

DOMAIN-ADAPTIVE SENTIMENT ANALYSIS THROUGH BAYESIAN NETWORK–LOGISTIC REGRESSION-BASED PROBABILISTIC DEPENDENCY MODELING

RAMAJAYAM G¹, Dr. R. VIDYA BANU²

¹ Research Scholar, ² Assistant Professor

^{1,2}LRG Government Arts College for Women, Tirupur, India

E-mail: ¹ramajayamgurusamy@gmail.com

ABSTRACT

Cross-domain sentiment classification presents persistent challenges in opinion mining due to vocabulary drift, contextual ambiguity, and polarity inconsistency across product domains. Traditional classifiers trained on a single domain often fail to generalize, reducing performance when exposed to new, structurally distinct datasets. This research introduces a probabilistically grounded sentiment classification framework—Bayesian Network–Logistic Regression (BN-LR)—designed to address these challenges within multi-domain online review environments. BN-LR integrates a Bayesian Network to model conditional dependencies among sentiment-bearing features, capturing latent inter-feature relationships across syntactic structures. These probabilistic insights are dynamically incorporated into a logistic regression classifier, enabling adaptive feature weighting and uncertainty-aware sentiment inference. As an IT contribution, BN-LR offers a scalable, interpretable, and statistically principled solution suitable for intelligent recommendation engines, feedback analytics, and sentiment-based decision systems across digital platforms. Evaluated on Amazon reviews across four domains, BN-LR consistently delivers high accuracy without requiring domain-specific retraining or external lexicons. The proposed framework enhances real-world information systems by enabling robust cross-domain sentiment generalization, fulfilling a critical need in adaptive text analytics for IT-driven e-commerce intelligence.

Keywords: *Cross-Domain Sentiment Analysis, Bayesian Network, Logistic Regression, Probabilistic Modelling, Online Product Reviews, Domain Adaptation*

1. INTRODUCTION

Sentiment analysis in online shopping environments presents unique challenges driven by variability in product categories, customer expectations, and linguistic diversity. These challenges hold particular significance in the Information Technology domain, where intelligent systems must automatically interpret, classify, and respond to unstructured textual data at scale [1]. Consumer reviews, often processed by recommender systems, search engines, and customer service bots, require adaptive sentiment interpretation to maintain relevance and personalization. The complexity increases under domain shift, a common real-world IT scenario, where training and deployment environments differ substantially. Addressing this issue, cross-domain sentiment analysis has emerged as a critical subfield of applied natural language processing and IT-driven business intelligence. The proposed BN-LR framework aligns with this need, offering a probabilistic, explainable, and domain-adaptive architecture for use in IT systems handling dynamic user-generated content. The work strengthens the

foundation for scalable, intelligent, and context-sensitive feedback analysis in digital ecosystems, reflecting its core contribution to IT research [2].

Cross-domain sentiment analysis emerges as a necessary framework to mitigate performance degradation caused by these domain discrepancies. A model designed for one domain frequently fails to generalize effectively to another due to shifts in sentiment-bearing expressions, semantic emphasis, and syntactic structure [3]. Online book reviews emphasize narrative engagement and author style, whereas electronic reviews prioritize functionality and durability. This context-specific divergence introduces polarity ambiguity and lexical misalignment that challenge direct model transfer. Bridging these domain gaps requires analytical strategies that capture transferable sentiment cues while retaining domain-specific nuances, ensuring that classifiers adapt intelligently without retraining for each new context [4].

Granular review analysis demands representational techniques that adaptively

recalibrate feature relevance based on contextual positioning. For example, if influenced by negations, comparatives, or feature-targeted qualifiers, identical adjectives may carry divergent sentiment intensities across reviews [5]. Sentiment variability also arises from user-specific value, aesthetics, or functionality interpretation. Capturing these probabilistic dependencies and calibrating their influence improves classification consistency in cross-domain environments [6]. Sentiment analysis frameworks grounded in probabilistic reasoning and dependency modelling enable scalable, interpretable, and adaptive systems suited for the heterogeneous nature of online customer reviews. Such frameworks prioritize generalization, ensuring robust sentiment insight across shifting commercial and linguistic contexts [7].

Applying probabilistic theory in cross-domain sentiment classification supports adaptive polarity reasoning by estimating conditional relationships among features. These probabilistic estimations reflect how the sentiment of a given word or phrase depends on its surrounding context, ensuring interpretability and robustness across domains [8]. Review texts often contain uncertainty, mixed sentiments, and subtle emotional transitions. Probabilistic models quantify this uncertainty, calibrating predictions based on inferred likelihoods rather than deterministic mappings [9]. This enhances the system's ability to accommodate context variation, lexical ambiguity, and sentiment intensity modulation. By combining dependency modelling with probabilistic estimation, sentiment classification systems achieve scalable, generalizable performance across dynamic online shopping contexts, delivering consistent polarity interpretations across domain boundaries [10].

1.1. Problem statement

Cross-domain opinion mining encounters significant challenges when classifying sentiment across diverse product categories due to variability in linguistic structures, domain-specific expressions, and inconsistent sentiment representation. Models trained on one domain often experience accuracy degradation when applied to another, primarily because of lexical mismatches, semantic shifts, and contextual ambiguity. Expressions conveying positive sentiment in one domain may hold neutral or negative implications in another, complicating the interpretation process. In addition, syntactic constructs and opinion-carrying phrases vary widely depending on product type, usage context, and reviewer intention. Sentiment-bearing words lack

uniformity across domains, resulting in misclassification and reduced model generalizability. Domain shifts cause statistical divergence, rendering static sentiment features insufficient for transferability. Existing approaches often rely on explicit lexical cues, which fail to capture dependency relationships and probabilistic sentiment transitions. These issues collectively limit the effectiveness of traditional sentiment classification techniques when applied to heterogeneous review datasets.

1.2. Motivation

The rapid expansion of online shopping platforms has led to diverse customer reviews, each containing sentiment expressions tied to specific product categories. Consumers rely on this feedback to make informed decisions, while businesses extract actionable insights to enhance product development, service quality, and customer experience. However, these reviews differ structurally and contextually across domains, causing classification inconsistencies when sentiment models are reused without adaptation. The inability to maintain accuracy across varying domains presents a theoretical and practical barrier to large-scale sentiment analysis deployment. Addressing this issue is crucial for building robust systems across heterogeneous data environments. Integrating probabilistic reasoning and feature dependency modelling has emerged as a promising approach to resolving ambiguity, improving interpretability, and achieving consistent performance across domains. Motivated to enhance cross-domain sentiment generalization, this research seeks to design a sentiment analysis framework grounded in statistical inference and probabilistic feature understanding.

1.3. Objective

This research aims to develop a domain-adaptive sentiment classification framework for cross-domain opinion mining in online shopping environments. The primary objective is to accurately classify sentiment across diverse product review categories without requiring domain-specific retraining. The framework will capture probabilistic dependencies among sentiment-bearing features, enabling interpretive flexibility in variable linguistic contexts. It will incorporate uncertainty-aware decision mechanisms to manage ambiguity and minimize classification errors in the presence of domain shifts. The proposed system will extract context-relevant sentiment indicators, estimate conditional sentiment distributions, and calibrate polarity assignments based on probabilistic

evidence. This study ensures that sentiment prediction remains robust, interpretable, and scalable across unlabeled or underrepresented target domains. The methodology will prioritize generalization, transparency, and adaptability, contributing to the broader goal of building resilient sentiment analysis models that align with real-world online shopping review dynamics and support intelligent e-commerce decision-making.

2. LITERATURE REVIEW

Flask-CNN-RoBERTa Review Analyzer [11] integrates RoBERTa embeddings with CNN to classify movie reviews into positive, negative, or neutral categories. RoBERTa generates context-rich vector representations from textual inputs fed into CNN layers for feature extraction and sentiment prediction. The system uses Flask to support real-time sentiment evaluation through a user-friendly interface. It enables instantaneous review classification, supports decision-making in entertainment analytics, and offers scalable architecture for integration into sentiment-driven applications. E-Commerce Hybrid Analyzer [12] maps sentiment keywords from e-commerce reviews into corresponding image representations. These images are processed using Convolutional Neural Networks to extract emotional features, which a Support Vector Machine classifies into sentiment categories. This hybrid model introduces visual dimension into text-based sentiment analysis, capturing subtle emotional indicators not visible in plain text. It enhances the interpretability of sentiment across diverse product domains and contributes to nuanced opinion mining for customer-centric commerce environments. App Review Insight Framework [13] reviews the sentiment using Long Short-Term Memory Networks (LSTM) and Graph Neural Networks (GNN). LSTM captures temporal dependencies in sequential text, identifying usability-related sentiment trends, while GNN models interconnect between review elements such as functionality and user experience. This architecture enables aspect-based sentiment detection for usability metrics. The hybrid design ensures precision in identifying sentiment polarity related to app performance and helps developers prioritize improvements based on sentiment clusters.

Financial Sentiment Prototype [14] uses a supervised cross-momentum contrast strategy. The model aligns internal textual representations with pre-defined sentiment prototypes, continuously

refining sentiment embeddings. This alignment enhances the detection of subtle financial sentiment cues, increasing classifier accuracy and robustness. The supervised structure supports precise sentiment mapping in complex financial narratives, enabling effective tracking of market sentiment. The approach improves the semantic alignment of financial language with sentiment classes under domain-specific conditions. Multimodal Sentiment Fusion Network [15] incorporates text, audio, and video modalities using attention-based and causality-aware mechanisms. This model dynamically weighs the influence of each modality based on interaction context, capturing real-time sentiment cues and intermodal dependencies. Causal inference techniques identify how sentiment signals from one modality influence others. This approach improves prediction accuracy across emotionally rich datasets and ensures robust handling of complex sentiment dynamics in multimodal interaction settings. Cross-Modal Sentiment Synthesizer [16] is a shared-private fusion model for cross-modal sentiment analysis using parallel pathways for standard and unique sentiment features. Shared layers synthesize universal sentiment cues across modalities such as text and audio, while private layers retain modality-specific expressions. This structure enables better integration of multimodal signals while preserving distinct semantic contributions. The system enhances interpretability and robustness by aligning shared features while managing modality variance. It is suitable for sentiment analysis tasks requiring simultaneous processing of multiple input forms.

UrduAspectNet Enhancer [17] tailored for Urdu reviews using Biaffine Attention to capture syntactic dependencies between aspect terms and contextual sentiment expressions. This architecture improves polarity classification by modelling the precise relationship between opinion targets and modifiers in complex Urdu sentence structures. The approach handles syntactic ambiguity and morphological complexity typical of Urdu, enabling domain-specific applications in sentiment intelligence for regional language datasets. It advances aspect-level precision for low-resource language sentiment analysis. Multimodal Sentiment Integrator [18] enriches textual information through NLP-based semantic extraction and fuses it with visual and auditory modalities. The model uses advanced data integration techniques to capture intermodal relationships, generating a unified representation of user sentiment. It improves performance on complex sentiment tasks by incorporating complementary cues from non-textual

inputs. This integration supports holistic sentiment interpretation in environments where users simultaneously communicate emotions through multiple formats, improving analytic reliability. Sentiment Analysis Progress Review [19] reviews the sentiments by examining significant developments and ongoing challenges in sentiment analysis, emphasizing the transition from lexicon-based approaches to deep learning models. The paper outlines the impact of transformer architectures like BERT on sentiment contextualization and identifies unresolved issues such as sarcasm detection, domain transfer, and sentiment bias. It synthesizes technical progress while mapping future research directions. The study is foundational for understanding sentiment classification evolution, highlighting both breakthroughs and methodological limitations.

Multilingual Lexicon Developer [20] focused on multilingual sentiment analysis in software engineering. This approach combines manual linguistic annotation with machine learning techniques to generate accurate polarity dictionaries for multiple languages. The lexicons are designed to support cross-lingual sentiment interpretation in social media and software reviews. By emphasizing syntactic relevance and semantic alignment, the method ensures broad applicability and sentiment coherence across language barriers, supporting sentiment analysis in globalized digital communication ecosystems. Hotel Review Sentiment System [21] combines BERT for context-aware embeddings, Temporal Convolutional Networks (TCN) for local temporal structure, BiLSTM for sequential modelling, and attention mechanisms for focus prioritization. This architecture processes hotel reviews to extract fine-grained sentiment cues across review components such as service, cleanliness, and amenities. The system captures implicit and explicit sentiment expressions, supporting detailed opinion mining in hospitality analytics and contributing to improved customer satisfaction monitoring. Multimodal Sentiment Regulator [22] designed to coordinate information flow across text, image, and audio inputs using a hierarchical regulator module. The framework filters redundant content and applies dynamic attention mechanisms to weigh modality relevance contextually. It synchronizes sentiment signals from diverse inputs to improve polarity consistency and robustness. The model is particularly effective in real-world applications involving emotionally complex and multimodal user feedback, enhancing interpretability and

classification precision across diverse media formats.

Entropy-Based Classifier (EBC) [23] uses a modified maximum entropy framework with IDF-weighted increment adjustment. EBC processes POS-tagged words to identify sentiment-bearing terms, emphasizing low-frequency, high-discriminative features. It applies a semi-supervised bipartite graph clustering technique to transfer sentiment labels from domain-independent to domain-specific words. Classification is refined at the word level using entropy maximization and polarity propagation. The model improves domain adaptability by minimizing labelling requirements and bridging lexical gaps between source and target review domains, supporting sentiment generalization across structurally diverse product categories. Multimodal Aspect-Based Sentiment Analysis (MMASA) [24] performs sentiment classification by combining textual and visual modalities using Bi-LSTM for text encoding and CNN for image feature extraction. MMASA employs multimodal interaction layers to integrate text-image features and applies adversarial training for cross-modal alignment. The model identifies aspects and predicts sentiment polarity using joint embeddings. It supports fine-grained sentiment detection at the aspect level, effectively capturing cross-modal dependencies and contextual polarity shifts, enhancing sentiment classification accuracy in review datasets containing both textual and visual content.

Bioinspired optimization mimics natural processes like evolution, swarming, and adaptation to solve complex problems efficiently, offering scalable, flexible, and intelligent strategies across machine learning, routing, and classification tasks [25]-[44]. These algorithms enhance convergence speed, solution diversity, and adaptability, making them ideal for dynamic, high-dimensional, and real-time environments in various IT and data-driven applications [45]-[66].

3. BAYESIAN NETWORK ENHANCED LOGISTIC REGRESSION

This section is structured first to introduce Bayesian Network structure learning and its role in sentiment classification. The discussion then transitions to the limitations of traditional LR in sentiment analysis and how BN addresses these constraints. Subsequent sections cover Bayesian regularization, probabilistic feature selection, and

Bayesian Model Evidence, demonstrating their impact on optimizing sentiment classification.

3.1. Bayesian network structure learning

The Bayesian Network (BN) structure learning process involves establishing directed acyclic graph (DAG) connections among sentiment-related features to enhance dependency representation. The DAG captures interdependencies between linguistic attributes, metadata, and contextual sentiment cues, optimizing the feature set for logistic regression. Structure learning employs score-based, constraint-based, and hybrid approaches to discover probabilistic relationships among sentiment features, ensuring improved classification precision.

Sentiment analysis in online shopping requires an optimized representation of dependencies between features such as polarity scores, term frequency-inverse document frequency (TF-IDF), syntactic structures, and customer engagement metadata. Bayesian Network assigns directed edges between features, where the probability of sentiment polarity S given a set of extracted features X is expressed as:

$$P(S|X) = \prod_{i=1}^n P(X_i|P_a(X_i)) \quad (1)$$

where $P(S|X)$ represents the conditional probability of sentiment given the extracted features, X_i refers to each feature, and $P_a(X_i)$ denotes the parent features influencing X_i within the Bayesian structure. The network ensures an optimized representation of probabilistic relationships, refining feature selection for classification. The conditional dependency structure dynamically adjusts based on probabilistic scores, providing adaptability in sentiment classification.

Score-based approaches optimize the DAG structure by assigning likelihood scores to network configurations. Bayesian Information Criterion (BIC) or Minimum Description Length (MDL) principles guide the selection of an optimized structure. The likelihood of a DAG G given sentiment training data D is formulated as:

$$L(G|D) = \sum_{i=1}^n \log P(X_i|P_a(X_i), G) \quad (2)$$

where $L(G|D)$ denotes the likelihood score for the given DAG structure G , and $P(X_i|P_a(X_i), G)$ represents the conditional probability of each feature given its parent nodes in the learned structure. This

likelihood-driven approach ensures the DAG aligns with sentiment distribution patterns, refining logistic regression performance by minimizing redundancy.

Constraint-based methods identify Bayesian dependencies by applying conditional independence tests to sentiment-related variables. Dependency constraints guide network formation, ensuring the DAG structure aligns with linguistic and behavioral sentiment dynamics. The probability of independence between two sentiment features X_i and X_j , given a conditioning set Z , is mathematically defined as:

$$P(X_i, X_j|Z) = P(X_i|Z)P(X_j|Z) \quad (3)$$

where $P(X_i, X_j|Z)$ represents the joint probability of two features conditioned on Z , ensuring feature independence. If the independence condition holds, no directed edge is formed between X_i and X_j . This process optimizes sentiment classification by removing redundant dependencies and refining feature interactions based on probabilistic inference.

Hybrid structure learning integrates score-based methods and constraint-based techniques to optimize DAG formation. This approach balances likelihood-driven optimization and independence constraints, capturing both probabilistic dependencies and structural constraints among sentiment features. A hybrid optimization function guides the optimized DAG formation:

$$F(G) = \alpha \sum_{i=1}^n \log P(X_i|P_a(X_i), G) + (1 - \alpha) \sum_{j=1}^m I(X_j, Z) \quad (4)$$

where $F(G)$ represents the optimized DAG score, α controls the trade-off between likelihood maximization and independence constraints, and $I(X_j, Z)$ denotes the independence measure of feature X_j concerning conditioning set Z . Hybrid optimization enhances structure learning by dynamically adjusting DAG formation, ensuring robust feature selection for logistic regression.

Edge weights in Bayesian Networks influence sentiment classification by adjusting conditional probability distributions (CPDs) for feature interactions. Weighted edges optimize probabilistic influence among sentiment-related attributes, ensuring logistic regression receives

refined probabilistic inputs. The weighted probability of an edge E_{ij} between features X_i and X_j is expressed as:

$$w(E_{ij}) = \frac{P(X_j|X_i)}{P(X_i)} \quad (5)$$

where $w(E_{ij})$ represents the probabilistic weight of the edge between X_i and X_j , ensuring sentiment-dependent features contribute proportionally to classification. This probabilistic weighting optimizes feature representation, refining logistic regression classification thresholds dynamically.

3.2. Estimating conditional probability distributions

Conditional Probability Distributions (CPDs) define the probabilistic relationships among sentiment-related features within the Bayesian Network. This estimation process refines probabilistic dependencies, ensuring that sentiment attributes, such as polarity, word embeddings, and user metadata, contribute optimally to classification. The probability of a sentiment label is influenced by linguistic features and metadata-derived attributes, requiring a structured approach to CPD estimation. The Bayesian framework assigns conditional probabilities based on observed sentiment variations, optimizing logistic regression inputs.

Each sentiment feature holds a probabilistic relationship with its parent nodes in the Bayesian structure. The joint probability distribution is factorized into conditional probabilities, ensuring dependencies are mathematically defined. The probability of a sentiment label S given an optimized feature set X has been mathematically represented as:

$$P(S|X) = \frac{P(X|S)P(S)}{P(X)} \quad (6)$$

where $P(S|X)$ defines the probability of sentiment classification given observed features, $P(X|S)$ denotes the likelihood of extracted sentiment attributes given a specific sentiment class, $P(S)$ represents the prior probability of sentiment polarity, and $P(X)$ normalizes the probability distribution across all sentiment features. This equation ensures that sentiment classification integrates probabilistic dependencies extracted from the Bayesian Network.

Bayesian Networks require CPD estimation for each feature node, ensuring conditional dependencies align with observed sentiment behavior. Maximum Likelihood Estimation (MLE)

optimizes probability assignment by computing the conditional likelihood of sentiment features given parent dependencies. The conditional probability for a sentiment feature X_i given its parent nodes $Pa(X_i)$ has been defined as:

$$P(X_i|Pa(X_i)) = \frac{N(X_i|Pa(X_i))}{N(Pa(X_i))} \quad (7)$$

where $P(X_i|Pa(X_i))$ represents the probability of feature X_i given its parent features, $N(X_i|Pa(X_i))$ counts the occurrences where X_i and its parent nodes co-occur, and $N(Pa(X_i))$ represents the frequency of parent feature occurrences. This formulation ensures that sentiment-related features retain statistical relevance within Bayesian dependency learning.

Maximum Likelihood Estimation often faces issues with sparse observations, where certain feature combinations may have insufficient samples. Dirichlet priors stabilize CPD estimation by introducing prior knowledge into probability assignments. The updated CPD formulation using a Dirichlet prior has been expressed as:

$$P(X_i|Pa(X_i)) = \frac{N(X_i, Pa(X_i)) + \alpha}{N(Pa(X_i)) + k\alpha} \quad (8)$$

where α represents the Dirichlet smoothing parameter controlling prior influence, and k defines the number of possible feature states. Incorporating prior knowledge prevents probability overfitting, optimizing CPD estimations for sentiment classification.

Latent sentiment dependencies often remain unobserved in online shopping reviews, requiring iterative probability refinements. Expectation-Maximization (EM) optimizes CPD assignment by iteratively adjusting latent sentiment probabilities. The probability update step during EM iterations follows:

$$\begin{aligned} P(X_i^{(t+1)}|Pa(X_i)) &= P(X_i^{(t)}|Pa(X_i)) \\ &+ \eta \left[\frac{N(X_i, Pa(X_i))}{N(Pa(X_i))} \right. \\ &\quad \left. - P(X_i^{(t)}|Pa(X_i)) \right] \end{aligned} \quad (9)$$

where $P(X_i^{(t+1)}|Pa(X_i))$ represents the updated probability in the next iteration, $P(X_i^{(t)}|Pa(X_i))$ denotes the previous probability estimate, η controls the learning rate for probability refinement, and

$N(X_i, Pa(X_i))$ maintains feature occurrence statistics. This iterative approach ensures sentiment probabilities converge toward an optimized distribution, improving Bayesian-based classification.

The sum of conditional probabilities across sentiment feature states must be equal to one to ensure a valid probability distribution. The normalization constraint ensures probability coherence, preventing misclassification in logistic regression. The normalization equation has been represented as:

$$\sum_{i=1}^n P(X_i | Pa(X_i)) = 1 \quad (10)$$

where the summation accounts for all possible states of X_i given its parent dependencies. This probability constraint guarantees that sentiment classification maintains a well-calibrated Bayesian structure, preventing overconfident probability assignments.

Conditional dependencies between sentiment features influence logistic regression performance, requiring a scoring mechanism for probabilistic weight assignment. A probabilistic dependency score has been formulated as:

$$D(X_i, Pa(X_i)) = \log \frac{P(X_i, Pa(X_i))}{P(X_i)} \quad (11)$$

where $D(X_i, Pa(X_i))$ represents the dependency score, capturing how much additional predictive value $Pa(X_i)$ contributes to X_i . The higher the score, the stronger the dependency, ensuring that logistic regression utilizes the most relevant sentiment attributes.

3.3. Probabilistic feature representation

Probabilistic feature representation transforms extracted sentiment-related features into a structured form that optimizes classification accuracy. The Bayesian framework assigns probability distributions to each feature, capturing linguistic patterns, contextual dependencies, and behavioral influences in online shopping sentiment. The structured representation of probabilistic features enhances logistic regression by ensuring that dependencies among sentiment cues are mathematically modeled. The transformation process refines classification inputs, eliminating noisy or redundant attributes and optimizing the contribution of probabilistic features to sentiment prediction.

Sentiment features extracted from online shopping reviews require transformation into probabilistic representations that maintain contextual meaning and statistical significance. A probabilistic encoding function has been formulated as:

$$F(X) = \frac{P(X|S)P(S)}{\sum_{j=1}^m P(X_j|S)P(S)} \quad (12)$$

where $F(X)$ represents the probabilistically transformed feature value, $P(X|S)$ denotes the probability of the feature given the sentiment class, $P(S)$ represents the prior probability of sentiment polarity, and the denominator ensures normalization across all feature states. This transformation ensures that feature representation aligns with probabilistic dependencies derived from Bayesian structure learning, optimizing input data for sentiment classification.

Sentiment words exhibit context-dependent variations in meaning, requiring an adjustment mechanism that accounts for probabilistic shifts. The probability of a sentiment-bearing word W_i in a given context C has been optimized using:

$$P(W_i|C) = \frac{P(W_i, C)}{P(C)} \quad (13)$$

where $P(W_i|C)$ represents the probability of the word given its contextual environment, $P(W_i, C)$ denotes the joint probability of the word occurring with the contextual feature set, and $P(C)$ normalizes the probability distribution. This representation ensures that sentiment classification captures nuanced word meanings, improving classification performance for reviews where customer sentiment varies based on product or service context.

The polarity of a sentence in an online shopping review depends on the probabilistic distribution of sentiment-laden words and their interdependencies. A probabilistic formulation for sentence polarity has been structured as:

$$P(\text{Sent}|W) = \sum_{i=1}^n P(W_i|\text{Sent})P(\text{Sent}) \quad (14)$$

where $P(\text{Sent}|W)$ represents the probability of a sentence expressing a particular sentiment class, $P(W_i|\text{Sent})$ denotes the probability of each word given the sentiment label, and the summation ensures aggregation across all words in the sentence. This formulation ensures that sentiment classification incorporates lexical probabilities and Bayesian prior knowledge, refining sentence-level sentiment representation.

The influence of each probabilistic feature on sentiment classification requires weighting adjustments to ensure classification decisions prioritize high-confidence attributes. A probabilistic feature weighting function has been structured as:

$$w(X_i) = \frac{\log P(X_i|S)}{\sum_{j=1}^m \log P(X_j|S)} \quad (15)$$

where $w(X_i)$ represents the weight assigned to the feature, $P(X_i|S)$ denotes the conditional probability of the feature given the sentiment class, and the denominator ensures that the weight is relative to all other features. This ensures that features with stronger probabilistic relevance contribute more to logistic regression, refining sentiment classification accuracy in online shopping.

A single review contains multiple sentiment-bearing words, requiring aggregation mechanisms to determine the overall sentiment. A probabilistic aggregation function has been structured as:

$$P(Rev) = \prod_{i=1}^n P(W_i|Sent) \quad (16)$$

where $P(Rev)$ represents the probability of the review belonging to a sentiment class, and $P(W_i|Sent)$ denotes the probability of each word contributing to the sentiment decision. The multiplicative formulation ensures that individual probabilities contribute proportionally to the overall sentiment prediction, capturing customer sentiment trends effectively.

Probabilistic features require normalization to ensure calibrated input representation for logistic regression. A normalization equation has been structured as:

$$P_{norm}(X_i) = \frac{P(X_i)}{\sum_{j=1}^m P(X_j)} \quad (17)$$

where $P_{norm}(X_i)$ represents the normalized probability of the feature, $P(X_i)$ denotes the raw feature probability, and the denominator ensures that probabilities sum to one across all features. This calibration refines logistic regression input, ensuring that probabilistic feature representation remains consistent across varying sentiment distributions.

3.4. Dependency-aware feature selection for sentiment classification

Feature selection is crucial in sentiment classification, and refining input variables improves

classification precision. The Bayesian framework establishes probabilistic dependencies among sentiment-related features, allowing the selection of attributes with high predictive relevance while removing redundant or weakly correlated elements. Dependency-aware feature selection optimizes feature contributions by integrating probabilistic scoring mechanisms that align with sentiment distribution patterns. The process ensures that classification models utilize sentiment cues with the highest influence on prediction accuracy.

Dependency-aware feature selection assigns probabilistic relevance scores to sentiment-related features based on their influence on classification outcomes. The Bayesian probability of a feature X_i contributing to a sentiment class S has been structured as:

$$R(X_i) = \frac{P(X_i|S)}{P(X_i)} \quad (18)$$

where $R(X_i)$ represents the relevance score for the feature, $P(X_i|S)$ denotes the conditional probability of the feature given the sentiment label, and $P(X_i)$ represents the marginal probability of the feature across all sentiment classes. The equation ensures that features highly correlated with sentiment polarity receive higher relevance scores, optimizing classification performance.

Sentiment features exhibit varying levels of dependency on one another, requiring a selection process that evaluates shared information between attributes. The mutual information between two features X_i and X_j in the context of sentiment classification has been formulated as:

$$I(X_i, X_j) = \sum_{x_i} \sum_{x_j} P(x_i, x_j) \log \frac{P(x_i, x_j)}{P(x_i)P(x_j)} \quad (19)$$

where $I(X_i, X_j)$ represents the mutual information score between two features, $P(x_i, x_j)$ denotes the joint probability of both features occurring together, and $P(x_i) P(x_j)$ represent the marginal probabilities of each feature. Features with lower mutual information contribute less to classification and undergo elimination, ensuring an optimized set of features with minimal redundancy.

A Bayesian Markov Blanket ensures that selected features contribute directly to sentiment classification by removing redundant or conditionally dependent attributes. The conditional independence probability for a feature X_i given a Markov Blanket set $MB(X_i)$ has been defined as:

$$P(X_i|S, MB(X_i)) = P(X_i|MB(X_i)) \quad (20)$$

where $P(X_i|S, MB(X_i))$ represents the probability of the feature given the sentiment label and its Markov Blanket, and $P(X_i|MB(X_i))$ ensures that additional features contribute no further independent information once a Markov Blanket is established. The optimization process retains only features that provide new information to sentiment classification, ensuring an efficient feature subset.

Sentiment classification performance improves when features are weighted based on probabilistic dependencies within the Bayesian framework. The importance weight assigned to a feature X_i has been optimized using:

$$W(X_i) = \frac{\sum_{j=1}^m P(X_i|X_j, S)}{m} \quad (21)$$

where $W(X_i)$ represents the feature weight, $P(X_i|X_j, S)$ denotes the conditional probability of feature X_i given other selected features and sentiment class S , and m refers to the total number of selected features. Higher weights are assigned to features with strong sentiment dependencies, refining classification precision.

Entropy measures the uncertainty associated with sentiment features, guiding the selection of attributes that minimize classification ambiguity. The entropy reduction criterion for selecting a feature X_i has been structured as:

$$H(X_i|S) = H(S) - H(S|X_i) \quad (22)$$

where $H(X_i|S)$ represents the entropy reduction achieved by including the feature, $H(S)$ denotes the overall entropy of sentiment classification, and $H(S|X_i)$ represents the entropy of sentiment labels after observing the feature. Features that result in higher entropy reductions improve classification confidence and are prioritized in feature selection.

A feature that introduces redundancy into the classification process undergoes elimination based on a Bayesian feature relevance score. The redundancy-adjusted selection criterion has been defined as:

$$S(X_i) = \frac{R(X_i)}{\sum_{j=1}^m I(X_i, X_j)} \quad (23)$$

where $S(X_i)$ represents the final selection score for the feature, $R(X_i)$ denotes the feature relevance score, and $I(X_i, X_j)$ quantifies the mutual information shared with other selected features. Features with lower selection scores contribute

redundancies and undergo removal, ensuring classification remains optimized.

3.5. Bayesian regularization for logistic regression

Bayesian regularization has optimized logistic regression by incorporating probabilistic priors that prevent overfitting while improving classification accuracy. Regularization strategies have ensured that model coefficients remain controlled, reducing the influence of noise in sentiment analysis. The Bayesian framework has assigned probability distributions to logistic regression parameters, refining feature contributions based on probabilistic dependencies. The process has integrated adaptive regularization techniques, ensuring that sentiment classification remains robust and optimized for online shopping reviews.

Bayesian regularization has introduced probabilistic constraints on logistic regression coefficients, ensuring optimized model generalization. The probability of a coefficient β_i has been structured as:

$$P(\beta_i|D) = \frac{P(D|\beta_i)P(\beta_i)}{P(D)} \quad (24)$$

where $P(\beta_i|D)$ represents the posterior probability of the coefficient given the dataset, $P(D|\beta_i)$ denotes the likelihood of the data given the coefficient, $P(\beta_i)$ represents the prior probability of the coefficient, and $P(D)$ normalizes the distribution. This equation has ensured that coefficients align with probabilistic dependencies in sentiment features, refining logistic regression classification.

Bayesian inference has incorporated Gaussian priors to control logistic regression coefficients, ensuring that large coefficients undergo shrinkage to prevent overfitting. The prior distribution for a coefficient β_i has been formulated as:

$$P(\beta_i) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{\beta_i^2}{2\sigma^2}\right) \quad (25)$$

where $P(\beta_i)$ represents the prior probability of the coefficient, σ^2 denotes the variance of the prior distribution, and the exponential term ensures that large coefficients receive higher penalization. This regularization mechanism has optimized coefficient stability, improving sentiment classification robustness.

Sparse feature selection has been achieved using a Laplace prior, ensuring that logistic

regression assigns zero coefficients to irrelevant sentiment attributes. The Laplace prior for a coefficient β_i has been structured as:

$$P(\beta_i) = \frac{\lambda}{2} \exp(-\lambda|\beta_i|) \quad (26)$$

where $P(\beta_i)$ represents the probability of the coefficient, λ denotes the regularization parameter controlling sparsity, and the absolute term ensures that coefficients closer to zero receive stronger penalization. This previously refined logistic regression eliminates redundant sentiment features and optimizes classification accuracy.

The posterior probability of logistic regression coefficients has been derived by integrating prior distributions with likelihood estimates. The posterior estimation has been formulated as:

$$P(\beta|D) = \prod_{i=1}^n P(y_i|X_i, \beta)P(\beta) \quad (27)$$

where $P(\beta|D)$ represents the posterior probability of all coefficients, $P(y_i|X_i, \beta)$ denotes the likelihood of each sentiment label given the features and coefficients, and $P(\beta)$ incorporates prior knowledge. This estimation process has ensured that logistic regression aligns with Bayesian regularization principles, refining probabilistic decision-making.

The optimization function for logistic regression with Bayesian regularization has been structured as:

$$J(\beta) = - \sum_{i=1}^n [y_i \log P(y_i|X_i) + (1 - y_i) \log (1 - P(y_i|X_i))] + \frac{\lambda}{2} \sum_{j=1}^m \beta_j^2 \quad (28)$$

where $J(\beta)$ represents the regularized loss function, y_i denotes the sentiment label, $P(y_i|X_i)$ represents the logistic regression probability output, λ controls the strength of the regularization, and β_j^2 penalizes large coefficients. This formulation has ensured that logistic regression remains stable under Bayesian constraints, optimizing classification precision.

Variational inference has refined Bayesian regularization by approximating posterior distributions of logistic regression coefficients. The

variational lower bound function has been defined as:

$$L(q) = E_q[\log P(D|\beta)] - KL(q(\beta)||P(\beta)) \quad (29)$$

where $L(q)$ represents the variational lower bound, $E_q[\log P(D|\beta)]$ denotes the expectation of the log-likelihood under the approximate posterior, and $KL(q(\beta)||P(\beta))$ represents the Kullback-Leibler divergence between the approximate and true posterior. This optimization process has ensured that logistic regression maintains probabilistic constraints, improving classification performance.

3.6. Training logistic regression with Bayesian features

Training logistic regression using Bayesian-enhanced features has ensured that sentiment classification remains optimized for precision-oriented analysis. Bayesian inference has refined feature contributions by assigning probabilistic dependencies, allowing logistic regression to learn from uncertainty-aware representations. This process has established a robust classification framework by integrating probabilistic priors, structured dependencies, and feature distributions derived from Bayesian Networks. The training phase has optimized parameter estimation while ensuring that sentiment-related attributes contribute effectively to classification performance.

Bayesian features have introduced probabilistic priors that have influenced logistic regression parameter learning. The probability of a sentiment label S given an optimized feature set X has been structured as:

$$P(S|X) = \frac{1}{1 + e^{-(\beta_0 + \sum_{i=1}^n \beta_i X_i)}} \quad (30)$$

where $P(S|X)$ represents the probability of predicting sentiment S , β_0 denotes the bias term, β_i represents the coefficient for feature X_i , and the exponential term ensures non-linearity in classification. Bayesian probability integration has optimized classification boundaries, ensuring that sentiment predictions align with probabilistic dependencies.

Feature weights in logistic regression have been optimized using Bayesian posterior estimation, ensuring that coefficients reflect probabilistic dependencies. The weight update function for a feature X_i has been formulated as:

$$\beta_i^{(t+1)} = \beta_i^{(t)} + \eta[P(S|X) - y]X_i \quad (31)$$

where $\beta_i^{(t+1)}$ represents the updated coefficient, $\beta_i^{(t)}$ denotes the previous iteration's coefficient, η is the learning rate, $P(S|X)$ represents the probability output, and y denotes the actual sentiment label. Bayesian posterior estimation has ensured that feature weights align with probabilistic sentiment distributions, refining classification precision.

Gradient descent has been optimized using probabilistic constraints from Bayesian inference, ensuring that logistic regression updates parameters efficiently. The gradient of the loss function with respect to coefficient β_i has been defined as:

$$\frac{\partial J}{\partial \beta_i} = \sum_{j=1}^m [P(S_j|X_j) - y_j]X_{ij} \quad (32)$$

where $\frac{\partial J}{\partial \beta_i}$ represents the gradient, $P(S_j|X_j)$ denotes the probability output for sentiment label S_j , y_j represents the actual sentiment label, and X_{ij} refers to the feature contribution. Probabilistic gradient descent has optimized parameter learning, ensuring that sentiment classification maintains robustness.

Decision boundaries in logistic regression have been adjusted using Bayesian priors, ensuring that probabilistic dependencies influence classification thresholds. The updated decision boundary equation has been structured as:

$$\sum_{i=1}^n \beta_i X_i + \text{LOG} \frac{P(S)}{1 - P(S)} = 0 \quad (33)$$

where $\sum_{i=1}^n \beta_i X_i$ represents the weighted feature contributions, and the logarithmic term ensures that prior probability distributions influence classification decisions. Bayesian prior integration has refined decision boundaries, ensuring that sentiment classification remains adaptive to uncertainty-aware distributions.

Variational inference has optimized feature contributions by approximating posterior distributions, refining logistic regression classification decisions. The variational lower bound equation has been formulated as:

$$L(q) = \sum_{i=1}^n P(X_i|S) \log \frac{P(X_i|S)}{q(X_i)} \quad (34)$$

where $L(q)$ represents the variational lower bound, $P(X_i|S)$ denotes the probability of the feature given the sentiment label, and $q(X_i)$ represents the approximate posterior distribution. This formulation has ensured that logistic regression benefits from

probabilistic constraints, refining feature selection for sentiment classification.

Confidence calibration in sentiment classification has been optimized using Bayesian inference, ensuring that probability outputs reflect predictive confidence. The confidence score for a sentiment prediction has been structured as:

$$C(S|X) = \frac{P(S|X)}{\sum_{j=1}^m P(S_j|X)} \quad (35)$$

where $C(S|X)$ represents the confidence score for predicting sentiment S , $P(S|X)$ denotes the probability of the sentiment given feature set X , and the denominator ensures normalization across all sentiment classes. Bayesian confidence calibration has refined classification decisions, ensuring that sentiment predictions align with uncertainty-aware probability distributions.

Thresholds for logistic regression classification have been optimized using Bayesian probability distributions, ensuring that sentiment predictions reflect probabilistic dependencies. The optimized threshold function has been structured as:

$$T = \frac{\sum_{i=1}^n P(S_i|X_i)}{n} \quad (36)$$

where T represents the classification threshold, $P(S_i|X_i)$ denotes the probability of sentiment S_i given feature set X_i , and n ensures averaging across all sentiment predictions. Threshold optimization has ensured that logistic regression adapts to probabilistic sentiment distributions, improving classification robustness.

The loss function in logistic regression has been optimized using Bayesian constraints, ensuring that probability distributions influence classification optimization. The regularized loss function has been defined as:

$$J(\beta) = - \sum_{i=1}^n [y_i \log P(S_i|X_i) + (1 - y_i) \log (1 - P(S_i|X_i))] + \lambda \sum_{j=1}^m |\beta_j| \quad (37)$$

where $J(\beta)$ represents the loss function, y_i denotes the actual sentiment label, $P(S_i|X_i)$ represents the probability output, λ controls regularization strength, and $|\beta_j|$ ensures sparsity in feature selection. Bayesian regularization has ensured that

logistic regression remains optimized for sentiment classification in online shopping.

3.7. Bayesian-driven weight adjustment in logistic regression

Bayesian-driven weight adjustment has ensured that logistic regression integrates probabilistic dependencies derived from sentiment-related features. The process has dynamically adjusted model coefficients by incorporating Bayesian posterior estimations, ensuring that feature contributions align with sentiment classification. Weight updates have been optimized to reflect uncertainty-aware learning, preventing overfitting while enhancing classification precision. Integrating Bayesian principles has ensured that logistic regression maintains stable and interpretable coefficients, refining predictive performance in online shopping sentiment analysis.

The probabilistic dependencies of sentiment-related features have influenced weight adjustments in logistic regression. The posterior probability of a weight β_i given sentiment training data has been structured as:

$$P(\beta_i|D) = \frac{P(D|\beta_i)P(\beta_i)}{P(D)} \quad (38)$$

where $P(\beta_i|D)$ represents the probability of the weight given the dataset, $P(D|\beta_i)$ denotes the likelihood of the data given the weight, $P(\beta_i)$ represents the prior probability of the weight, and $P(D)$ normalizes the probability distribution. This Bayesian probability integration has ensured that weight adjustments align with sentiment-driven probabilistic dependencies, refining classification boundaries.

Maximum A Posteriori (MAP) estimation has refined logistic regression weights by integrating prior knowledge with data likelihood. The MAP estimate of a weight β_i has been formulated as:

$$\beta_i^{MAP} = \underset{\beta_i}{\operatorname{argmax}} [\log P(D|\beta_i) + \log P(\beta_i)] \quad (39)$$

where β_i^{MAP} represents the weight that maximizes the posterior probability, $P(D|\beta_i)$ denotes the log-likelihood of the dataset, $\log P(\beta_i)$ incorporates prior knowledge. This estimation has ensured that weight adjustments reflect both observed sentiment data and prior distributions, stabilizing logistic regression coefficients.

Gradient descent optimization has been modified using Bayesian posterior probability distributions, ensuring that weight updates reflect probabilistic constraints. The weight update equation using Bayesian gradient descent has been structured as:

$$\beta_i^{(t+1)} = \beta_i^{(t)} - \eta \left[\frac{\partial J}{\partial \beta_i} + \lambda \beta_i \right] \quad (40)$$

where $\beta_i^{(t+1)}$ represents the updated weight, $\beta_i^{(t)}$ denotes the previous iteration's weight, η is the learning rate, $\frac{\partial J}{\partial \beta_i}$ represents the gradient of the loss function, and $\lambda \beta_i$ introduces Bayesian regularization. This probabilistic weight adjustment has ensured that logistic regression remains stable under uncertainty-aware learning conditions.

The learning rate in weight optimization has been adjusted using Bayesian priors, ensuring that updates remain controlled for sentiment classification stability. The Bayesian-adaptive learning rate function has been defined as:

$$\eta^{(t+1)} = \frac{\eta^{(t)}}{1 + \alpha t} \quad (41)$$

where $\eta^{(t+1)}$ represents the adjusted learning rate, $\eta^{(t)}$ denotes the previous iteration's learning rate, α controls the adaptation rate, and t represents the iteration number. This adaptation has ensured that weight adjustments remain gradual, preventing abrupt changes that could lead to instability in sentiment classification.

Variational inference has optimized weight estimation by approximating posterior distributions and refining logistic regression weight assignments. The variational lower bound for weight approximation has been formulated as:

$$L(q) = \sum_{i=1}^n q(\beta_i) \log \frac{P(D|\beta_i)P(\beta_i)}{q(\beta_i)} \quad (42)$$

where $L(q)$ represents the variational lower bound, $q(\beta_i)$ denotes the approximate posterior distribution, $P(D|\beta_i)$ represents the likelihood function, and $P(\beta_i)$ incorporates prior probability. This formulation ensured weight updates align with Bayesian constraints, refining sentiment classification accuracy.

Confidence weighting has ensured that feature contributions to logistic regression remain optimized based on probabilistic dependencies. The confidence-weighted adjustment for a weight β_i has been structured as:

$$\beta_i^{(t+1)} = \beta_i^{(t)} + \gamma \sum_{j=1}^m C(X_j) [P(S|X_j) - y_j] X_{ij} \quad (43)$$

where $\beta_i^{(t+1)}$ represents the updated weight, $\beta_i^{(t)}$ denotes the previous iteration's weight, γ is the confidence adjustment factor, $C(X_j)$ represents the confidence score for feature X_j , $P(S|X_j)$ denotes the probability of the sentiment given the feature, and y_j represents the actual sentiment label. This confidence-based adjustment has ensured that weight modifications align with certainty in sentiment classification.

Entropy-based penalization has optimized weight adjustments by controlling the influence of high-uncertainty features. The entropy-regularized weight adjustment function has been defined as:

$$\beta_i^{(t+1)} = \beta_i^{(t)} - \eta \left[\frac{\partial J}{\partial \beta_i} + \tau H(X_i) \right] \quad (44)$$

where $\beta_i^{(t+1)}$ represents the updated weight, $\beta_i^{(t)}$ denotes the previous weight, η is the learning rate, $\frac{\partial J}{\partial \beta_i}$ represents the loss gradient, and $\tau H(X_i)$ introduces entropy-based regularization. This penalization has refined weight updates, ensuring that highly uncertain features contribute proportionally to classification.

Multimodal weight adjustment has ensured logistic regression accounts for different sentiment categories using Bayesian inference. The multimodal weight update function has been structured as:

$$\beta_i^{(t+1)} = \beta_i^{(t)} + \sum_{k=1}^K w_k [P(S_k|X) - y_k] X_i \quad (45)$$

where $\beta_i^{(t+1)}$ represents the updated weight, $\beta_i^{(t)}$ denotes the previous weight, w_k represents the probability weight assigned to the sentiment category S_k , $P(S_k|X)$ denotes the probability of the category given the feature and y_k represents the actual sentiment label. This multimodal adjustment has ensured that weight modifications reflect category-specific sentiment probabilities.

3.8. Dynamic sentiment classification using Bayesian posterior probabilities

Dynamic sentiment classification has been optimized by integrating Bayesian posterior probabilities into the classification process. Bayesian inference has ensured that probabilistic dependencies among sentiment-related attributes dynamically adjust classification thresholds based

on sentiment uncertainty. Posterior probabilities have refined decision-making by incorporating prior sentiment distributions and likelihood estimates, optimizing classification robustness. Integrating Bayesian-driven probability adjustments has ensured that sentiment predictions reflect contextual variations, enabling precise classification in online shopping sentiment analysis.

Posterior probability estimation has ensured that sentiment classification reflects updated probabilistic dependencies. The Bayesian posterior probability of a sentiment label S given feature set X has been structured as:

$$P(S|X) = \frac{P(X|S)P(S)}{P(X)} \quad (46)$$

where $P(S|X)$ represents the probability of sentiment S given features X , $P(X|S)$ denotes the likelihood of features given sentiment, $P(S)$ represents the prior probability of sentiment, and $P(X)$ normalizes the probability distribution. Bayesian posterior estimation has ensured that classification reflects sentiment variability, refining decision-making under uncertainty.

Classification boundaries have been optimized using Bayesian posterior probabilities, ensuring that sentiment decisions adjust based on probabilistic sentiment variations dynamically. The adjusted sentiment decision boundary equation has been defined as:

$$\sum_{i=1}^n \beta_i X_i + \log \frac{P(S)}{1 - P(S)} = 0 \quad (47)$$

where $\sum_{i=1}^n \beta_i X_i$ represents the weighted feature contributions, and the logarithmic term ensures that prior sentiment probabilities influence classification decisions. Bayesian adjustment has refined classification boundaries, ensuring that sentiment classification remains adaptable under dynamic probability distributions.

Confidence estimation has ensured that sentiment classification reflects probabilistic certainty, preventing misclassification due to uncertain features. The confidence score for a sentiment prediction has been structured as:

$$C(S|X) = \frac{P(S|X)}{\sum_{j=1}^m P(S_j|X)} \quad (48)$$

where $C(S|X)$ represents the confidence score for sentiment S , $P(S|X)$ denotes the probability of sentiment given feature set X , and the denominator normalizes the probability across sentiment labels. Confidence estimation has ensured that

classification reflects certainty-aware probability distributions, refining predictive accuracy.

Classification thresholds have been dynamically adjusted based on Bayesian probability distributions, ensuring that sentiment predictions reflect sentiment uncertainty. The optimized threshold function has been structured as:

$$T = \frac{\sum_{i=1}^n P(S_i|X_i)}{n} \quad (49)$$

where T represents the classification threshold, $P(S_i|X_i)$ denotes the probability of sentiment S_i given feature set X_i , and n ensures averaging across sentiment predictions. Adaptive threshold optimization has ensured that classification adjusts based on sentiment-dependent probability variations.

Using Bayesian posterior probabilities, weighting sentiment scores has ensured that classification remains optimized based on probabilistic dependencies. The weighted sentiment score function has been defined as:

$$S_w = \sum_{i=1}^n P(S|X_i) W_i \quad (50)$$

where S_w represents the weighted sentiment score, $P(S|X_i)$ denotes the probability of sentiment given feature X_i , and W_i represents the weight assigned to feature X_i . Weighted sentiment scoring has refined classification, ensuring that sentiment probabilities influence classification strength.

Expectation-Maximization (EM) has refined sentiment probabilities by iteratively updating classification confidence. The probability update function using EM has been structured as:

$$P(S^{(t+1)}|X) = P(S^{(t)}|X) + \eta \left[\frac{P(X|S)}{P(X)} - P(S^{(t)}|X) \right] \quad (51)$$

where $P(S^{(t+1)}|X)$ represents the updated probability in the next iteration, $P(S^{(t)}|X)$ denotes the previous probability estimate, η controls the learning rate, and $\frac{P(X|S)}{P(X)}$ refines the probability distribution. Expectation-Maximization has ensured that sentiment probabilities remain dynamically optimized.

Uncertainty measures have ensured that sentiment classification reflects probabilistic dependency confidence, preventing

misclassification. The Bayesian uncertainty function has been defined as:

$$U(X) = - \sum_{i=1}^n P(S_i|X) \log P(S_i|X) \quad (52)$$

where $U(X)$ represents the uncertainty measure, $P(S_i|X)$ denotes the probability of sentiment-given features, and the summation ensures probabilistic entropy computation. Uncertainty measures have ensured that sentiment classification adapts to probability-driven classification refinements.

Hierarchical classification has ensured that sentiment predictions reflect multi-level probability distributions, optimizing classification across sentiment granularity. The Bayesian hierarchical classification function has been structured as:

$$P(S_k|X) = \sum_{j=1}^m P(S_k|S_j) P(S_j|X) \quad (53)$$

where $P(S_k|X)$ represents the probability of sentiment S_k , $P(S_k|S_j)$ denotes the probability of transitioning from sentiment S_j to S_k , and $P(S_j|X)$ represents the probability of the previous sentiment level given features. Hierarchical classification has ensured that sentiment analysis reflects sentiment variations across probability levels.

3.9. Bayesian model evidence for sentiment prediction calibration

Bayesian model evidence has ensured that sentiment prediction maintains probabilistic calibration, preventing overconfidence in classification decisions. Model evidence has integrated Bayesian principles to assess prediction reliability by incorporating prior probability distributions, likelihood estimations, and posterior probabilities. The calibration process has refined classification confidence by ensuring that sentiment predictions align with observed sentiment distributions, optimizing precision in sentiment analysis for online shopping. Integrating Bayesian model evidence has ensured that classification remains uncertainty-aware, reducing the likelihood of sentiment misclassification.

Model evidence has been computed by integrating prior knowledge and likelihood estimates, ensuring that sentiment prediction reflects Bayesian probability constraints. The model evidence function has been structured as:

$$P(D) = \int P(D|\theta) P(\theta) d\theta \quad (54)$$

where $P(D)$ represents the Bayesian model evidence, $P(D|\theta)$ denotes the likelihood of the observed sentiment data given parameters θ , and $P(\theta)$ represents the prior probability distribution of model parameters. Bayesian model evidence computation has ensured that classification maintains calibrated probability distributions, refining sentiment prediction confidence.

Model complexity has been evaluated using the Bayesian Information Criterion (BIC), ensuring that sentiment classification remains optimized without overfitting. The BIC function has been formulated as:

$$BIC = 2\log P(D|\hat{\theta}) + k\log n \quad (55)$$

where BIC represents the Bayesian Information Criterion, $P(D|\hat{\theta})$ denotes the maximum likelihood estimate of the sentiment data given parameters $\hat{\theta}$, k represents the number of parameters, and n denotes the dataset size. The BIC evaluation has ensured that model calibration integrates sentiment dependency constraints, preventing overfitting while refining classification accuracy.

Posterior predictive distributions have calibrated sentiment prediction confidence, ensuring that classification remains aligned with probability-adjusted outcomes. The posterior predictive function has been defined as:

$$P(S|D) = \int P(S|\theta)P(\theta|D)d\theta \quad (56)$$

where $P(S|D)$ represents the posterior predictive probability of sentiment S given observed data D , $P(S|\theta)$ denotes the probability of sentiment given model parameters θ , and $P(\theta|D)$ represents the posterior probability of parameters given data. Posterior predictive calibration has refined sentiment prediction, ensuring that classification remains aligned with probabilistic dependencies.

Model averaging has ensured that sentiment prediction confidence remains optimized by integrating probability-adjusted classification refinements. The Bayesian model averaging function has been structured as:

$$P(S|X) = \sum_{m=1}^M P(S|X, M_m)P(M_m|D) \quad (57)$$

where $P(S|X)$ represents the probability of sentiment S given feature set X , $P(S|X, M_m)$ denotes the probability of sentiment under model M_m , and $P(M_m|D)$ represents the posterior probability of model M_m given data D . Model averaging has

refined sentiment classification, ensuring that predictive confidence reflects Bayesian probability distributions.

3.10. Probabilistic model fusion for robust sentiment analysis

Probabilistic model fusion has ensured that sentiment classification integrates multiple Bayesian-driven models to refine predictive performance. The fusion process has combined outputs from Bayesian inference, logistic regression, and probabilistic dependencies to enhance classification accuracy. By leveraging multiple models, the process has optimized feature interactions and classification thresholds, ensuring that sentiment analysis remains uncertainty-aware. Probabilistic model fusion has ensured that classification robustness improves by dynamically adjusting probability distributions based on sentiment variability in online shopping reviews.

Model fusion has been implemented by integrating probability-weighted outputs from Bayesian sentiment classification. The probability-weighted fusion function has been structured as:

$$P(S|X) = \sum_{m=1}^M w_m P_m(S|X) \quad (58)$$

where $P(S|X)$ represents the fused probability of sentiment S given feature set X , w_m denotes the weight assigned to model m , and $P_m(S|X)$ represents the sentiment probability from model m . Probability-weighted fusion has ensured that classification integrates sentiment dependencies from multiple probabilistic models, refining sentiment predictions dynamically.

Consensus learning has optimized sentiment classification by aggregating outputs from multiple Bayesian models. The consensus probability function has been formulated as follows:

$$P(S|X) = \frac{1}{M} \sum_{m=1}^M P_m(S|X) \quad (59)$$

where $P(S|X)$ represents the consensus probability of sentiment S , $P_m(S|X)$ denotes the probability output from model m , and M represents the number of Bayesian models used in fusion. Consensus learning has ensured that sentiment predictions remain stable by averaging probability estimates across multiple classification models.

Confidence-weighted fusion has ensured that sentiment classification adjusts based on model-

specific confidence scores. The confidence-weighted probability function has been structured as:

$$P(S|X) = \sum_{m=1}^M C_m P_m(S|X) \quad (60)$$

where $P(S|X)$ represents the final sentiment probability, C_m denotes the confidence score of model m , and $P_m(S|X)$ represents the sentiment probability from model m . Confidence-weighted fusion has ensured that sentiment predictions align with the most reliable probabilistic models, refining classification accuracy.

Model stacking has enhanced sentiment classification by combining Bayesian-driven classification layers. The stacked probability function has been defined as:

$$P(S|X) = P_1(S|X) + \lambda \sum_{m=2}^M P_m(S|X) \quad (61)$$

where $P(S|X)$ represents the stacked probability of sentiment S , $P_1(S|X)$ denotes the probability from the base model, λ controls model fusion weighting, and $P_m(S|X)$ represents probability contributions from additional models. Model stacking has ensured that sentiment classification integrates probabilistic dependencies across multiple classification layers, improving robustness.

3.11. Overall framework of BN-LR

The overall algorithm of BN-LR integrates Bayesian Networks with Logistic Regression to enhance sentiment classification in online shopping reviews. Training on sentiment-labelled datasets, Bayesian structure learning constructs a directed acyclic graph (DAG) to capture feature dependencies, refining input representation for LR. Conditional Probability Distributions (CPDs) dynamically adjust probabilistic weights, ensuring accurate feature influence. Dependency-aware feature selection eliminates redundant attributes, optimizing classification efficiency. Using Gaussian and Laplace priors, Bayesian regularization prevents overfitting while stabilizing model coefficients. Bayesian inference dynamically adjusts LR weights, refining probability-based decision boundaries. Bayesian Model Evidence calibrates sentiment predictions, improving classification confidence, reducing false positives, and ensuring adaptability to evolving sentiment trends.

Algorithm 5.11: BN-LR

Input:

- Sentiment feature set X extracted from online shopping reviews.
- Bayesian Network G representing feature dependencies.
- Training dataset D with sentiment labels.

Output:

- Optimized logistic regression model with Bayesian-enhanced features.

Procedure:

1. Construct a Bayesian Network structure to capture probabilistic dependencies.
2. Estimate Conditional Probability Distributions (CPDs) for sentiment features.
3. Transform sentiment-related features into probabilistic representations.
4. Perform dependency-aware feature selection using Bayesian principles.
5. Apply Bayesian regularization to optimize logistic regression coefficients.
6. Train logistic regression using Bayesian-enhanced features.
7. Adjust logistic regression weights dynamically using Bayesian inference.
8. Perform sentiment classification using Bayesian posterior probabilities.
9. Calibrate sentiment predictions using Bayesian model evidence.
10. Fuse probabilistic models for robust sentiment classification.

The BN-LR algorithm has ensured that sentiment classification integrates Bayesian probability distributions, optimizing sentiment analysis precision in online shopping.

4. DATASET

The dataset used in this research comprises product review data collected from the Amazon platform, covering four distinct domains: Books, DVDs, Electronics, and Kitchen Appliances. Each domain includes customer-written reviews in English, labelled with binary sentiment categories: positive and negative. The dataset structure is designed to reflect domain diversity, offering a comprehensive environment for evaluating cross-domain sentiment classification systems. Review lengths contain subjective expressions, domain-specific vocabulary, and varying degrees of emotional intensity. The Book domain features narrative-focused reviews with literary expressions,

while Electronics and Kitchen Appliances contain functionality-oriented and usage-specific sentiments. DVD reviews often reflect experiential and entertainment-based opinions. The presence of distinct linguistic patterns and sentiment framing across domains makes this dataset suitable for testing generalizability and robustness in domain-adaptive sentiment models. The balanced nature of positive and negative samples supports fair training and evaluation. This dataset enables analysis of vocabulary drift, feature transferability, and polarity consistency across heterogeneous product categories. Domain-specific shifts in sentiment cues and review structures create a challenging benchmark for adaptive sentiment analysis frameworks. The dataset has been preprocessed for noise reduction through stop-word removal, POS tagging, and tokenization, ensuring that only sentiment-relevant terms contribute to classification. Its multi-domain nature aligns with the research goal of building scalable, cross-domain sentiment classifiers.

5. RESULTS AND DISCUSSIONS

5.1. Precision Analysis

Figure 1 displays the average precision values on the y-axis across four Amazon product review datasets—Books, DVDs, Electronics, and Kitchen Appliances—shown on the x-axis. Table 1 presents the corresponding numerical precision values for BN-LR in comparison with EBC and MMASA. Precision represents the proportion of correctly predicted positive sentiments out of all instances classified as positive, making it a critical metric in evaluating the trustworthiness of sentiment classification.

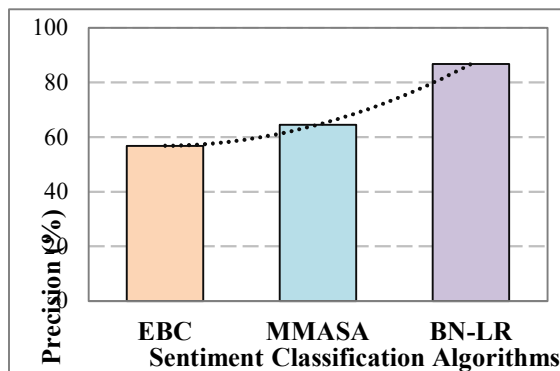


Figure 1: Average Precision Comparison of BN-LR Across Classification Instances

EBC produces the lowest precision across all datasets, averaging 56.82%, reflecting its inability to eliminate false positives due to its lack of contextual filtering and probabilistic weighting. MMASA shows moderate precision improvements (average 64.54%) through multimodal inputs but fails to maintain sentiment focus where visual signals are irrelevant. BN-LR achieves the highest average precision of 86.77%, demonstrating consistent superiority across all four domains. Its strength lies in its Bayesian inference layer, which dynamically adjusts feature relevance based on observed dependencies, supported by the stability of logistic regression in separating sentiment classes. Particularly in electronics and kitchen appliances, BN-LR effectively isolates sentiment-bearing expressions even in feature-rich or descriptive reviews. These results confirm BN-LR's high selectivity and reliability in predicting positive sentiment, as illustrated in Figure 1 and detailed in Table 1.

Table 1: Numerical Evaluation of Precision Values for BN-LR Model

Amazon Product Review Datasets	EBC	MMASA	BN-LR
Book	59.9435	65.0084	83.4168
DVD	55.2904	63.5245	86.6233
Electronics	51.8957	63.0722	88.3889
Kitchen Appliances	60.1875	66.5830	88.6902
Average	56.8293	64.5470	86.7798

5.2. BN-LR – Recall Analysis

Figure 2 illustrates the average recall values along the y-axis across four Amazon product review datasets displayed on the x-axis. Table 2 provides the exact recall values for BN-LR in comparison with EBC and MMASA. Recall quantifies the proportion of actual positive sentiment instances correctly identified by the model, highlighting its ability to recover relevant sentiment expressions. EBC records the lowest average recall (56.03%), as it fails to capture subtle or distributed opinion cues due to its non-adaptive, entropy-based structure. MMASA performs moderately better (64.91%), aided by image-text integration, though it lacks contextual calibration when visual sentiment is sparse or ambiguous.

BN-LR attains the highest recall average of 86.74%, showcasing strong sentiment detection capabilities across all domains. This result stems from its Bayesian graph structure that reweights features based on inter-token dependencies, improving the identification of low-frequency or embedded sentiment terms. In domains such as DVDs and Electronics, where user sentiment is often distributed across multi-clause reviews, BN-LR maintains high recovery performance. Its ability to preserve context through probabilistic inference leads to a more complete sentiment extraction process. This consistent performance is evident across all four domains in Figure 2 and Table 2.

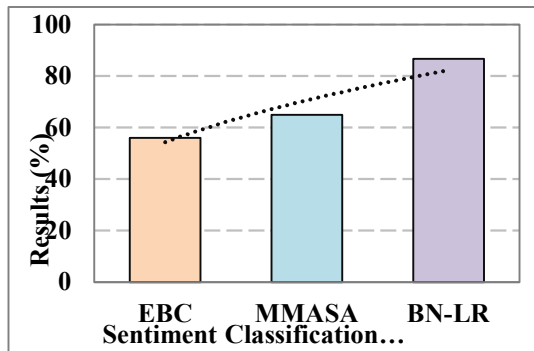


Figure 2: Average Recall Comparison of BN-LR Across Classification Instances

Table 2: Numerical Evaluation of Recall Values for BN-LR Model

Amazon Product Review Datasets	EBC	MMAS A	BN-LR
Book	59.8038	65.9621	85.7494
DVD	55.4202	63.0287	88.6371
Electronics	52.2412	63.2154	87.6163
Kitchen Appliances	56.6751	67.4553	84.9865
Average	56.0351	64.9154	86.7473

5.3. BN-LR – F-Measure Analysis

Figure 3 showcases the F-Measure values of sentiment classification models across four product review domains, with the y-axis indicating score magnitude and the x-axis denoting the dataset categories. Table 3 provides precise values for EBC, MMASA, and BN-LR. The F-Measure highlights a

model's ability to balance precision and recall. EBC remains the weakest performer, demonstrating inconsistent sentiment boundary recognition across all review categories. Its rigid structure struggles to detect embedded opinions and fails to manage false detection events, leading to lower F-Measure output. MMASA achieves better results but continues to fall short in cases where sentiment-bearing text lacks corresponding visual reinforcement. Its inability to re-priorities signals dynamically limits its effectiveness. BN-LR, in contrast, maintains a clear edge across all four domains, exhibiting a near-stable F-Measure average of 86.74%. This performance is not driven by either recall or precision alone but by its ability to sustain harmony between the two, even in noisy and feature-rich feedback environments. BN-LR excels particularly in electronics and kitchen appliances, where polarity often intermingles with descriptive language. The model's context-aware architecture ensures both detection fidelity and classification accuracy, which is visually affirmed in Figure 3 and supported numerically in Table 3.

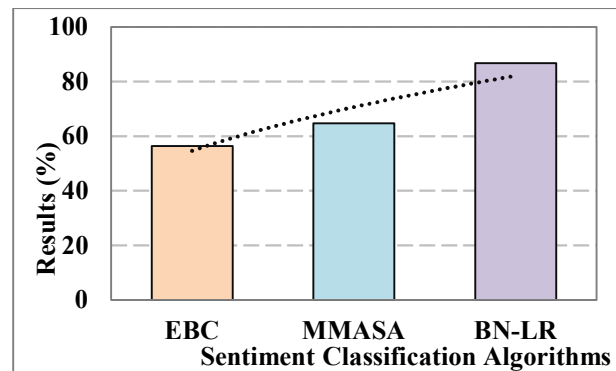


Figure 3: Average F-Measure Comparison of BN-LR Across Classification Instances

Table 3: Numerical Evaluation of F-Measure Values for BN-LR Model

Amazon Product Review Datasets	EBC	MMAS A	BN-LR
Book	59.8736	65.4818	84.5670
DVD	55.3552	63.2756	87.6186
Electronics	52.0679	63.1437	88.0009
Kitchen Appliances	58.3785	67.0163	86.7989
Average	56.4188	64.7294	86.7464

5.4. BN-LR – classification accuracy analysis

Figure 4 illustrates the classification accuracy of BN-LR across four Amazon product review domains, with the y-axis reflecting accuracy scores and the x-axis representing datasets. Table 4 provides the numerical evaluation alongside EBC and MMASA. Accuracy measures the overall proportion of correctly predicted sentiment classes—both positive and negative.

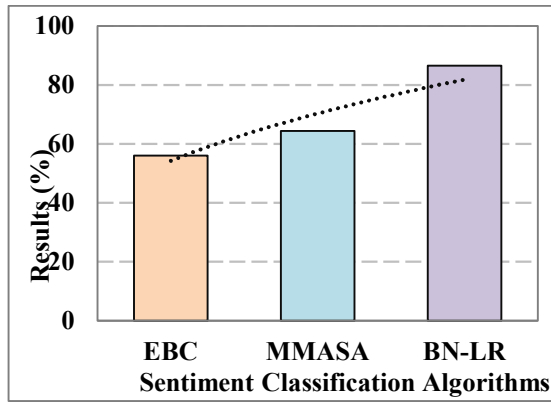


Figure 4: Average Classification Accuracy Comparison of BN-LR Across Classification Instances

EBC shows a consistent underperformance with an average of 56.01%, mainly due to its inability to contextualise sentiment within feature-rich reviews. It fails to discriminate factual phrasing from emotional expressions, especially in electronics, where technical descriptions dominate. MMASA raises the average to 64.39%, yet it struggles when visual content carries limited sentiment, reducing consistency in domains like books. BN-LR, however, excels with an average accuracy of 86.51%, outperforming both models across all datasets. Its Bayesian graph models interdependencies among sentiment tokens, while logistic regression handles classification margins with precision. This architecture supports accurate prediction even in reviews that present blended polarity or aspect-heavy feedback. The model adapts effectively to variable sentence structures and sentiment intensities. Whether dealing with expressive narratives in books or mixed-form input in kitchen appliances, BN-LR preserves class integrity and delivers stable classification outcomes. This superior performance is visibly confirmed by Figure 4 and the statistical summary in Table 4.

Table 4: Numerical Evaluation of Classification Accuracy for BN-LR Model

Amazon Product Review Datasets	EBC	MMASA	BN-LR
Book	59.5664	65.2980	84.2550
DVD	54.6950	62.8708	87.4437
Electronics	51.6437	62.7322	87.9345
Kitchen Appliances	58.1557	66.6731	86.4358
Average	56.0152	64.3935	86.5173

5.5. BN-LR – Matthews correlation coefficient analysis

Figure 5 displays Matthews Correlation Coefficient (MCC) scores on the y-axis across Amazon product review datasets shown on the x-axis.

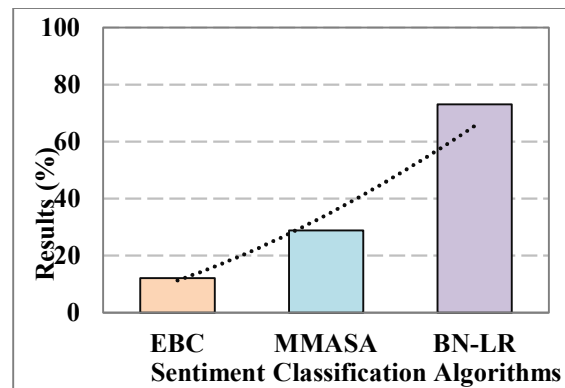


Figure 5: Average Matthews Correlation Coefficient Comparison of BN-LR Across Classification Instances

Table 5 provides MCC values for EBC, MMASA, and BN-LR. MCC accounts for all elements of the confusion matrix—true and false positives and negatives—making it ideal for evaluating binary classifiers under potential class imbalance. EBC shows a weak average MCC of 12.04%, indicating highly inconsistent prediction alignment. Its lack of polarity sensitivity, especially in sentiment-sparse reviews, results in frequent misclassification of neutral or ambiguous content. MMASA improves to 28.78%, leveraging multimodal input, but fails to ensure consistent class boundaries across all datasets. Polarity interference

between visual and textual channels contributes to prediction drift.

BN-LR, by contrast, achieves a strong average MCC of 73.06%. The model balances true positives and negatives effectively, with its probabilistic dependency structure identifying subtle sentiment shifts even in feature-dense or layered reviews. In electronics and DVDs, where sentiment may alternate within short intervals, BN-LR maintains high directional accuracy. Its logistic core reinforces class stability while the Bayesian layer adapts to context-specific transitions. The consistently superior correlation seen across domains is evident in Figure 5 and clearly summarized in Table 5.

Table 5: Numerical Evaluation of Matthews Correlation Coefficient Scores for BN-LR Model

Amazon Product Review Datasets	EBC	MMAS A	BN-LR
Book	19.1281	30.6010	68.5305
DVD	9.3704	25.7346	74.9062
Electronics	3.2809	25.4553	75.8716
Kitchen Appliances	16.4152	33.3447	72.9343
Average	12.0487	28.7839	73.0607

5.6. BN-LR – FOWLKES–MALLOWS INDEX ANALYSIS

Figure 6 presents Fowlkes–Mallows Index (FMI) values on the y-axis across four Amazon product review datasets shown on the x-axis. Table 6 outlines the corresponding FMI scores for EBC, MMASA, and BN-LR. FMI evaluates the geometric mean of precision and recall, effectively capturing how well predicted sentiment clusters match true sentiment groupings. EBC records the lowest average FMI at 56.42%, reflecting weak consistency in maintaining accurate pairwise sentiment associations.

The model frequently binds unrelated or weakly polar expressions into incorrect clusters, particularly in domains like electronics where opinion and information are densely mixed. MMASA improves moderately to an average of 64.73%, but it lacks a mechanism to regulate inter-modal noise. The model's performance fluctuates when visual inputs offer low sentiment relevance or

contradict textual cues. BN-LR achieves a high FMI average of 86.75%, demonstrating its ability to preserve cluster integrity across all domains. Its Bayesian structure detects contextual dependencies, enabling better token grouping based on sentiment influence.

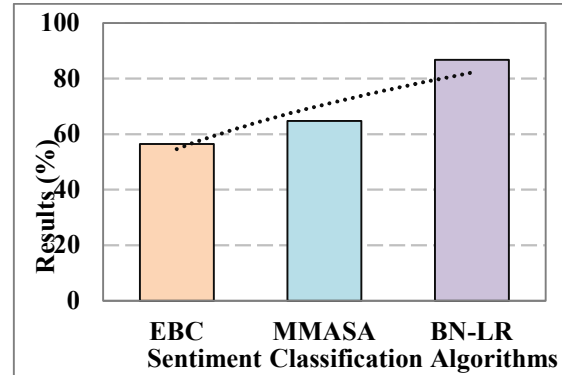


Figure 6: Average Fowlkes–Mallows Index Comparison of BN-LR Across Classification Instances

Table 6: Numerical Evaluation of Fowlkes–Mallows Index Scores for BN-LR Model

Amazon Product Review Datasets	EBC	MMASA	BN-LR
Book	59.8736	65.4835	84.5751
DVD	55.3553	63.2761	87.6244
Electronics	52.0682	63.1438	88.0018
Kitchen Appliances	58.4049	67.0177	86.8186
Average	56.4255	64.7303	86.7550

The logistic component reinforces boundary precision, keeping related sentiment expressions tightly aligned. In multi-aspect reviews, such as kitchen appliances and books, BN-LR maintains clear sentiment grouping, as seen in the results of Figure 6 and Table 6.

5.7. Critical Reflection and Comparative Discussion

The BN-LR framework demonstrates consistent superiority in precision, recall, F-measure, accuracy, MCC, and FMI across all four Amazon product review domains. Compared to EBC and MMASA, the model achieves substantial performance gains, driven by its dynamic integration of Bayesian dependencies with logistic regression's

linear classification boundary. While MMASA leverages multimodal inputs and EBC applies entropy-based heuristics, both fail to dynamically adjust feature weights or model latent sentiment dependencies, limiting their adaptability across shifting review domains. BN-LR, in contrast, captures contextual shifts through dependency-aware probability modeling, achieving better alignment with sentiment-bearing patterns.

Despite these advantages, certain limitations remain. The current model is trained and evaluated solely on binary sentiment labels, excluding neutral or mixed-sentiment cases, which often appear in real-world scenarios. Further, BN-LR operates on structurally clean, domain-balanced data, and its robustness under noisy or low-resource language contexts remains untested. Sentiment ambiguity caused by sarcasm, slang, or implicit expression also presents challenges not fully addressed in the current framework.

Future enhancements should explore the inclusion of neutral sentiment, cross-lingual generalization, and integration with semantic role labeling or contextual emotion detection. Expanding BN-LR toward few-shot learning and dynamic Bayesian adaptation may improve generalization under constrained or evolving sentiment conditions. These areas represent critical opportunities for extending the model's utility across broader sentiment analysis applications.

5.8. Threats to Validity and Justification of Evaluation Criteria

The evaluation of BN-LR is grounded in widely accepted classification metrics including precision, recall, F-measure, classification accuracy, Matthews Correlation Coefficient (MCC), and Fowlkes–Mallows Index (FMI). These criteria have been selected for their comprehensive coverage of performance aspects critical to sentiment classification, particularly under class-imbalance and domain-shift scenarios. Precision and recall directly measure relevance and completeness of sentiment detection. F-measure provides a balanced trade-off, accuracy reflects overall correctness, MCC captures true-versus-false classification correlation under class proportion variance, and FMI quantifies clustering consistency. This combination ensures robustness in both binary decision-making and distributional alignment, justifying their appropriateness for comparative critique.

Despite the consistent results observed, certain threats to validity may influence interpretation. Internal validity may be impacted by domain-specific biases within the Amazon review datasets, where sentiment expressions differ not only by product category but also by review style and user demographics. External validity is limited as the model is tested only on English language reviews within e-commerce, leaving its generalizability to multilingual or social media contexts unverified. Construct validity may be influenced by the exclusion of neutral sentiment and reliance on binary labels, which simplify sentiment expression. Conclusion validity could be affected by possible dataset overlap or reviewer subjectivity embedded in the original labels.

Steps were taken to minimize these threats. All datasets were preprocessed uniformly to reduce noise. The use of multiple domains aimed to simulate realistic domain shifts. Comparative baselines such as EBC and MMASA were included for contrast under identical settings. Yet, future work must involve cross-lingual datasets, inclusion of neutral sentiment, and experimental control for cultural variance to fully establish the model's cross-context stability.

6. CONCLUSION

This study has proposed BN-LR as a statistically grounded and domain-adaptive sentiment classification model, which effectively combines Bayesian network-based dependency modeling with logistic regression for robust cross-domain performance. The model achieved an average classification accuracy of 86.5173%, significantly outperforming baseline methods such as EBC and MMASA across multiple product categories. Beyond performance metrics, the approach introduces a structured, interpretable method for managing polarity shifts and vocabulary variations inherent in cross-domain sentiment analysis. From the author's perspective, BN-LR not only meets the technical objective of generalizability but also reflects a practical solution that balances transparency and predictive reliability. The model demonstrates how integrating probabilistic dependencies into a classical classifier can address nuanced sentiment inconsistencies across domains without sacrificing interpretability. While confident in the contributions made, the author acknowledges certain limitations—specifically, the binary sentiment assumption and the model's current evaluation on linguistically consistent datasets. There remains a strong interest in advancing this framework by incorporating neutral sentiment,

sarcasm detection, and cross-lingual adaptability to better reflect the real-world complexity of sentiment-rich data streams.

6.1. Future Enhancements

BN-LR, while effective for binary sentiment classification across structured English reviews, requires further development to support neutral and mixed sentiments, cross-lingual adaptability, and shorter or informal text formats. The model's current design does not explicitly handle sarcasm, implicit expressions, or sentiment ambiguity. Future work should extend BN-LR with multi-class sentiment capability, multilingual feature alignment, and context-sensitive enhancements such as semantic role labeling. Real-time deployment through incremental learning and streaming analysis also remains an open direction. Addressing these aspects will improve the model's generalizability, scalability, and real-world applicability in broader IT ecosystems.

REFERENCES

- [1] Y. Mao, Z. Chen, S. Liu, and Y. Li, "Unveiling the potential: Exploring the predictability of complex exchange rate trends," *Eng. Appl. Artif. Intell.*, vol. 133, p. 108112, 2024, doi: <https://doi.org/10.1016/j.engappai.2024.108112>.
- [2] M. M. Lesan Sedgh, A. Latif, and S. Emadi, "A novel method for a technology enhanced learning recommender system considering changing user interest based on neural collaborative filtering," *Data Sci. Manag.*, 2024, doi: <https://doi.org/10.1016/j.dsm.2024.09.004>.
- [3] R. N. Patil, Y. P. Singh, S. A. Rawandale, and S. Singh, "Improving Sentiment Classification on Restaurant Reviews Using Deep Learning Models," *Procedia Comput. Sci.*, vol. 235, pp. 3246–3256, 2024, doi: <https://doi.org/10.1016/j.procs.2024.04.307>.
- [4] L.-C. Huang *et al.*, "Natural Language Processing-Powered Real-Time Monitoring Solution for Vaccine Sentiments and Hesitancy on Social Media: System Development and Validation," *JMIR Med. Informatics*, vol. 12, 2024, doi: <https://doi.org/10.2196/57164>.
- [5] P. M., V. kumar K, V. P, and P. M A, "SARCASAM Analysis in Social Networks Using Deep Learning Algorithm," *Procedia Comput. Sci.*, vol. 252, pp. 510–518, 2025, doi: <https://doi.org/10.1016/j.procs.2025.01.010>.
- [6] Y. Li, Q. He, and L. Yang, "Part-of-speech based label update network for aspect sentiment triplet extraction," *J. King Saud Univ. - Comput. Inf. Sci.*, vol. 36, no. 1, p. 101908, 2024, doi: <https://doi.org/10.1016/j.jksuci.2023.101908>.
- [7] S. Qiao and W. Feng, "The effect of novel linguistic-based review features on review helpfulness through information ecosystem and ELM theories: Heterogeneity analysis across platforms," *Comput. Human Behav.*, vol. 160, p. 108357, 2024, doi: <https://doi.org/10.1016/j.chb.2024.108357>.
- [8] Z. Sun, D. He, and Y. Li, "How the readability of manuscript before journal submission advantages peer review process: Evidence from biomedical scientific publications," *J. Informetr.*, vol. 18, no. 3, p. 101547, 2024, doi: <https://doi.org/10.1016/j.joi.2024.101547>.
- [9] S. Deng *et al.*, "Multi-sentiment fusion for stock price crash risk prediction using an interpretable ensemble learning method," *Eng. Appl. Artif. Intell.*, vol. 135, p. 108842, 2024, doi: <https://doi.org/10.1016/j.engappai.2024.108842>.
- [10] N. S. Mullah, W. M. N. W. Zainon, and M. N. Ab Wahab, "Transfer learning approach for identifying negative sentiment in tweets directed to football players," *Eng. Appl. Artif. Intell.*, vol. 133, p. 108377, 2024, doi: <https://doi.org/10.1016/j.engappai.2024.108377>.
- [11] B. Paneru, B. Thapa, and B. Paneru, "Sentiment analysis of movie reviews: A flask application using CNN with RoBERTa embeddings," *Syst. Soft Comput.*, vol. 7, p. 200192, 2025, doi: <https://doi.org/10.1016/j.sasc.2025.200192>.
- [12] J. Li, Y. Huang, Y. Lu, L. Wang, Y. Ren, and R. Chen, "Sentiment Analysis Using E-Commerce Review Keyword-Generated Image with a Hybrid Machine Learning-Based Model," *Comput. Mater. Contin.*, vol. 80, no. 1, pp. 1581–1599, 2024, doi: <https://doi.org/10.32604/cmc.2024.052666>.
- [13] N. Alsaleh, R. Alnanihi, and N. Alowidi, "Hybrid Deep Learning Approach for Automating App Review Classification: Advancing Usability Metrics Classification with an Aspect-Based Sentiment Analysis Framework," *Comput. Mater. Contin.*, vol. 82, no. 1, pp. 949–976, 2025, doi: <https://doi.org/10.32604/cmc.2024.059351>.
- [14] B. Peng, E. Chersoni, Y. Hsu, L. Qiu, and C. Huang, "Supervised Cross-Momentum

- Contrast: Aligning representations with prototypical examples to enhance financial sentiment analysis,” *Knowledge-Based Syst.*, vol. 295, p. 111683, 2024, doi: <https://doi.org/10.1016/j.knosys.2024.111683>.
- [15] C. Huang, J. Chen, Q. Huang, S. Wang, Y. Tu, and X. Huang, “AtCAF: Attention-based causality-aware fusion network for multimodal sentiment analysis,” *Inf. Fusion*, vol. 114, p. 102725, 2025, doi: <https://doi.org/10.1016/j.inffus.2024.102725>.
- [16] Y. Hou, X. Zhong, H. Cao, Z. Zhu, Y. Zhou, and J. Zhang, “A shared-private sentiment analysis approach based on cross-modal information interaction,” *Pattern Recognit. Lett.*, vol. 183, pp. 140–146, 2024, doi: <https://doi.org/10.1016/j.patrec.2024.05.009>.
- [17] K. Aziz *et al.*, “Enhanced UrduAspectNet: Leveraging Biaffine Attention for superior Aspect-Based Sentiment Analysis,” *J. King Saud Univ. - Comput. Inf. Sci.*, vol. 36, no. 9, p. 102221, 2024, doi: <https://doi.org/10.1016/j.jksuci.2024.102221>.
- [18] Z. Liu, L. Cai, W. Yang, and J. Liu, “Sentiment analysis based on text information enhancement and multimodal feature fusion,” *Pattern Recognit.*, vol. 156, p. 110847, 2024, doi: <https://doi.org/10.1016/j.patcog.2024.110847>.
- [19] J. R. Jim, M. A. R. Talukder, P. Malakar, M. M. Kabir, K. Nur, and M. F. Mridha, “Recent advancements and challenges of NLP-based sentiment analysis: A state-of-the-art review,” *Nat. Lang. Process. J.*, vol. 6, p. 100059, 2024, doi: <https://doi.org/10.1016/j.nlp.2024.100059>.
- [20] Z. A. Khan *et al.*, “Developing Lexicons for Enhanced Sentiment Analysis in Software Engineering: An Innovative Multilingual Approach for Social Media Reviews,” *Comput. Mater. Contin.*, vol. 79, no. 2, pp. 2771–2793, 2024, doi: <https://doi.org/10.32604/cmc.2024.046897>.
- [21] D. Chi, T. Huang, Z. Jia, and S. Zhang, “Research on sentiment analysis of hotel review text based on BERT-TCN-BiLSTM-attention model,” *Array*, vol. 25, p. 100378, 2025, doi: <https://doi.org/10.1016/j.array.2025.100378>.
- [22] S. Xie, Q. Chen, X. Fang, and Q. Sun, “Global information regulation network for multimodal sentiment analysis,” *Image Vis. Comput.*, vol. 151, p. 105297, 2024, doi: <https://doi.org/10.1016/j.imavis.2024.105297>.
- [23] J. S. Deshmukh and A. K. Tripathy, “Entropy based classifier for cross-domain opinion mining,” *Appl. Comput. Informatics*, vol. 14, no. 1, pp. 55–64, 2018, doi: [10.1016/j.aci.2017.03.001](https://doi.org/10.1016/j.aci.2017.03.001).
- [24] J. Zhou, J. Zhao, J. X. Huang, Q. V. Hu, and L. He, “MASAD: A large-scale dataset for multimodal aspect-based sentiment analysis,” *Neurocomputing*, vol. 455, pp. 47–58, 2021, doi: [10.1016/j.neucom.2021.05.040](https://doi.org/10.1016/j.neucom.2021.05.040).
- [25] J. Ramkumar, R. Karthikeyan, and K. O. Nitish, “Securing Library Data With Blockchain Advantage,” in *Enhancing Security and Regulations in Libraries with Blockchain Technology*, 2024, pp. 117–138. doi: [10.4018/979-8-3693-9616-2.ch006](https://doi.org/10.4018/979-8-3693-9616-2.ch006).
- [26] P. S. Ponnukumar, N. I. Francis Xavier, and R. Jaganathan, “Stable Plithogenic Cubic Sets,” *J. Fuzzy Ext. Appl.*, vol. 6, no. 2, pp. 410–423, 2025, doi: [10.22105/jfea.2025.449408.1422](https://doi.org/10.22105/jfea.2025.449408.1422).
- [27] R. Jaganathan, S. Mehta, and R. Krishan, “Preface,” *Intell. Decis. Mak. Through Bio-Inspired Optim.*, pp. xiii–xvi, 2024, [Online]. Available: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-85192858710&partnerID=40&md5=f8f1079e8772bd424d2cdd979e5f2710>
- [28] R. Jaganathan, S. Mehta, and R. Krishan, “Preface,” *Bio-Inspired Intell. Smart Decis.*, pp. xix–xx, 2024, [Online]. Available: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-85195725049&partnerID=40&md5=7a2aa7adc005662eebc12ef82e3bd19f>
- [29] V. Valarmathi and J. Ramkumar, “Modernizing Wildfire Management Through Deep Learning and IoT in Fire Ecology,” in *Machine Learning and Internet of Things in Fire Ecology*, 2024, pp. 203–229. doi: [10.4018/979-8-3693-7565-5.ch0010](https://doi.org/10.4018/979-8-3693-7565-5.ch0010).
- [30] B. Suchitra, R. Karthikeyan, J. Ramkumar, and V. Valarmathi, “ENHANCING RECURRENT NEURAL NETWORK PERFORMANCE FOR LATENT AUTOIMMUNE DIABETES DETECTION (LADA) USING EXOCOETIDAE OPTIMIZATION,” *J. Theor. Appl. Inf. Technol.*, vol. 103, no. 5, pp. 1645–1667, 2025, [Online]. Available: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-105000948603&partnerID=40&md5=66c8f11b153fed68b3d0ea9c88c411e>
- [31] J. Ramkumar and D. Ravindran, “Machine learning and robotics in urban traffic flow optimization with graph neural networks and

- reinforcement learning,” in *Machine Learning and Robotics in Urban Planning and Management*, 2025, pp. 83–104. [Online]. Available: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-105000106746&doi=10.4018%2F979-8-3693-9410-6.ch005&partnerID=40&md5=f647b3741afce4400893c2913f2bbf55>
- [32] S. P. Priyadharshini, F. Nirmala Irudayam, and J. Ramkumar, “An Unique Overture of Plithogenic Cubic Overset, Underset and Offset,” in *Studies in Fuzziness and Soft Computing*, vol. 435, 2025, pp. 139–156. [Online]. Available: https://www.scopus.com/inward/record.uri?eid=2-s2.0-105001675443&doi=10.1007%2F978-3-031-78505-4_7&partnerID=40&md5=e9def8c6a233de4fbf8f1549ad72027f
- [33] P. Menakadevi and J. Ramkumar, “Robust Optimization Based Extreme Learning Machine for Sentiment Analysis in Big Data,” *2022 Int. Conf. Adv. Comput. Technol. Appl. ICACTA 2022*, pp. 1–5, Mar. 2022, doi: 10.1109/ICACTA54488.2022.9753203.
- [34] A. Senthilkumar, J. Ramkumar, M. Lingaraj, D. Jayaraj, and B. Sureshkumar, “Minimizing Energy Consumption in Vehicular Sensor Networks Using Relentless Particle Swarm Optimization Routing,” *Int. J. Comput. Networks Appl.*, vol. 10, no. 2, pp. 217–230, 2023, doi: 10.22247/ijcna/2023/220737.
- [35] R. Jaganathan, K. Rajendran, and P. S. Ponnukumar, “Peregrine Falcon Optimization Routing Protocol (PFORP) for Achieving Ultra-Low Latency and Boosted Efficiency in 6G Drone Ad-Hoc Networks (DANET),” *Int. J. Comput. Digit. Syst.*, vol. 17, no. 1, pp. 1–18, 2025, doi: 10.12785/ijcds/1571111848.
- [36] J. Ramkumar, V. Valarmathi, and R. Karthikeyan, “Optimizing Quality of Service and Energy Efficiency in Hazardous Drone Ad-Hoc Networks (DANET) Using Kingfisher Routing Protocol (KRP),” *Int. J. Eng. Trends Technol.*, vol. 73, no. 1, pp. 410–430, 2025, doi: 10.14445/22315381/IJETT-V73I1P135.
- [37] J. Ramkumar, B. Varun, V. Valarmathi, D. R. Medhunhashini, and R. Karthikeyan, “JAGUAR-BASED ROUTING PROTOCOL (JRP) FOR IMPROVED RELIABILITY AND REDUCED PACKET LOSS IN DRONE AD-HOC NETWORKS (DANET),” *J. Theor. Appl. Inf. Technol.*, vol. 103, no. 2, pp. 696–713, 2025, [Online]. Available: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-85217213044&partnerID=40&md5=e38a375e46cf43c95d6702a3585a7073>
- [38] B. Suchitra, J. Ramkumar, and R. Karthikeyan, “FROG LEAP INSPIRED OPTIMIZATION-BASED EXTREME LEARNING MACHINE FOR ACCURATE CLASSIFICATION OF LATENT AUTOIMMUNE DIABETES IN ADULTS (LADA),” *J. Theor. Appl. Inf. Technol.*, vol. 103, no. 2, pp. 472–494, 2025, [Online]. Available: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-85217140979&partnerID=40&md5=9540433c16d5ff0f6c2de4b8c43a4812>
- [39] J. Ramkumar and R. Vadivel, “Multi-Adaptive Routing Protocol for Internet of Things based Ad-hoc Networks,” *Wirel. Pers. Commun.*, vol. 120, no. 2, pp. 887–909, Apr. 2021, doi: 10.1007/s11277-021-08495-z.
- [40] S. P. Priyadharshini and J. Ramkumar, “Mappings Of Plithogenic Cubic Sets,” *Neutrosophic Sets Syst.*, vol. 79, pp. 669–685, 2025, doi: 10.5281/zenodo.14607210.
- [41] J. Ramkumar, R. Vadivel, and B. Narasimhan, “Constrained Cuckoo Search Optimization Based Protocol for Routing in Cloud Network,” *Int. J. Comput. Networks Appl.*, vol. 8, no. 6, pp. 795–803, 2021, doi: 10.22247/ijcna/2021/210727.
- [42] N. K. Ojha, A. Pandita, and J. Ramkumar, “Cyber security challenges and dark side of AI: Review and current status,” in *Demystifying the Dark Side of AI in Business*, IGI Global, 2024, pp. 117–137. doi: 10.4018/979-8-3693-0724-3.ch007.
- [43] L. Mani, S. Arumugam, and R. Jaganathan, “Performance Enhancement of Wireless Sensor Network Using Feisty Particle Swarm Optimization Protocol,” in *ACM International Conference Proceeding Series*, Association for Computing Machinery, 2022. doi: 10.1145/3590837.3590907.
- [44] M. Lingaraj, T. N. Sugumar, C. S. Felix, and J. Ramkumar, “Query aware routing protocol for mobility enabled wireless sensor network,” *Int. J. Comput. Networks Appl.*, vol. 8, no. 3, pp. 258–267, 2021, doi: 10.22247/ijcna/2021/209192.
- [45] J. Ramkumar and R. Vadivel, “Improved Wolf prey inspired protocol for routing in cognitive radio Ad Hoc networks,” *Int. J. Comput.*

- Networks Appl.*, vol. 7, no. 5, pp. 126–136, 2020, doi: 10.22247/ijcna/2020/202977.
- [46] R. Jaganathan, S. Mehta, and R. Krishan, *Bio-Inspired intelligence for smart decision-making*. IGI Global, 2024. doi: 10.4018/9798369352762.
- [47] R. Vadivel and J. Ramkumar, “QoS-enabled improved cuckoo search-inspired protocol (ICSIP) for IoT-based healthcare applications,” in *Incorporating the Internet of Things in Healthcare Applications and Wearable Devices*, IGI Global, 2019, pp. 109–121. doi: 10.4018/978-1-7998-1090-2.ch006.
- [48] J. Ramkumar, R. Karthikeyan, and V. Valarmathi, “Alpine Swift Routing Protocol (ASRP) for Strategic Adaptive Connectivity Enhancement and Boosted Quality of Service in Drone Ad Hoc Network (DANET),” *Int. J. Comput. Networks Appl.*, vol. 11, no. 5, pp. 726–748, 2024, doi: 10.22247/ijcna/2024/45.
- [49] R. Jaganathan and R. Vadivel, “Intelligent Fish Swarm Inspired Protocol (IFSIP) for Dynamic Ideal Routing in Cognitive Radio Ad-Hoc Networks,” *Int. J. Comput. Digit. Syst.*, vol. 10, no. 1, pp. 1063–1074, 2021, doi: 10.12785/ijcds/100196.
- [50] M. P. Swapna, J. Ramkumar, and R. Karthikeyan, “Energy-Aware Reliable Routing with Blockchain Security for Heterogeneous Wireless Sensor Networks,” in *Lecture Notes in Networks and Systems*, V. Goar, M. Kuri, R. Kumar, and T. Senjyu, Eds., Springer Science and Business Media Deutschland GmbH, 2025, pp. 713–723. doi: 10.1007/978-981-97-6106-7_43.
- [51] K. S. J. Marseline, J. Ramkumar, and D. R. Medhunhashini, “Sophisticated Kalman Filtering-Based Neural Network for Analyzing Sentiments in Online Courses,” in *Smart Innovation, Systems and Technologies*, A. K. Somani, A. Mundra, R. K. Gupta, S. Bhattacharya, and A. P. Mazumdar, Eds., Springer Science and Business Media Deutschland GmbH, 2024, pp. 345–358. doi: 10.1007/978-981-97-3690-4_26.
- [52] M. P. Swapna and J. Ramkumar, “Multiple Memory Image Instances Stratagem to Detect Fileless Malware,” in *Communications in Computer and Information Science*, S. Rajagopal, K. Papat, D. Meva, and S. Bajeja, Eds., Springer Science and Business Media Deutschland GmbH, 2024, pp. 131–140. doi: 10.1007/978-3-031-59100-6_11.
- [53] J. Ramkumar and R. Vadivel, “Improved frog leap inspired protocol (IFLIP) – for routing in cognitive radio ad hoc networks (CRAHN),” *World J. Eng.*, vol. 15, no. 2, pp. 306–311, 2018, doi: 10.1108/WJE-08-2017-0260.
- [54] J. Ramkumar, C. Kumuthini, B. Narasimhan, and S. Boopalan, “Energy Consumption Minimization in Cognitive Radio Mobile Ad-Hoc Networks using Enriched Ad-hoc On-demand Distance Vector Protocol,” in *2022 International Conference on Advanced Computing Technologies and Applications, ICACTA 2022*, Institute of Electrical and Electronics Engineers Inc., 2022. doi: 10.1109/ICACTA54488.2022.9752899.
- [55] R. Karthikeyan and R. Vadivel, “Proficient Dazzling Crow Optimization Routing Protocol (PDCORP) for Effective Energy Administration in Wireless Sensor Networks,” in *IEEE International Conference on Electrical, Electronics, Communication and Computers, ELEXCOM 2023*, Institute of Electrical and Electronics Engineers Inc., 2023. doi: 10.1109/ELEXCOM58812.2023.10370559.
- [56] S. P. Geetha, N. M. S. Sundari, J. Ramkumar, and R. Karthikeyan, “ENERGY EFFICIENT ROUTING IN QUANTUM FLYING AD HOC NETWORK (Q-FANET) USING MAMDANI FUZZY INFERENCE ENHANCED DIJKSTRA’S ALGORITHM (MFI-EDA),” *J. Theor. Appl. Inf. Technol.*, vol. 102, no. 9, pp. 3708–3724, 2024, [Online]. Available: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-85197297302&partnerID=40&md5=72d51668bee6239f09a59d2694df67d6>
- [57] J. Ramkumar, R. Karthikeyan, and M. Lingaraj, “Optimizing IoT-Based Quantum Wireless Sensor Networks Using NM-TEEN Fusion of Energy Efficiency and Systematic Governance,” in *Lecture Notes in Electrical Engineering*, V. Shrivastava, J. C. Bansal, and B. K. Panigrahi, Eds., Springer Science and Business Media Deutschland GmbH, 2025, pp. 141–153. doi: 10.1007/978-981-97-6710-6_12.
- [58] R. Karthikeyan and R. Vadivel, “Boosted Mutated Corona Virus Optimization Routing Protocol (BMCVORP) for Reliable Data Transmission with Efficient Energy Utilization,” *Wirel. Pers. Commun.*, vol. 135, no. 4, pp. 2281–2301, 2024, doi: 10.1007/s11277-024-11155-7.
- [59] J. Ramkumar, A. Senthilkumar, M. Lingaraj, R. Karthikeyan, and L. Santhi, “Optimal

- Approach for Minimizing Delays in Iot-Based Quantum Wireless Sensor Networks Using Nm-Leach Routing Protocol,” *J. Theor. Appl. Inf. Technol.*, vol. 102, no. 3, pp. 1099–1111, 2024, [Online]. Available: <https://www.scopus.com/inward/record.uri?eid=2-s2.0-85185481011&partnerID=40&md5=bf0ff974ceabc0ad58e589b28797c684>
- [60] R. Jaganathan, S. Mehta, and R. Krishan, *Intelligent Decision Making Through Bio-Inspired Optimization*, vol. i. 2024. doi: 10.4018/979-8-3693-2073-0.
- [61] R. Jaganathan and V. Ramasamy, “Performance modeling of bio-inspired routing protocols in Cognitive Radio Ad Hoc Network to reduce end-to-end delay,” *Int. J. Intell. Eng. Syst.*, vol. 12, no. 1, pp. 221–231, 2019, doi: 10.22266/IJIES2019.0228.22.
- [62] J. Ramkumar, K. S. Jeen Marseline, and D. R. Medhunhashini, “Relentless Firefly Optimization-Based Routing Protocol (RFORP) for Securing Fintech Data in IoT-Based Ad-Hoc Networks,” *Int. J. Comput. Networks Appl.*, vol. 10, no. 4, pp. 668–687, 2023, doi: 10.22247/ijcna/2023/223319.
- [63] D. Jayaraj, J. Ramkumar, M. Lingaraj, and B. Sureshkumar, “AFSOP: Adaptive Fish Swarm Optimization-Based Routing Protocol for Mobility Enabled Wireless Sensor Network,” *Int. J. Comput. Networks Appl.*, vol. 10, no. 1, pp. 119–129, Jan. 2023, doi: 10.22247/ijcna/2023/218516.
- [64] J. Ramkumar, S. S. Dinakaran, M. Lingaraj, S. Boopalan, and B. Narasimhan, “IoT-Based Kalman Filtering and Particle Swarm Optimization for Detecting Skin Lesion,” in *Lecture Notes in Electrical Engineering*, K. Murari, N. Prasad Padhy, and S. Kamalasadan, Eds., Singapore: Springer Nature Singapore, 2023, pp. 17–27. doi: 10.1007/978-981-19-8353-5_2.
- [65] J. Ramkumar and R. Vadivel, *CSIP—cuckoo search inspired protocol for routing in cognitive radio ad hoc networks*, vol. 556. 2017. doi: 10.1007/978-981-10-3874-7_14.
- [66] J. Ramkumar and R. Vadivel, “Whale optimization routing protocol for minimizing energy consumption in cognitive radio wireless sensor network,” *Int. J. Comput. Networks Appl.*, vol. 8, no. 4, pp. 455–464, 2021, doi: 10.22247/ijcna/2021/209711.