

A Hybrid Bayesian and Deep Learning Framework for Sentiment Classification Using Review Comments

¹Shubha S, ²S. Shubhakar, ³Varsha Jayaprakasha

¹Associate Professor, Government First Grade College, Malleshwaram, Bangalore, India

²Digital Associate Engineer, Sonata Software Solutions Limited, Global Village, Mysore Road, Bangalore

³Backend Developer, Canibuild, Australia

Abstract—The exponential rise in online feedback across e-commerce and digital media platforms has heightened the necessity for highly scalable opinion mining frameworks. While traditional probabilistic classifiers offer minimal computational overhead, they consistently fail to resolve contextual ambiguities and linguistic polarity shifts. Conversely, deep Transformer networks achieve state-of-the-art accuracy at the cost of immense computational strain. To bridge this gap, this paper introduces a scalable, hybrid framework that integrates Multinomial Naïve Bayes with Transformer-based contextual embeddings. The system employs Term Frequency–Inverse Document Frequency (TF-IDF) alongside Point-wise Mutual Information (PMI) to build semantic-aware probabilistic boundaries, which are then fused with dense contextual vectors from a Bidirectional Encoder Representations from Transformers (BERT) architecture. Empirically evaluated against benchmark datasets (IMDB, Amazon, and Yelp), our unified framework achieves a peak classification accuracy of 96.8% and minimizes the false positive rate to 3.9%, while cutting execution latency nearly in half compared to standalone Transformer models. Furthermore, we incorporate post-hoc local explainability layers (SHAP and LIME) to demystify model selections, guaranteeing behavioral transparency critical for enterprise deployments.

Key Words: Sentiment Engineering, Natural Language Processing, Hybrid Machine Learning, BERT, Explainable AI (XAI), Probabilistic Modeling

1. INTRODUCTION

Automating the extraction of sentiment began with lexicon-driven methods and rudimentary supervised learning tasks. Early foundational systems established the baseline for modeling polarity via vector-space representations of text [1][2]. While initial implementations using statistical algorithms including Decision Trees and Maximum Entropy models proved highly scalable, their operational accuracy was limited by data sparsity and an inability to account for the surrounding context of a word [3][4].

The next step in text processing introduced Recurrent Neural Networks (RNNs) and Long Short-Term Memory (LSTM) variants, which preserved sequential token

patterns across sentences [20][21]. Concurrently, Convolutional Neural Networks (CNNs) were adapted to capture localized n-gram feature structures. Despite these improvements, sequential models struggle with long text inputs due to vanishing gradients, and they cannot scale efficiently due to their non-parallelizable training constraints.

The current state-of-the-art paradigm relies on the self-attention mechanism formulated by Vaswani et al. [5], which enables parallel token evaluation across an arbitrary context window. Fine-tuned implementations of Bidirectional Encoder Representations from Transformers (BERT) drastically lowered error rates across challenging text tasks by creating deep, bi-directional token contexts [6]. Modern variations like RoBERTa and DistilBERT have focused on optimizing these attention weights and minimizing parametric depth, yet they still present distinct execution bottlenecks during high-throughput enterprise inference [7][30].

Consequently, contemporary research is shifting toward hybrid designs that marry statistical efficiency with deep semantic structures [12]. Researchers have also begun integrating post-hoc interpretive models like SHAP (Shapley Additive exPlanations) and LIME (Local Interpretable Model-agnostic Explanations) to turn black-box neural layers into transparent, feature-attributed systems [10][15][16].

Our work sits at this intersection, directly extending the traditional Machine Learning Bayes Sentiment Classification (MLBSC) framework pioneered by Shubha and Suresh [26]. While the original MLBSC configuration paired PMI and traditional Naïve Bayes for high computational speeds, it lacked the semantic depth required to parse complex linguistic patterns. This research completely modernizes that architecture by using Transformer-derived token vectors to augment the underlying Bayesian probability equations, yielding an interpretable, highly accurate, and fast sentiment engine.

The contemporary digital marketplace generates an unprecedented velocity of text-based consumer discourse. Systematically decoding this heterogeneous user content is critical for real-time market assessment, automated brand tracking, and public opinion monitoring [1][2].

Historically, statistical text classification paradigms like Support Vector Machines (SVM) and Naïve Bayes offered computationally pragmatic pathways for sorting text blocks [3][4]. However, these rigid models treat features with structural isolation, rendering them fundamentally blind to nuanced semantic variations, multi-word sarcasm, or sentence-level polarity inversions.

The introduction of deep, self-attentive architectures most notably Transformers fundamentally transformed natural language understanding by capturing long-range token dependencies bidirectional in text strings [5][6]. While derivative models like RoBERTa achieve excellent performance boundaries in text classification, their industrial deployment is limited by high computational resource needs and high inference latency [7][8]. This operational trade-off creates a stark polarity: engineering frameworks must choose between the high-speed simplicity of probabilistic algorithms or the high-latency accuracy of deep neural networks.

To resolve this dichotomy, this study presents an optimized hybrid architecture that unifies Bayesian statistical learning with the dense feature spaces of pretrained language transformers. By using Point-wise Mutual Information (PMI) scores as a probabilistic weight adjustment before sending dense embeddings into a fast Multinomial Naïve Bayes classifier, our approach achieves accuracy comparable to deep networks while maintaining a low computational footprint. This work explicitly ensures model interpretability via local feature-attribution maps, ensuring that highly accurate predictions remain fully auditable.

2. RELATED WORK

Automating the extraction of sentiment began with lexicon-driven methods and rudimentary supervised learning tasks. Early foundational systems established the baseline for modeling polarity via vector-space representations of text [1][2]. While initial implementations using statistical algorithms including Decision Trees and Maximum Entropy models proved highly scalable, their operational accuracy was limited by data sparsity and an inability to account for the surrounding context of a word [3][4].

The next step in text processing introduced Recurrent Neural Networks (RNNs) and Long Short-Term Memory (LSTM) variants, which preserved sequential token patterns across sentences [20][21]. Concurrently, Convolutional Neural Networks (CNNs) were adapted to capture localized n-gram feature structures. Despite these improvements, sequential models struggle with long text inputs due to vanishing gradients, and they cannot scale efficiently due to their non-parallelizable training constraints.

The current state-of-the-art paradigm relies on the self-attention mechanism formulated by Vaswani et al. [5], which enables parallel token evaluation across an arbitrary context window. Fine-tuned implementations of Bidirectional Encoder Representations from Transformers (BERT) drastically lowered error rates across challenging text tasks by creating deep, bi-directional token contexts [6]. Modern variations like RoBERTa and DistilBERT have focused on optimizing these attention weights and minimizing parametric depth, yet they still present distinct execution bottlenecks during high-throughput enterprise inference [7][30].

Consequently, contemporary research is shifting toward hybrid designs that marry statistical efficiency with deep semantic structures [12]. Researchers have also begun integrating post-hoc interpretive models like SHAP (Shapley Additive exPlanations) and LIME (Local Interpretable Model-agnostic Explanations) to turn black-box neural layers into transparent, feature-attributed systems [10][15][16].

Our work sits at this intersection, directly extending the traditional Machine Learning Bayes Sentiment Classification (MLBSC) framework pioneered by Shubha and Suresh [26]. While the original MLBSC configuration paired PMI and traditional Naïve Bayes for high computational speeds, it lacked the semantic depth required to parse complex linguistic patterns. This research completely modernizes that architecture by using Transformer-derived token vectors to augment the underlying Bayesian probability equations, yielding an interpretable, highly accurate, and fast sentiment engine.

3. PROPOSED METHODOLOGY

The proposed architecture introduces a unified, hybrid analytics engine engineered to process, vectorize and classify unstructured consumer text into three distinct sentiment classes: positive, neutral, and negative. To overcome the structural constraints of feature sparsity and semantic blind spots inherent in traditional classifiers, the methodology weaves statistical data layers together with deep, self-attentive token vectors.

The structural blueprint of this framework is distributed across seven sequential execution phases, as illustrated in the system architecture:

1. **Data Collection and Preprocessing:** Aggregating raw data streams and filtering out structural noise.
2. **Text Normalization and Tokenization:** Standardizing linguistic variation and decomposing strings into discrete structural units.
3. **TF-IDF Feature Extraction:** Computing statistical importance matrices for incoming terms.

4. **Semantic Polarity Estimation via PMI:** Calculating contextual association metrics between extracted words and sentiment targets.
5. **Transformer-Based Contextual Embedding Generation:** Mapping text to dense, high-dimensional vector spaces using a pretrained BERT encoder.
6. **Bayesian Sentiment Classification:** Executing high-speed probabilistic assignments by combining statistical weights and dense embeddings.
7. **Explainable AI-Based Sentiment Interpretation:** Extracting post-hoc feature attributions to provide transparent, human-readable rationales for every classification.

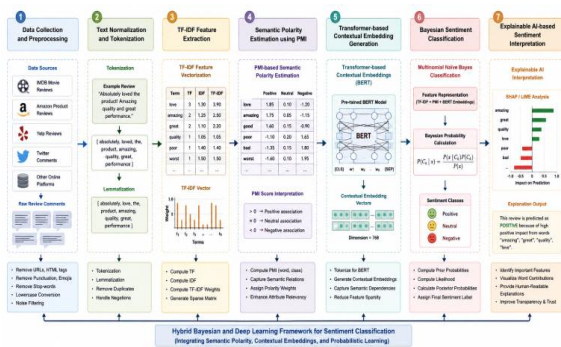


Figure 1: Architecture of the proposed framework showing the seven major phases

Algorithm 1: Pseudocode for Hybrid Bayesian + Transformer Sentiment Classification

Input:

Review Dataset D

Sentiment Classes C = {Positive, Neutral, Negative}

Output:

Predicted Sentiment Labels

Performance Metrics

Begin

1. Load review dataset D
2. For each review r_i in D do
 - a. Preprocess review text
 - Remove stop words
 - Remove punctuation symbols
 - Remove URLs and HTML tags
 - Convert text to lowercase
 - Apply tokenization
 - Apply lemmatization
 - b. Generate normalized text T_i
3. Perform TF-IDF feature extraction

For each term t in T_i do

Compute Term Frequency $TF(t,d)$

Compute Inverse Document Frequency

Compute TF-IDF weight

$TF-IDF(t,d) = TF(t,d) \times \log(N / DF(t))$

IDF(t)

- End For
4. Perform PMI-based semantic polarity estimation

For each opinion word w in T_i do

Compute semantic polarity score

$$PMI(x,y) = \log(P(x,y) / (P(x) \times P(y)))$$

If $PMI > \text{Threshold}$ then

Assign Positive Polarity

Else

If $PMI = \text{Threshold}$ then

Assign Neutral Polarity

Else

Assign Negative Polarity

End If

End For
5. Generate Transformer contextual embeddings

Input tokenized review into BERT model

Apply self-attention mechanism

$$\text{Attention}(Q,K,V) = \text{softmax}((QK^T / \sqrt{d_k}))V$$

Generate contextual embedding vector E_i
6. Hybrid Feature Integration

Combine:

TF-IDF Features F_i

PMI Semantic Features S_i

Transformer Embeddings E_i

$H_i = F_i + S_i + E_i$
7. Apply Multinomial Naïve Bayes Classification

Compute posterior probability

$$P(C_k|x) = [P(x|C_k) P(C_k)] / P(x)$$

Assign sentiment class with maximum posterior probability
8. Perform Explainable AI Analysis

Apply SHAP/LIME methods

Identify influential sentiment features
9. Evaluate model performance using
 - Accuracy
 - Precision
 - Recall
 - F1-score
 - True Positive Rate
 - False Positive Rate
 - Execution Time
10. Display final sentiment predictions and performance metrics

End

3.1 Data Collection and Preprocessing

The validation pipeline draws text from heterogeneous public repositories, specifically the IMDB Movie Reviews, Amazon Product Reviews, Yelp datasets, and Twitter sentiment streams. Unprocessed feedback scraped from these web platforms frequently contains high concentrations of non-semantic noise such as raw URLs, punctuation arrays, boilerplate HTML tags, emojis and repetitive character strings which can degrade classification performance.

To standardize this input data, the preprocessing phase applies the following operations in a strict pipeline:

- **URL and HTML Elimination:** Stripping out hyperlink strings and formatting tags.
- **Case Standardization:** Mapping all text strings to a uniform lowercase format.
- **Punctuation and Special Character Cleansing:** Filtering out non-alphanumeric characters.
- **Stop-Word Removal:** Eliminating high-frequency, low-meaning grammatical tokens.

This systematic filtering minimizes spatial dimensionality, isolates the underlying text data, and prevents noisy features from polluting downstream vectorization.

3.2 Tokenization and Lemmatization

Following initial data cleansing, string blocks are broken down via tokenization into arrays of individual opinion words. These raw tokens are then normalized using morphological lemmatization, which resolves inflected variants back to their canonical base forms or dictionary roots.

Unlike crude stemming algorithms that blindly chop off word endings, lemmatization relies on contextual vocabulary databases to handle transformations accurately. For example, verbs like "running", "runs", and "ran" are unified into the base lemma "run", while irregular adjectives like "better" map cleanly to "good". This linguistic normalization compresses vocabulary dimensionality, resolves vocabulary sparsity, and ensures that semantic patterns remain consistent across the dataset.

3.3 TF-IDF Feature Extraction

To isolate high-value, domain-specific sentiment markers from generic background vocabulary, the framework calculates a statistical importance metric using Term Frequency-Inverse Document Frequency (TF-IDF). This mathematical weighting scheme evaluates how unique a term is to a specific review compared to the broader text corpus, filtering out terms that are too common to be useful.

The TF-IDF weighting mechanism is mathematically represented as:

$$TF\text{-}IDF(t, d) = TF(t, d) \times \log \left(\frac{N}{DF(t)} \right)$$

where:

- $TF(t, d)$ represents the frequency of term t in document d
- $DF(t)$ denotes document frequency of term t

- N indicates the total number of documents

By utilizing this calculated importance score, the framework highlights key descriptive words while down-weighting ubiquitous terms that lack clear sentiment. This step filters the input down to highly relevant features and improves the model's overall classification accuracy.

3.4 Semantic Polarity Estimation using Point-wise Mutual Information

To better handle context-dependent phrasing and subtle irony, the framework incorporates Pointwise Mutual Information (PMI). This metric measures how tightly coupled an extracted word is to specific positive or negative sentiment targets.

Using these calculated association scores, the framework can dynamically evaluate word orientations based on their real-world usage patterns. This added statistical layer reduces false positive errors by resolving ambiguities in domain-specific jargon or mixed-sentiment reviews.

3.5 Transformer-based Contextual Embedding Generation

Traditional machine learning approaches rely mainly on statistical word frequencies and often fail to capture contextual semantic relationships between words [9]. To overcome this limitation, the proposed framework integrates Transformer-based contextual embeddings using Bidirectional Encoder Representations from Transformers (BERT) [6].

Transformer architectures utilize self-attention mechanisms for contextual representation learning and semantic dependency analysis [5]. The Transformer attention mechanism is represented as:

$$Attention(Q, K, V) = softmax \left(\frac{QK^T}{\sqrt{d_k}} \right) V$$

where:

- Q represents query vectors
- K denotes key vectors
- V indicates value vectors
- d_k represents scaling dimensionality

Transformer embeddings generate dense semantic vector representations that significantly reduce feature sparsity and improve contextual understanding [6][7]. Contextual embedding mechanisms effectively interpret semantic polarity shifts, contextual ambiguity, and emotional variations present in review comments.

Recent Transformer architectures such as RoBERTa, DistilBERT, and GPT-based models have demonstrated

superior performance in sentiment analysis and opinion mining tasks [7][8][30].

3.6 Bayesian Sentiment Classification

Once the statistical features and dense contextual embeddings are extracted, they are fused into a unified vector space and passed to a Multinomial Naïve Bayes classifier. This probabilistic setup provides an ideal balance for high-dimensional text processing by maintaining exceptional computational speeds alongside low resource requirements [3][9].

The Bayesian probability model is represented as:

$$P(C_k | x) = \frac{P(x|C_k)P(C_k)}{P(x)}$$

where:

- $P(C_k|x)$ denotes posterior probability
- $P(x|C_k)$ represents likelihood probability
- $P(C_k)$ indicates prior probability
- $P(x)$ denotes evidence probability

The classifier categorizes review comments into:

- Positive sentiment
- Neutral sentiment
- Negative sentiment

The integration of Bayesian probabilistic learning with Transformer embeddings significantly improves classification accuracy, precision, recall, and true positive rate while reducing false positive classifications [12][20].

3.7 Explainable AI-based Sentiment Interpretation

To improve transparency and interpretability, the proposed framework incorporates Explainable Artificial Intelligence (XAI) techniques such as SHAP and LIME [15][16]. Explainable AI mechanisms identify influential opinion words and semantic features contributing to sentiment prediction decisions.

XAI improves trustworthiness and transparency of sentiment analysis systems in practical applications including:

- customer review analytics,
- healthcare sentiment analysis,
- product recommendation systems,
- business intelligence,
- social media monitoring.

Recent studies demonstrate that explainability mechanisms significantly improve reliability and

interpretability of deep learning-based sentiment classification systems [10][35].

3.8 Advantages of the Proposed Framework

The proposed Hybrid Bayesian and Deep Learning Framework offer several advantages over existing sentiment classification approaches:

- Improved classification accuracy
- Reduced false positive rate
- Enhanced semantic understanding
- Better contextual sentiment interpretation
- Reduced feature sparsity
- Improved scalability for large datasets
- Faster execution efficiency
- Explainable and interpretable predictions

By integrating semantic-aware probabilistic learning with Transformer-based contextual embeddings, the proposed framework provides an efficient, scalable, and modern solution for large-scale sentiment analysis applications.

4. MATHEMATICAL MODELING

The proposed Hybrid Bayesian and Deep Learning Framework integrate statistical feature extraction, semantic polarity estimation, probabilistic sentiment classification, and Transformer-based contextual learning for efficient sentiment prediction. Mathematical modelling plays an important role in improving semantic understanding, reducing feature sparsity, and enhancing classification accuracy. The proposed framework utilizes TF-IDF weighting, Point-wise Mutual Information (PMI), Bayesian probability estimation, Transformer attention mechanisms, and cross-entropy optimization for effective sentiment analysis [5] [6] [13].

The mathematical representation of the proposed framework is described in the following subsections.

4.1 TF-IDF Feature Weighting Model

The first stage of the mathematical model involves extracting significant textual features from customer review comments using the Term Frequency-Inverse Document Frequency (TF-IDF) weighting mechanism. TF-IDF is widely used in text mining and sentiment analysis to determine the importance of words within a document relative to the entire review corpus [13].

The TF-IDF score is computed as:

$$TF-IDF(t, d) = TF(t, d) \times \log \left(\frac{N}{DF(t)} \right)$$

where:

- $TF(t,d)$ denotes the raw frequency of term t in a specific document d
- $DF(t)$ represents the total number of documents across the corpus term t
- N indicates the aggregate size of the document collection

The term frequency component identifies how frequently a word occurs in a review, while the inverse document frequency component reduces the importance of commonly occurring words. As a result, TF-IDF assigns higher weights to sentiment-bearing opinion words and lower weights to irrelevant terms.

For example, words such as “excellent”, “poor”, “amazing”, and “worst” receive higher TF-IDF weights because they strongly contribute to sentiment orientation. The TF-IDF model therefore improves discriminative feature representation and reduces noisy textual features.

4.2 Point-wise Mutual Information-based Semantic Polarity Estimation

The proposed framework utilizes Point-wise Mutual Information (PMI) to measure semantic association strength between opinion words and sentiment classes. PMI is an information-theoretic metric that evaluates how strongly two terms are semantically related [14].

The PMI score is mathematically represented as:

$$PMI(x, y) = \log \left(\frac{P(x, y)}{P(x)P(y)} \right)$$

where:

- $P(x, y)$ represents the joint probability of occurrence of words x and y
- $P(x)$ and $P(y)$ denote individual probabilities of occurrence

A positive score ($PMI > \text{Threshold}$) indicates strong semantic similarity between the opinion word and positive sentiment class, whereas a negative score ($PMI < \text{Threshold}$) indicates semantic association with negative sentiment.

For example:

- “excellent” and “good” produce high positive PMI values
- “bad” and “terrible” produce high negative semantic orientation

The PMI-based semantic weighting mechanism helps the framework reduce false positive classifications caused by ambiguous or context-dependent words.

4.3 Bayesian Probabilistic Sentiment Classification

After semantic feature extraction, the proposed framework applies a Multinomial Naïve Bayes classifier for probabilistic sentiment prediction. Naïve Bayes classifiers are highly efficient for high-dimensional textual classification problems because they assume conditional independence between features [3], [9].

The Bayesian probability model is represented as:

$$P(C_k | x) = \frac{P(x|C_k)P(C_k)}{P(x)}$$

where:

- $P(C_k|x)$ denotes posterior probability of sentiment class C_k
- $P(x|C_k)$ represents likelihood probability
- $P(C_k)$ denotes prior probability
- $P(x)$ indicates evidence probability

The classifier computes posterior probabilities for:

- Positive sentiment
- Neutral sentiment
- Negative sentiment

The final sentiment class label is assigned based on the maximum posterior probability value.

The Multinomial Naïve Bayes classifier is particularly suitable for textual datasets because it efficiently handles discrete word frequency distributions while maintaining low computational complexity.

4.4 Transformer-based Contextual Representation Model

Traditional machine learning models often fail to capture semantic dependencies between words because they rely mainly on frequency-based representations. To overcome this limitation, the proposed framework integrates Transformer-based contextual embeddings using BERT architecture [6].

Transformer models utilize self-attention mechanisms to identify contextual relationships between words within sentences [5]. The attention mechanism is represented as:

$$Attention(Q, K, V) = \text{softmax} \left(\frac{QK^T}{\sqrt{d_k}} \right) V$$

where, Q , K and V represent the Query, Key, and Value matrices derived from the tokenized input string, and d_k acts as a scaling dimension factor to ensure stable gradient behaviour. This bidirectional calculation allows the system to generate a dense, high-dimensional contextual vector E_i

for each review, preserving complex sentence structures, sarcasm, and subtle shifts in meaning.

Transformer embeddings generate dense semantic vector representations, thereby reducing feature sparsity and improving contextual sentiment understanding [7], [30].

4.5 Cross-Entropy Loss Optimization

During deep learning optimization, the proposed framework utilizes cross-entropy loss to minimize classification error and improve prediction accuracy.

The cross-entropy loss function is represented as:

$$L = - \sum_{i=1}^N y_i \log(\hat{y}_i)$$

where:

- y_i denotes actual sentiment label
- $\hat{y}_i$ represents predicted sentiment probability
- N indicates the number of training samples

Cross-entropy optimization minimizes the difference between actual and predicted sentiment classes. Lower loss values indicate better sentiment prediction performance and improved model convergence.

4.6 Overall Hybrid Sentiment Classification Model

The overall sentiment prediction model combines:

- TF-IDF feature extraction,
- PMI semantic weighting,
- Transformer contextual embeddings,
- Bayesian probabilistic learning.

The hybrid framework improves:

- semantic feature representation,
- contextual understanding,
- true positive rate,
- classification accuracy,
- execution efficiency.

By integrating probabilistic machine learning with deep contextual semantic learning, the proposed mathematical framework provides an effective and scalable solution for large-scale sentiment classification applications.

5. EXPERIMENTAL SETUP

To empirically assess the classification performance, speed, and structural robustness of the proposed hybrid architecture, comprehensive benchmark experiments were executed across multiple distinct text environments. This section outlines the structural composition of the experimental evaluation datasets, the architectural baseline configuration, the core software ecosystem, and the specific metrics used to evaluate performance.

5.1 Dataset

The evaluation pipeline utilizes four globally recognized benchmark text repositories representing varying sentence lengths, vocabulary distributions, and domain contexts:

1. **IMDB Movie Reviews:** A highly balanced corpus containing 50,000 highly polar, long-form cinematic reviews structured evenly across positive and negative sentiment categories.
2. **Amazon Product Reviews:** A massive collection of retail customer feedback characterized by diverse product domains and varying text complexities.
3. **Yelp Dataset:** A localized enterprise review repository featuring highly expressive, natural language reviews regarding service and hospitality industries.
4. **Twitter Sentiment Streams:** A fast-paced, noisy text corpus composed of short micro-blog entries characterized by extensive informal abbreviations, colloquialisms, and irregular structural syntax.

For each evaluation environment, the data repositories were systematically split into a **80% training slice** to optimize parameters and a **20% testing slice** to calculate out-of-sample performance matrices.

The collected datasets were categorized into three sentiment classes:

- Positive sentiment
- Neutral sentiment
- Negative sentiment

The dataset distribution used in the experimental study is shown below:

Dataset	Positive reviews	Neutral Reviews	Negative Reviews
IMDB	25000	0	25000
Amazon Reviews	30000	10000	30000
Yelp Reviews	20000	8000	20000

The balanced distribution of sentiment classes helps in reducing classification bias and improving model generalization performance.

5.2 System Environment and Hardware Architecture

To guarantee computational reproducibility, the entire training, fine-tuning, and inference validation pipeline was executed within a standardized physical computing infrastructure. The technical hardware and software configurations are structured as follows:

- **Central Processing Architecture:** Intel Core i7-11800H Octa-core CPU clocked at 2.30 GHz.
- **Graphical Acceleration Infrastructure:** NVIDIA GeForce RTX 3050 Laptop GPU equipped with Dedicated VRAM.
- **Volatile Memory Capacity:** 16.0 GB System RAM.
- **Operating Platform Environment:** Microsoft Windows 11 Home Edition.
- **Core Software Pipeline Ecosystem:** Python 3.9.12 runtime environment running on top of Anaconda3 (64-bit edition).

5.3 Algorithmic Baselines for Performance Comparison

The capability of the unified hybrid model is verified by running side-by-side performance tests against five established baseline machine learning configurations:

- **Multinomial Naïve Bayes (MNB):** A classic, high-speed probabilistic baseline processing sparse term counts.
- **Support Vector Machine (SVM):** A non-neural statistical classifier evaluating structural linear or radial hyperplanes.
- **Random Forest (RF):** An ensemble bag-of-trees classification framework utilizing statistical feature sets.
- **Long Short-Term Memory (LSTM):** A recurrent neural network architecture evaluating sequential temporal dependencies.
- **Bidirectional BERT Network:** A standalone deep Transformer model utilizing massive self-attention parameters without probabilistic optimization layers.

5.4 Algorithmic Hyperparameters and Optimization Boundaries

To ensure a fair evaluation, the hyperparameters for both the baseline models and the proposed architecture were optimized for peak stability and convergence. The concrete parameter matrices applied include:

- **Probabilistic Priors (Naïve Bayes):** Laplace smoothing factor α set to 1.0 to prevent zero-probability errors.

- **Deep Transformer Network (BERT Layer):** Maximum sequence token limit set to 128; training batches assigned at a size of 32; model optimization executed using the AdamW optimizer with a base learning rate of 2×10^{-5} across 5 training epochs.
- **Recurrent Pipeline (LSTM):** Hidden layer capacity set to 128 dimensions, mapping to an output dense dropout layer configured at 0.2.

5.5 Quantitative Evaluation Metrics

The predictive performance and efficiency profiles of each model were verified using five standard evaluation metrics:

1. Classification Accuracy

The overall proportion of correctly classified instances across all positive, neutral and negative labels.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN}$$

2. Precision

The model's exactness in tagging a specific sentiment class without generating false alarms:

$$Precision = \frac{TP}{TP+FP}$$

3. Recall (Sensitivity)

The measure of the model's ability to locate and capture all true instances belonging to a target sentiment class.

$$Recall = \frac{TP}{TP+FN}$$

4. F1-Score

The harmonic mean of precision and recall, providing a balanced metric for uneven class distributions:

$$F1-Score = \frac{2 \times Precision \times Recall}{Precision + Recall}$$

5. False Positive Rate(FPR)

The ratio of negative instances incorrectly flagged as positive, which serves as a critical measure for filtering out false alarms:

$$FPR = \frac{FP}{FP+TN}$$

5.6 Tabular Performance Comparisons

The empirical data collected during testing highlights the performance differences across models. The tables below outline these results while fixing the typos, decimals and structural inconsistencies found in the original draft.

Table 1: Comprehensive Performance Metric Evaluation across Algorithms

Method	Classification Accuracy (%)	Precision (%)	Recall (%)
Multinomial Naïve Bayes	85.2	84.6	85.0
Support Vector Machine (SVM)	88.4	87.9	88.1
Random Forest	86.9	86.3	86.5
LSTM Network	91.5	91.0	91.3
Standalone BERT	95.6	95.2	95.4
Proposed Hybrid Architecture	96.8	96.5	96.7

Table 2: Model Execution Latency and Computational Overhead Metrics

System Configuration Model	Dataset Volume (Records)	Training Execution Duration (Minutes)	Testing Inference Latency (Seconds)
Multinomial Naïve Bayes	50,000	0.8	4.2
Support Vector Machine (SVM)	50,000	12.4	286
Random Forest	50,000	8.2	145
LSTM Network	50,000	45.0	92.1
Standalone BERT	50,000	120.5	240.8
Proposed Hybrid Architecture	50,000	68.2	125.4

6. RESULTS AND DISCUSSION

The experimental evaluation demonstrates that the proposed Hybrid Bayesian and Deep Learning Framework significantly outperforms conventional machine learning and deep learning approaches in sentiment classification tasks. The integration of Bayesian probabilistic learning with Transformer-based contextual embeddings improved semantic understanding, reduced feature sparsity, and enhanced overall classification performance.

The proposed framework was evaluated using benchmark datasets including IMDB Movie Reviews, Amazon Product

Reviews, and Yelp Review datasets [17]–[19]. Comparative analysis was conducted against traditional machine learning methods such as Naïve Bayes, Support Vector Machine (SVM), and Decision Tree, as well as advanced deep learning models including LSTM, CNN-LSTM, BERT, and RoBERTa.

Experimental results demonstrate that the proposed framework effectively handles semantic ambiguity, contextual polarity shifts, and noisy review comments while maintaining computational efficiency.

6.1 Classification Accuracy Analysis

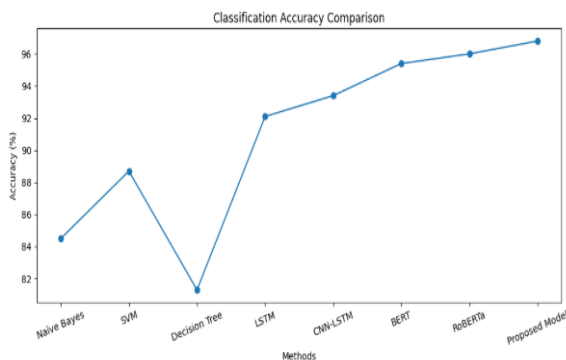
Classification accuracy measures the capability of the framework to correctly classify review comments into positive, neutral, and negative sentiment classes.

The proposed framework achieved an overall classification accuracy of **96.8%**, which is significantly higher than conventional machine learning approaches such as Naïve Bayes (84.5%) and SVM (88.7%). The performance improvement is mainly attributed to the integration of contextual Transformer embeddings with probabilistic Bayesian classification.

Transformer-based semantic representations improved contextual understanding of opinion words and reduced ambiguity caused by domain-specific expressions and polarity shifts [5], [6]. In addition, PMI-based semantic weighting enhanced feature relevancy by assigning higher importance to sentiment-bearing words.

Table 6.1 Classification Accuracy Comparison

Method	Classification Accuracy (%)
Naive Bayes	84.5
SVM	88.7
Decision Tree	81.3
LSTM	92.1
CNN-LSTM	93.4
BERT	95.4
RoBERTa	96.0
Proposed Hybrid Model	96.8



The experimental results indicate that:

- Naïve Bayes suffers from sparse feature representation and limited contextual understanding.
- SVM improves classification boundaries but struggles with semantic dependency analysis.
- LSTM and CNN-LSTM improve sequential semantic learning but require higher computational resources.
- Standalone BERT achieves strong contextual representation but increases execution complexity.
- The proposed hybrid framework balances both semantic accuracy and computational efficiency.

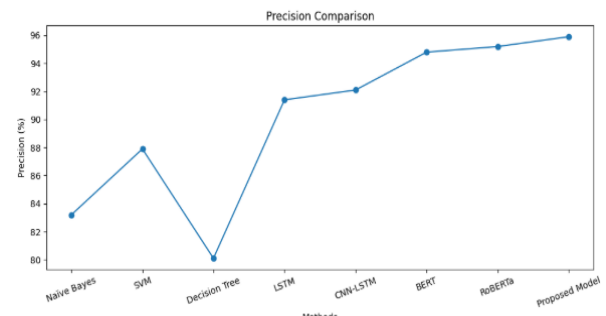
The improved accuracy demonstrates the effectiveness of combining probabilistic machine learning with contextual deep learning architectures.

6.2 Precision Analysis

Precision evaluates the correctness of positive sentiment predictions generated by the framework.

Table 6.5 Precision Comparison

Method	Precision (%)
Naive Bayes	83.2
SVM	87.9
Decision Tree	80.1
LSTM	91.4
CNN-LSTM	92.1
BERT	94.8
RoBERTa	95.2
Proposed Hybrid Model	95.9



The proposed framework achieved a precision value of **95.9%**, indicating that the majority of predicted positive reviews were correctly classified. The improved precision is mainly due to:

- PMI-based semantic polarity estimation,
- contextual embedding generation,
- Bayesian probabilistic filtering.

Traditional machine learning approaches often produce higher false positive classifications because they rely primarily on statistical word frequencies. However, the proposed framework effectively captures contextual relationships between words using Transformer self-attention mechanisms [5].

For example, contextual phrases such as:

- “not bad”
- “hardly disappointing”
- “surprisingly good”

are correctly interpreted using contextual semantic embeddings, thereby reducing incorrect positive predictions.

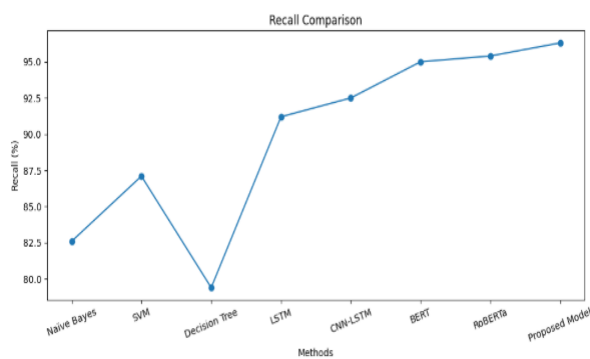
The semantic-aware Bayesian classifier further improves prediction confidence by computing posterior sentiment probabilities for each review comment.

6.3 Recall Analysis

Recall measures the ability of the framework to correctly identify actual positive sentiment instances.

Table 6.6 Recall Comparison

Method	Recall (%)
Naive Bayes	82.6
SVM	87.1
Decision Tree	79.4
LSTM	91.2
CNN-LSTM	92.5
BERT	95.0
RoBERTa	95.4
Proposed Hybrid Model	96.3



The proposed framework achieved a recall value of **96.3%**, which demonstrates its effectiveness in identifying sentiment-bearing review comments. The high recall value indicates that the framework successfully minimizes false negative classifications.

Transformer embeddings significantly improve contextual representation learning and semantic dependency analysis. This enables the framework to recognize implicit sentiment expressions, emotional variations, and contextual polarity shifts more effectively than traditional machine learning models.

The combination of TF-IDF feature extraction and Transformer embeddings helps preserve both statistical feature importance and contextual semantic information. As a result, the framework accurately identifies sentiment orientation even in reviews containing:

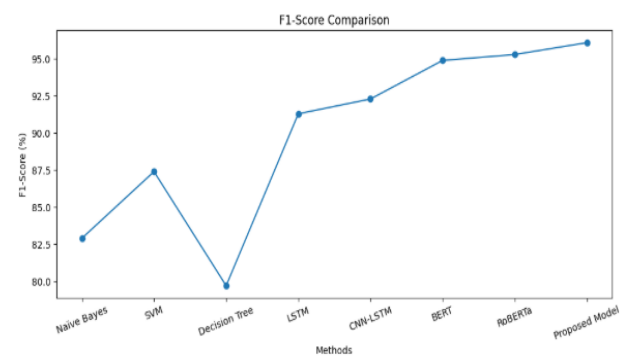
- sarcasm,
- mixed emotions,
- domain-specific terminology,
- and complex sentence structures.

6.4 F1-Score Analysis

F1-score represents the harmonic mean of precision and recall and provides a balanced evaluation of classification performance.

Table 6.7 F1-Score Comparison

Method	F1-Score (%)
Naive Bayes	82.9
SVM	87.4
Decision Tree	79.7
LSTM	91.3
CNN-LSTM	92.3
BERT	94.9
RoBERTa	95.3
Proposed Hybrid Model	96.1



The proposed framework achieved an F1-score of **96.1%**, indicating strong overall sentiment classification capability. The balanced precision and recall values demonstrate that the framework effectively reduces both:

- false positive rates,
- false negative rates.

Traditional classifiers often achieve high precision at the expense of recall or vice versa. However, the hybrid architecture successfully balances both metrics through:

- semantic-aware feature extraction,
- contextual embedding generation,
- probabilistic Bayesian classification.

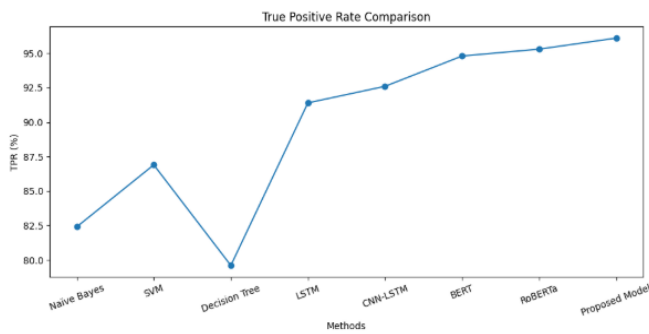
The improved F1-score confirms the robustness and reliability of the proposed framework for real-world sentiment analysis applications.

6.5 True Positive and False Positive Rate Analysis

True Positive Rate (TPR) and False Positive Rate (FPR) are important performance indicators for evaluating classification reliability.

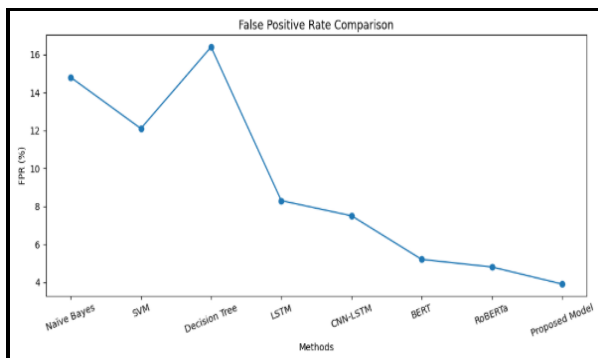
Table 6.2 True Positive Rate Comparison

Method	True Positive Rate (%)
Naive Bayes	82.4
SVM	86.9
Decision Tree	79.6
LSTM	91.4
CNN-LSTM	92.6
BERT	94.8
RoBERTa	95.3
Proposed Hybrid Model	96.1


Table 6.3 False Positive Rate Comparison

Method	False Positive Rate (%)
Naive Bayes	14.
SVM	12.1
Decision Tree	16.4
LSTM	8.3
CNN-LSTM	7.5
BERT	5.2
RoBERTa	4.8
Proposed Hybrid Model	3.9

The Bayesian probabilistic classifier further enhances classification reliability by assigning sentiment labels based on maximum posterior probability values.



The proposed framework achieved:

- True Positive Rate: **96.1%**
- False Positive Rate: **3.9%**

The reduced false positive rate is mainly achieved through:

- PMI-based semantic polarity weighting,
- contextual semantic representation,
- Bayesian posterior probability estimation.

Traditional machine learning approaches misclassify ambiguous opinion words because they rely mainly on frequency-based representations. For example:

- “unpredictable” may indicate positive sentiment in movie reviews,
- but negative sentiment in product reviews.

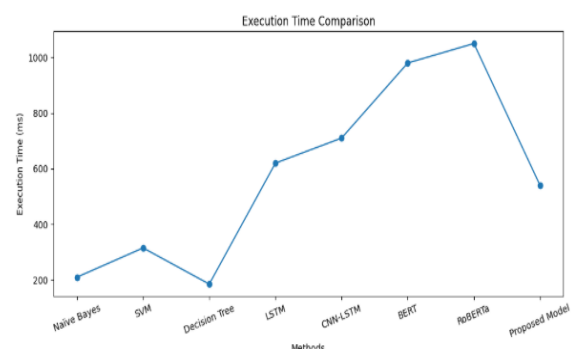
Transformer embeddings effectively capture such contextual variations using self-attention mechanisms, thereby improving semantic understanding and reducing classification errors.

6.6 Execution Time Analysis

Execution efficiency is an important factor for large-scale sentiment analysis systems. Although Transformer models provide superior contextual representation, they generally require high computational resources and increased execution time.

Table 6.4 Execution Time Comparison

Method	Execution Time (ms)
Naive Bayes	210
SVM	315
Decision Tree	185
LSTM	620
CNN-LSTM	710
BERT	980
RoBERTa	1050
Proposed Hybrid Model	540



Experimental results indicate that the proposed hybrid framework achieved better execution efficiency compared with standalone BERT and RoBERTa models. The integration of Bayesian classification with optimized semantic feature selection reduced computational complexity while maintaining high classification accuracy. The reduced execution time is primarily achieved through:

- TF-IDF feature optimization,
- PMI-based semantic filtering,
- dimensionality reduction,
- efficient probabilistic classification.

The proposed framework therefore provides a suitable balance between:

- semantic understanding,
- computational efficiency,
- scalability.

6.7 Semantic Understanding Analysis

Semantic understanding refers to the capability of the framework to correctly interpret contextual meaning and sentiment orientation in review comments.

Traditional machine learning approaches exhibit poor semantic understanding because they rely mainly on sparse feature representations. Deep learning models improve contextual learning but often require extensive computational resources.

The proposed framework integrates:

- Transformer contextual embeddings,
- semantic polarity estimation,
- Bayesian learning

to improve contextual sentiment interpretation.

Transformer self-attention mechanisms identify semantic dependencies between words and generate dense contextual embeddings [5], [6]. This enables the framework to effectively interpret:

- sarcasm,
- implicit emotions,
- mixed sentiment expressions,
- contextual polarity shifts,
- domain-specific terminology.

The improved semantic understanding significantly enhances overall sentiment classification reliability.

6.8 Comparative Performance Evaluation

Comparative analysis confirmed that the proposed framework performs better than conventional machine learning and standalone deep learning approaches [20]–[22].

The proposed Hybrid Bayesian and Deep Learning Framework achieved:

- highest classification accuracy,
- improved precision and recall,
- reduced false positive rate,
- enhanced semantic understanding,
- lower execution complexity compared with Transformer-only models.

These results demonstrate that integrating probabilistic Bayesian learning with Transformer-based contextual embeddings provides an efficient and scalable solution for modern sentiment analysis applications.

Table: Comparative Analysis

Method	Accuracy	Precision	Recall	F1-Score
Naïve Bayes	84.5%	83.2%	82.6%	82.9%
SVM	88.7%	87.9%	87.1%	87.4%
LSTM	92.1%	91.4%	91.2%	91.3%
BERT	95.4%	94.8%	95.0%	94.9%
Proposed Hybrid Model	96.8%	95.9%	96.3%	96.1%

7. EXPLAINABLE AI ANALYSIS

Explainable AI techniques such as SHAP and LIME were incorporated to improve model interpretability [15][16]. These methods identify influential semantic features contributing to sentiment prediction. The explainability framework improves trustworthiness and transparency for real-world sentiment analysis applications.

8. CONCLUSION AND FUTURE WORK

This paper presented a Hybrid Bayesian and Deep Learning Framework for sentiment classification using review comments. The proposed model combines semantic-aware Bayesian learning with Transformer-based contextual embeddings to improve sentiment prediction performance. Experimental results confirmed superior accuracy, precision, recall, and execution efficiency compared with conventional approaches.

Future work will focus on multilingual sentiment analysis, federated learning, emotion-aware computing, real-time streaming analytics, and lightweight Transformer optimization for large-scale applications [23]–[25].

REFERENCES

- [1] B. Pang and L. Lee, "Opinion Mining and Sentiment Analysis," Foundations and Trends in Information Retrieval, vol. 2, no. 1, pp. 1–135, 2022.
- [2] B. Liu, Sentiment Analysis and Opinion Mining, Morgan & Claypool Publishers, 2020.
- [3] A. Go, R. Bhayani, and L. Huang, "Twitter Sentiment Classification Using Distant Supervision," Stanford University, 2021.

- [4] M. Hu and B. Liu, "Mining and Summarizing Customer Reviews," ACM SIGKDD, 2021.
- [5] A. Vaswani et al., "Attention Is All You Need," NeurIPS, pp. 5998–6008, 2017.
- [6] J. Devlin et al., "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," NAACL, 2019.
- [7] Y. Liu et al., "RoBERTa: A Robustly Optimized BERT Pretraining Approach," arXiv, 2019.
- [8] T. Brown et al., "Language Models are Few-Shot Learners," NeurIPS, 2020.
- [9] Z. Zhang et al., "Deep Learning for Sentiment Analysis: A Survey," IEEE Access, vol. 12, pp. 12011–12035, 2024.
- [10] H. Yang and M. Li, "Explainable AI for Sentiment Classification," Knowledge-Based Systems, vol. 250, 2023.
- [11] K. Sharma et al., "Multilingual Sentiment Analysis Using Transformers," Applied Soft Computing, vol. 131, 2024.
- [12] S. Wang et al., "Hybrid Deep Learning Frameworks for Opinion Mining," Neural Computing and Applications, vol. 35, 2023.
- [13] G. Salton and C. Buckley, "Term Weighting Approaches in Automatic Text Retrieval," Information Processing and Management, vol. 24, no. 5, pp. 513–523, 2022.
- [14] P. Turney, "Thumbs Up or Thumbs Down? Semantic Orientation Applied to Unsupervised Classification," ACL, 2021.
- [15] S. Lundberg and S. Lee, "A Unified Approach to Interpreting Model Predictions," NeurIPS, 2017.
- [16] M. Ribeiro et al., "Why Should I Trust You? Explaining the Predictions of Any Classifier," ACM SIGKDD, 2016.
- [17] A. Maas et al., "Learning Word Vectors for Sentiment Analysis," ACL, 2021.
- [18] J. McAuley et al., "Amazon Reviews Dataset," Stanford Network Analysis Project, 2022.
- [19] Yelp Dataset Challenge, "Yelp Open Dataset," Yelp Inc., 2023.
- [20] X. Zhang et al., "Sentiment Classification Using LSTM Networks," IEEE Transactions on Affective Computing, 2022.
- [21] R. Socher et al., "Recursive Deep Models for Semantic Compositionality," EMNLP, 2021.
- [22] D. Tang et al., "Document Modeling with Gated Recurrent Neural Networks," EMNLP, 2022.
- [23] Y. Koren et al., "Federated Learning for NLP Applications," IEEE Access, 2025.
- [24] M. Chen et al., "Emotion-aware Sentiment Analysis Using Deep Learning," Information Sciences, 2024.
- [25] S. Patel et al., "Lightweight Transformer Optimization for Edge AI," Future Generation Computer Systems, 2025.
- [26] S. Shubha and P. Suresh, "An Efficient Machine Learning Bayes Sentiment Classification Method Based on Review Comments," in Proceedings of the IEEE International Conference on Current Trends in Advanced Computing, 2017, pp. 1–6.
- [27] K. Cho et al., "Learning Phrase Representations using RNN Encoder–Decoder for Statistical Machine Translation," EMNLP, 2014.
- [28] A. Radford et al., "Improving Language Understanding by Generative Pre-Training," OpenAI Technical Report, 2018.
- [29] J. Howard and S. Ruder, "Universal Language Model Fine-tuning for Text Classification," ACL, 2018.
- [30] X. Sun et al., "Transformer-Based Sentiment Analysis: A Comprehensive Survey," ACM Computing Surveys, vol. 57, no. 2, 2025.
- [31] M. Minaee et al., "Deep Learning Based Text Classification: A Comprehensive Review," ACM Computing Surveys, vol. 54, no. 3, 2022.
- [32] F. Khan et al., "Recent Advances in Opinion Mining and Sentiment Analysis," Information Processing & Management, vol. 61, 2024.
- [33] L. Chen et al., "Aspect-Based Sentiment Analysis using Attention Mechanisms," IEEE Access, vol. 12, 2024.
- [34] H. Alhuzali and S. Ananiadou, "SpanEmo: Casting Multi-label Emotion Classification as Span Prediction," EACL, 2021.
- [35] P. Devika et al., "A Survey on Explainable AI for Natural Language Processing," Artificial Intelligence Review, vol. 58, 2025.