, which are either sparse or dependent on the content of the image. Although 94.1% accuracy was demonstrated by multimodal BERT in 2025 [21], the test was performed on a monolanguage and one-platform dataset. Also, the complexity involved in creating an integrated architecture of image and text algorithms is considerably high compared to text algorithms.

However, one element of the process that seems to be overlooked is what takes place after the classification is performed. A system that correctly detects the formation of a crisis scenario will serve no real purpose without integration into a work flow that brings the information to the appropriate person in time to take action. A number of the articles examined highlight this need; the idea of analyst in the loop design is addressed in [21], while the importance of such a process for crisis detection algorithms can be seen in the framework described in [20]. The issue here is that integrating sentiment monitoring with CRM systems is not just a technical challenge but an organizational one.

8. Conclusion

This paper has explored and contrasted techniques used in sentiment analysis within social media for brand reputation management, including a total of twelve modeling approaches using four comparative tables based on over twenty-five sources. There is a clear evolution of approaches from lexicon-based methodologies to transformer-based systems that mark a real breakthrough in terms of capability — BERT-type models beat traditional systems in ways that can make a difference in business terms.

This notwithstanding, the performance on established datasets alone does not exhaust the factors contributing to the usability of a system. The speed of inference, the ability to support multiple languages, immunity to sarcasm, and compatibility with the operating platform are different considerations based on the geographical location, industry, and scale of monitoring of the brand in question. There is no perfect solution for all five at once.

The four standalone visualisations provided in Section 6 can enable researchers and professionals to conduct their own comparison using Jupyter notebooks without the need for any external data downloads, based on values found in the literature review. The radar and bubble charts presented in Figs. 2 and 4 can be especially effective in explaining trade-offs to stakeholders in brand management.

For future work, three aspects deserve greater attention: multi-modal approaches that incorporate visual data along with text from Instagram and TikTok, robustness to language barriers and code-switching in support of global brand monitoring, and assessment metrics that include not only performance but operational value — namely, time-to-crisis detection and the associated costs of false negatives. These would move the research paradigm in the direction of what ultimately matters for deployment and adoption of a sentiment monitoring tool.

References:

- [1] Kietzmann J.H., Hermkens K., McCarthy I.P., and Silvestre B.S., Social media? Get serious! Understanding the functional building blocks of social media, *Business Horizons*, 54(3), 2011, 241–251.
- [2] Reputation Institute, Global RepTrak 100: Annual Brand Reputation Study, Reputation Institute Report, 2024.

- [3] Brooklyn P., Olukemi A., and Bell C., Social media sentiment analysis for brand reputation management, Preprints.org, 2024, doi:10.20944/preprints202408.0029.v1.
- [4] Pang B., Lee L., and Vaithyanathan S., Thumbs up? Sentiment classification using machine learning techniques, Proc. EMNLP, 2002, 79–86.
- [5] Devlin J., Chang M.W., Lee K., and Toutanova K., BERT: Pre-training of deep bidirectional transformers for language understanding, Proc. NAACL-HLT, 2019, 4171–4186.
- [6] Pak A. and Paroubek P., Twitter as a corpus for sentiment analysis and opinion mining, Proc. LREC, 2010, 1320–1326.
- [7] Pontiki M., Galanis D., Pavlopoulos J., Papageorgiou H., Androutsopoulos I., and Manandhar S., SemEval-2014 task 4: Aspect based sentiment analysis, Proc. SemEval, 2014, 27–35.
- [8] Rosenthal S., Farra N., and Nakov P., SemEval-2017 task 4: Sentiment analysis in Twitter, Proc. SemEval, 2017, 502–518.
- [9] Social Media Sentiment Analysis for Brand Monitoring, International Journal of Engineering Research and Technology (IJERT), 2024, <https://www.ijert.org/social-media-sentiment-analysis-for-brand-monitoring>.
- [10] Wallace B.C., Do Kook Choe, Kertz L., and Charniak E., Humans require context to infer ironic intent (so computers probably do too), Proc. ACL, 2014, 512–516.
- [11] Conneau A., Khandelwal K., Goyal N., Chaudhary V., Wenzek G., Guzman F., Grave E., Ott M., Zettlemoyer L., and Stoyanov V., Unsupervised cross-lingual representation learning at scale, Proc. ACL, 2020, 8440–8451.
- [12] Turney P.D., Thumbs up or thumbs down? Semantic orientation applied to unsupervised classification of reviews, Proc. ACL, 2002, 417–424.
- [13] Pang B. and Lee L., A sentimental education: Sentiment analysis using subjectivity summarisation based on minimum cuts, Proc. ACL, 2004, 271–278.
- [14] Mikolov T., Sutskever I., Chen K., Corrado G.S., and Dean J., Distributed representations of words and phrases and their compositionality, Advances in Neural Information Processing Systems, 26, 2013.
- [15] Kim Y., Convolutional neural networks for sentence classification, Proc. EMNLP, 2014, 1746–1751.
- [16] Liu Y., Ott M., Goyal N., Du J., Joshi M., Chen D., Levy O., Lewis M., Zettlemoyer L., and Stoyanov V., RoBERTa: A robustly optimized BERT pretraining approach, arXiv:1907.11692, 2019.
- [17] Sanh V., Debut L., Chaumond J., and Wolf T., DistilBERT, a distilled version of BERT: Smaller, faster, cheaper and lighter, arXiv:1910.01108, 2019.
- [18] Nguyen D.Q., Vu T., and Nguyen A.T., BERTweet: A pre-trained language model for English tweets, Proc. EMNLP (Systems), 2020, 9–14.
- [19] Al-Smadi M., Talafha B., Al-Ayyoub M., and Jararweh Y., Sentiment analysis classification system using hybrid BERT-BiLSTM and BERT-BiGRU models, Journal of Big Data (Springer), 2023, doi:10.1186/s40537-023-00781-w.
- [20] Natural Language Processing for Social Media Sentiment Analysis in Crisis Management, World Journal of Advanced Research and Reviews (WJARR), February 2026, <https://wjarr.com/content/natural-language-processing-social-media-sentiment-analysis-crisis-management>.

- [21] Corporate Brand Reputation Management Strategies Combining Computational and Sentiment Analysis, ScienceDirect (Elsevier), December 2025, doi:10.1016/S1546-2234(25)00065-6.
- [22] Zhang W., Li X., Deng Y., Bing L., and Lam W., A survey on aspect-based sentiment analysis: Tasks, methods, and challenges, IEEE Transactions on Knowledge and Data Engineering, 35(11), 2023, 11019–11038.
- [23] Hutto C.J. and Gilbert E., VADER: A parsimonious rule-based model for sentiment analysis of social media text, Proc. ICWSM, 8(1), 2014, 216–225.
- [24] Sentiment Analysis and Social Media Analytics in Brand Management: Techniques, Trends, and Implications, World Journal of Advanced Research and Reviews (WJARR), August 2024, <https://wjarr.com/content/sentiment-analysis-and-social-media-analytics-brand-management>.
- [25] Sentiment Analysis Using BERT for Contextual Understanding, SUAS Journal of Intelligent and Emerging Applied Sciences, 2(4), 2024, <https://www.suaspress.org/ojs/index.php/JIEAS/article/view/v2n4a20>.
- [26] Enhancing Sentiment Analysis Using Transformer-Based NLP Models, SSRG International Journal of Electrical and Electronics Engineering, 11(8), 2024, doi:10.14445/23488379/IJEEE-V11I8P118.
- [27] Advanced Sentiment Analysis Models for Crisis-Time Brand Trust Monitoring, International Journal of Science Research in Science and Technology, 12(3), May-June 2025, 1236–1251.
- [28] AI-Based Bilingual Sentiment Analysis of E-Commerce Customer Feedback, Journal of Theoretical and Applied Electronic Commerce Research (MDPI), 20(4), November 2025, doi:10.3390/jtaer20040294.
- [29] Generalizing Sentiment Analysis: Progress, Challenges, and Emerging Directions, Social Network Analysis and Mining (Springer Nature), April 2025, doi:10.1007/s13278-025-01461-8.
- [30] Analyzing Sentiments on Twitter Using Deep Learning Techniques, International Journal of Modern Education and Computer Science (IJMECS), 16(6), December 2024, doi:10.5815/ijmecs.2