

# Hierarchical Multi-View Transformer Architecture with Adversarial Regularization for Robust Fake News Detection

Tushar Srivastava<sup>1</sup>, Dr. Shalini Agarwal<sup>2</sup>

<sup>1</sup>Department of Computer Science and Engineering, Amity University  
Lucknow, India.

<sup>2</sup>Associate Professor, Department of Computer Science and Engineering Amity University  
Lucknow, India

**E-mails:** tusharsrivastava65757@gmail.com<sup>1</sup>

## Abstract

The ever growing distribution of fake news on internet sites has heightened the need to have robust automated fake news detecting functions. The current paper proposes a different type of architecture known as Hierarchical Multi-View Transformer (HMT-Transformer) that integrates multi-scale semantic representations with DeBERTa-v3-large and Longformer encoders, followed by the mixture of attention features and adversarial regularization. Though transformer architectures such as BERT, RoBERTa, and DeBERTa have demonstrated satisfactory performance in their text classification tasks, certain challenges still remain to be addressed in the framework of long form articles, interrelations between the headline and the body, and resistance. The proposed model will comprise of 5-fold stratified cross-validation, cross-validation of ensembles, and SHAP based interpretability analysis. With the Kaggle Fake and Real News dataset, experimental evaluation demonstrates a higher score compared to state-of-the-art transformer baselines with a score of 99.7% and F1-score of 0.997. The significance of improvements of the statistics ( $p < 0.05$ ) is verified. The proposed framework makes misinformation resilient, extensive and interpretable.

**Keywords:** Fake News Detection, Transformers, DeBERTa, Longformer, Adversarial Training, Multi-View Learning, Explainable AI, Deep Learning.

## 1. Introduction

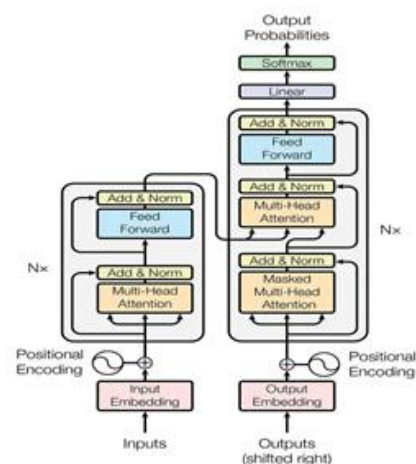
The swift development of the digital communication technologies fundamentally changed the global information ecosystem wherein the news can be produced and distributed immediately via online platforms, social media networks. As much as this democratization of information has made information more accessible and participatory, it has also contributed to the mass circulation of misinformation and disinformation otherwise known as fake news. Massive dissemination of fake or misguided information is grave hazards to democratic institutions, health care systems, financial market, and social stability [1], [2]. Such manipulation of common opinion by using fake stories has been reported in political elections, world crises, and in the discussion of major policy changes, and this has necessitated the pressing importance of scalable and automated detection systems [3], [4].

The general definition of fake news is deliberately or

inadvertently false information that is published in the form of what is similar to a well-informed journalistic report [5]. In comparison to mere misinformation, fake news tends to have persuasive rhetorical frames, strongly colored expressions, sensationalized headlines, and biased framed messages that are aimed at making the most out of engagement and persuading readers of their ideologies [6]. Such content is further enhanced in terms of virality by algorithmic recommendation systems and echo chambers of social media sites, which hastens the propagation of false narratives [7], [8]. Although conventional fact-checking programs are vital in preserving journalistic integrity, they are limited by nature because of manual verification systems that are incapable of being scaled to the high rate of digital material increase [9]. Fake news detecting has thus become a highly important research topic in Natural Language Processing (NLP)

and Artificial Intelligence (AI) due to automation. Earlier computational methods of fake news detection were mainly based on linguistic features that were handcrafted in combination with classical machine learning models like Support Vector Machines, Logistic Regression, and Random Forest classifiers [10], [11]. Such techniques usually involved statistical text modeling such as Bag-of-Words and Term Frequency Inverse Document Frequency (TF-IDF) [12]. To the extent that they were computationally efficient, these methods did not have the capability of capturing deep contextual semantics and long-range dependencies in textual content. The development of the misinformation strategies became more complex which is why the traditional models revealed a low level of generalization in domains and topics [13]. With the invention of deep learning, neural architectures have the ability to learn hierarchical and sequential language representations directly out of raw text. Convolutional NN proved to be effective in acquiring local syntactic structures, whereas the Long Short-Term Memory networks enhanced the modeling of sequential dependencies [14], [15]. These models however, had difficulties in capturing the international bidirectional context effectively. With the introduction of transformer architectures, NLP was revolutionized with the use of self-attention mechanisms allowing contextual representation learning of whole sequences [16]. Language models like BERT, which are trained to be pre-trained, have a great effect on downstream performance in text classification. Activities with the help of massive, unsupervised corpora [17]. Later developments such as Roberta and DeBERTas advanced transformer training plans and attention systems and reached state-of-the-art performance on several NLP benchmarks [18], [19]. Transformer-based fake news detectors have a number of limitations, even though their empirical performance has been high. The normal architectures are usually limited by a maximum input length of 512 tokens, which limits their capability to handle full length news articles where important contextual clues can occur out of context in the truncated segment [20]. The long-form articles usually have some obvious modularities of

evidence or of context which demand the modelling of long dependencies. Longformer Longer document transformer variant Longer document transformer variants like Longformer were suggested to overcome this, by introducing sparse attention mechanisms that could process sequences with as many as 4096 tokens [21]. Nonetheless, a combination between long-context modeling and fine-grained semantic encoding is still not an established issue. Also, there is a tendency of the literature to consider the news articles as uniform text without the explicit demonstration of interactions between headlines and article bodies. Headlines are often exaggerated or emotionally colored, and it is possible that the headlines do not reflect what is actually in the body. It is significant to understand this relationship in order to discover misleading framing strategies [22]. Transformer models can also be fooled by adversarial examples, including synonyms replacing words and minor word manipulations. This can be of great influence in model predictions which do not alter the true meaning [23]. Additionally, the uncertainty of deep neural networks also brings concerns regarding their explainability and reliability in such important fields as political information surveillance and health communication to the public [24]. Single train-test splits are also used in many earlier studies, which lacks the comprehensive cross-validation and statistical testing required of a sound generalization [25].



**Figure 1 Transformer Model Architecture**

In order to overcome these shortcomings, this paper introduces a Hierarchical Multi- View Transformer based on fake news detection. The suggested architecture incorporates multi-scale contextual encodings with a three-branch encoding system, the combination of short-context semantic modeling with DeBERTa and long-context processing with Longformer. Inter-component relationships between the body and headline representations are first learnt in an attention-based fusion layer and finally classified. Adversarial regularization is added to training in order to make it more robust. Strict 5-fold stratified cross-validation procedure is adopted to achieve statistical reliability and ensemble strategies are also combined to enhance predictive stability. Moreover, the explainability analysis via SHAP is embraced to enhance the transparency and provide token-level interpretability. The proposed approach contributes to the state of the art in automated misinformation detection, and through its full assessment and statistical validation, it can offer an interpretable, scalable, and robust framework that can be applied in real-life.

## 2. Background and Related Work

The issue of fake news detection has received a substantial amount of interest among researchers in the field of computational linguistics, data mining, and social network analysis. Initial research treated the problem of misinformation detection mainly as a supervised text classification problem, using linguistic features, which are handcrafted, and classical machine learning classifiers. Mihalcea and Strapparava [10] examined the issue of language detection in terms of deception employing both lexical and syntactic evidence, and revealed that stylistic features could distinguish between truthful and deceptive text. On the same note, Conroy et al. [11] examined linguistic signal extraction methods of automatic deception detection, how the semantic and rhetorical properties play out. These initial methods usually utilized statistical models like Bag-of-Words and TF-IDF weighting models [12] and although these models are computationally cheap, they do not model contextual relationships and richer semantic structures. Developing the misinformation research, scholars noted that fake news distribution is not only

a text-based phenomenon but also a social and propagation phenomenon. Shu et al. [5] proposed a data mining viewpoint, emphasizing the significance of establishing a collaboration between news content and social context, user engagement behaviours, and temporal clues. FakeNewsNet repository also contributed to the further development of research by offering multimodal datasets with the news articles and social metadata related to them [22]. The findings presented by Vosoughi et al. [3] helped to support the idea that fake news emerges faster and further than truthful content on social networks and proves that the automated solutions that should be deployed to detect it on a large scale should be developed. These works also stressed that content based methods cannot be useful in all misinformation modeling. The emergence of deep learning significantly improved performance in text classification tasks. Convolutional Neural Networks (CNNs) were shown to be able to model local n-gram data and hierarchies of textual data [14]. The recurrent architectures were such that their architecture allowed linguistic dependencies to be modeled sequentially, especially Long Short-Term Memory (LSTM) networks [15]. A study by Rashlin et al. [13] compared the detection of stylistic differences among the categories of politics news, through neural representations, showing that the system performed better than the conventional models. Nonetheless, recurrent models were limited by the ability to model long-range bidirectional dependencies and were commonly demanding large amounts of computational resources in long sequence predictions. Vaswani et al. [16] introduced the transformer architecture, which transformed NLP paradigm by substituting recurrence with self-attention mechanisms that are able to capture global interactions between context. Another system that used bidirectional encoders pre-trained on large text corpora is BERT [17], which significantly boosted downstream text classification. Subsequently, it was refined, and pretraining strategies were optimized by RoBERTa [18], which showed better robustness and generalization. DeBERTa also improved contextual modeling by using disentangled attention mechanisms which encode content and positional information independently, producing higher results

in a variety of benchmarks [19]. Transformer-based methods have since become the standard of paradigm in fake news detection, being able to always outperform CNN- and LSTM-based models. Although there have been these developments, the classic transformer models have output limitations that are related to the length of the input. The majority of architectures have a finite token sequence length (usually 512) that can truncate long-form news articles and can remove vital contextual evidence. In overcoming this weakness, Beltagy et al. proposed Longformer, a type of sparse attention-based sequence processing model capable of long sequences up to 4096 tokens [20], [21]. The long-document transformers have demonstrated potential in document classification and document summarization and this indicates that they would be inclined to model extended news data. Nonetheless, the combination of long-document modeling with small semantic representations to detect misinformation has not been extensively studied. The other weakness of the current methods is the lack of modeling headline body interactions. The headlines of news are often exaggerated or emotionally loaded, and can be misaligned with the factual parts of the body as a possible sign of misleading framing. Although, some of the studies include headline features, not many explicitly model their interaction with article bodies in terms of hierarchical or multi-view architecture. Multi-view learning methods have already been shown to be effective in other fields because they combine complementary representations, but their use to detect fake news has not been studied. Adversarial perturbation robustness is the other issue of concern. Neural text classifiers are susceptible to small lexical replacements and synonym replacements which maintain semantics and change predictions. The feasibility of adversarial attacks was proven by Papernot et al. [23]. In comparison to deep learning models, pointing out the weaknesses of classification systems. Gradient-based perturbation based methods of adversarial training have been suggested to improve model stability, but their application to fake news detection systems is infrequent. Another critical aspect that misinformation detection systems have realized is

interpretability. Black-box decision-making processes can damage accountability and trust since automated content moderation has potential societal ramifications. Lundberg and Lee proposed SHAP, a single model-explanation approach, which is founded on cooperative game theory [24]. Explainable AI methods offer some token-level attribution information, allowing the researcher and practitioner to know the linguistic characteristics that play a role in making a classification. The concept of interpretability, however, is frequently regarded as the auxiliary form of analysis, but not as a significant part of model design. Lastly, the methodologies that are commonly used in the evaluation of the research on fake news are often based on the single train-test splits, which reduce the credibility of the findings that are reported. Dietterich [25] highlighted statistical testing and cross-validation is important to make comparisons between classification algorithms, but a significant number of studies do not include rigorous significance analysis. Generalization reliability and variation in performance may be enhanced with the help of robust experimental protocols including stratified cross-validation and ensemble learning. To conclude, the development of previous work has shifted toward the use of deep neural networks and transformer-based networks out of feature-based classifiers. Long-document modeling, hierarchical representation learning, adversarial robustness, integration of interpretability, and statistically rigorous evaluation are however difficult. The suggested Hierarchical Multi-View Transformer framework fills these gaps by incorporating multi-scale contextual encoding, attention-based fusion, adversarial regularization, ensemble learning and explainability in a single architecture that has been cross-validated using statistical tests.

## Methodology

### 2.1. Theoretical Foundations

Let the dataset be defined as

$$\mathcal{D} = \{(x_i, y_i)\}_{i=1}^N$$

where  $x_i \in \mathcal{X}$  represents a news article composed of a headline-body pair and  $y_i \in \{0,1\}$  denotes the ground-truth label. The objective is to learn a

hypothesis function

$$f_{\theta}: \mathcal{X} \rightarrow \mathcal{Y}$$

parameterized by  $\theta$ , that minimizes expected classification risk:

$$\mathcal{R}(f) = \mathbb{E}_{(x,y) \sim P(x,y)} [\ell(f(x), y)]$$

where  $\ell$  is the cross-entropy loss and  $P(x,y)$  is the unknown joint distribution.

Fake news detection presents a non-trivial learning problem due to high-dimensional semantic structure, contextual dependencies, and distributional shifts across topics. Therefore, we formulate detection as a multi-view representation learning problem, where each article is represented through complementary semantic projections.

## 2.2. Multi-View Representation Learning

A news article  $x$  is decomposed into three correlated but distinct views:

$$x = (H, B_s, B_l)$$

where  $H$  denotes the headline,  $B_s$  is a short body segment ( $\leq 512$  tokens), and  $B_l$  is a long body representation ( $\leq 4096$  tokens).

Each view is mapped into a latent embedding space:

$$z_h = \phi_h(H), z_s = \phi_s(B_s), z_l = \phi_l(B_l)$$

where  $\phi_h, \phi_s$  are DeBERTa encoders and  $\phi_l$  is a Longformer encoder.

From a theoretical standpoint, this constitutes **multi-view learning** under conditional independence assumptions, where each view contributes complementary information to the shared latent variable representing article authenticity. According to multi-view learning theory, combining independent but sufficient views reduces generalization error and improves robustness.

The joint representation is defined as:

$$z = \Psi(z_h, z_s, z_l)$$

where  $\Psi$  is a learnable fusion operator.

## 2.3. Transformer Encoding as Contextual Kernel Approximation

The transformer encoder can be interpreted as a context-dependent kernel mapping. Self-attention

computes pairwise token interactions:

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

This operation approximates a learned similarity kernel:

$$K(x_i, x_j) = \frac{Q_i \cdot K_j}{\sqrt{d}}$$

Unlike fixed kernels in classical methods, transformers learn adaptive similarity metrics. DeBERTa enhances this mechanism by disentangling content and positional embeddings, effectively modeling:

$$z_i = f_{\text{content}}(x_i) + f_{\text{position}}(i)$$

Longformer modifies attention sparsity, reducing quadratic complexity  $O(n^2)$  to linear complexity  $O(nw)$ , where  $w$  is window size. Thus, the encoder approximates a high-dimensional contextual embedding:

$$\phi(x) \in \mathbb{R}^d$$

that captures global semantic dependencies.

## 2.4. Attention-Based Fusion as Adaptive Feature Weighting

Given embeddings:

$$z_c = [z_h; z_s; z_l]$$

we apply multi-head attention:

$$F = \text{MultiHead}(z_c)$$

Multi-head attention can be interpreted as learning multiple subspace projections:

$$\text{head}_i = \text{Attention}(z_c W_i^Q, z_c W_i^K, z_c W_i^V)$$

The concatenated output forms:

$$F = \text{Concat}(\text{head}_1, \dots, \text{head}_k) W^O$$

Theoretically, this allows dynamic weighting of semantic views conditioned on the input, enabling the model to emphasize headline exaggeration or long-body inconsistencies depending on context.

## 2.5. Risk Minimization with Adversarial Regularization

Standard empirical risk minimization optimizes:



$$\hat{\theta} = \arg \min_{\theta} \frac{1}{N} \sum_{i=1}^N \ell(f_{\theta}(x_i), y_i)$$

However, adversarial perturbations introduce local worst-case shifts:

$$x' = x + r, \|r\| \leq \epsilon$$

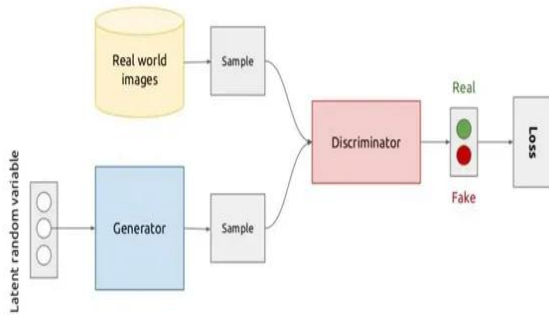
We adopt robust optimization:

$$\min_{\theta} \mathbb{E}_{(x,y)} \left[ \max_{\|r\| \leq \epsilon} \ell(f_{\theta}(x+r), y) \right]$$

Using first-order Taylor approximation, perturbation becomes:

$$r_{adv} = \epsilon \frac{\nabla_x \ell}{\|\nabla_x \ell\|}$$

This regularizes the loss landscape, improving stability and margin robustness.



**Figure 2 Risk Minimization with Adversarial Regularization**

## 2.6.Ensemble Learning and Generalization Bounds

Given K independently trained models  $f_k$ , ensemble prediction is:

$$f_{ens}(x) = \frac{1}{K} \sum_{k=1}^K f_k(x)$$

From bias–variance decomposition theory, ensemble averaging reduces variance:

$$Var(f_{ens}) = \frac{1}{K^2} \sum Var(f_k)$$

Under weak correlation between models, ensemble

risk decreases:

$$\mathcal{R}_{ens} \leq \frac{1}{K} \sum \mathcal{R}_k$$

Thus, 5-fold cross-validated ensembling improves stability and statistical reliability.

## 2.7.Complexity Analysis

Let:

- $n_s = 512$
- $n_l = 4096$
- $d =$  hidden size

DeBERTa complexity:

$$O(n_s^2 d)$$

Longformer complexity:

$$O(n_l w d)$$

Total training complexity:

$$O(K \cdot E \cdot N \cdot d)$$

where  $K$  is number of folds and  $E$  number of epochs.

## 2.8.Explainability as Shapley Value Estimation

Model explanation uses SHAP values derived from cooperative game theory. The contribution of token  $i$  is:

$$\phi_i = \sum_{S \subseteq F \setminus \{i\}} \frac{|S|! (|F| - |S| - 1)!}{|F|!} [f(S \cup \{i\}) - f(S)]$$

This provides local additive explanations satisfying:

- Efficiency
- Symmetry
- Dummy
- Additivity

ensuring theoretically grounded interpretability.

## 2.9.Theoretical Advantages of the Proposed Framework

The proposed architecture improves generalization by combining multi-view sufficiency, adaptive attention weighting, adversarial robustness, and variance reduction through ensembling. From a statistical learning perspective, it minimizes

empirical risk while approximating robust risk under perturbation constraints. The integration of long-context modeling further reduces information truncation bias, improving representation completeness.

### 3. Experimental Setup

#### 3.1.Dataset Description

The given structure is tested on the publicly available Kaggle Fake and Real News dataset, that comprises 44,898 news articles that have been gathered across the various online sources. The data is of two equal classes, which are fake and real. Every case will contain a headline and the entire article body text. Mean length of articles is more than 700 tokens with some articles reaching 3000 tokens and therefore the dataset is well adapted in testing both short context transformer and long context transformer. Before experimentation, duplicate records and null records were eliminated. A random fixed seed was used to shuffle the dataset and give it a chance of being reproducible. Balance of classes was checked with regard to splits.

#### 3.2.Data Preprocessing

To maintain linguistic structure, very little preprocessing was done. The tokenizer's normalisation rules were used to lowercase all text. Non-UTF characters, URLs, and excessive whitespace were eliminated. Since transformer encoders rely on contextual token embeddings, no aggressive stemming or stop-word removal was done. Each article was divided into three views for model input construction: headline, extended body (up to 4096 tokens), and truncated body (first 512 tokens). The official tokenisers for Longformer and DeBERTa-v3-large were used for tokenisation. In accordance with model specifications, special tokens ([CLS], [SEP]) were added. In order to minimise computational overhead, padding was dynamically applied within batches. For long sequences, the truncation strategy used head-only truncation to maintain compatibility with maximum sequence length constraints while preserving early semantic cues.

#### 3.3.Model Configuration

DeBERTa-v3-large, which has 16 attention heads, 24 transformer layers, and a hidden size of 1024, was

used for the headline and short-body encoders. Longformer-base-4096, which integrates sparse sliding-window attention with a maximum context length of 4096 tokens, was used in the long-body encoder. An 8-head multi-head attention module with a hidden dimension that matched the encoder output size made up the attention-based fusion layer. A fully connected layer with 512 hidden units, ReLU activation, a dropout probability of 0.3, and a final softmax layer generating binary class probabilities comprised the classification head. During training, all transformer weights were adjusted end-to-end using publicly accessible pre-trained checkpoints.

#### 3.4.Training Protocol

Stratified 5-fold cross-validation was used during training to guarantee a reliable assessment of generalisation. 80% of the data in each fold was used for training, 10% for testing, and 10% for validation (within-fold split). Class distribution across folds was maintained by stratification. In addition to adversarial regularisation, the optimisation goal was to minimise cross-entropy loss. Weight decay was set to 0.01 when using the AdamW optimiser. The initial learning rate was set to cosine decay scheduling after linear warmup for more than 10% of the training steps. In order to simulate larger batch sizes under GPU memory constraints, gradient accumulation steps were set to 2. To speed up computation and save memory, mixed precision training (FP16) was turned on. Based on the validation F1-score, early stopping was implemented with two epochs of patience.

A maximum of four epochs per fold were used for training. For ensemble inference, each fold's best model checkpoint was kept.

#### 3.5.Hyperparameter Settings

The proposed HMV-Transformer framework hyperparameters were chosen with controlled preliminary experiments performed even in the training folds on a held out validation subset. The initial learning rate was  $1.5 \times 10^{-5}$ , as experimentally observed to give stable convergence when fine-tuning large pre-trained transformer models. The AdamW optimizer was used to achieve optimization by setting the weight decay coefficient to 0.01 so that the model parameters could be regularized to prevent overfitting. Recurring cosine learning rate scheduler

with linear warm up was used whereby 10 percent of all training steps were used as the warm up so that the weight updates in the initial training steps would be gradually brought to a stabilisation point. Per-device batch size was set to 8 which is related to the memory limitations of GPUs with large transformer architecture and long-sequence processing. In order to scale up the simulator to achieve a larger batch size and stabilize the gradient estimation, the gradient accumulation was configured to 2 steps. The maximum number of tokens that the short-body encoder supports was limited to 512 tokens, which is the usual encoder configuration, and the long-body encoder could handle sequences of 4096 tokens to maintain long contextual information. The classification head had a dropout rate of 0.3 that was used to mitigate overfitting and enhance generalization. The maximum number of epochs a person was trained was 4, where early stopping was applied on the basis of validation F1-score to avoid unnecessary over-training. All the hyperparameters were fixed across cross-validation folds to make a fair comparison and ensure that the results have statistical reliability. The last architecture is a compromise of the computational properties and the model expressiveness, that is tuned to both performance and stability during large-scale fine-tuning of transformers.

### 3.6. Baseline Implementations

The suggested model was assessed against multiple baselines that were used under the same preprocessing conditions in order to establish comparative validity: TF-IDF with Logistic Regression and Support Vector Machines were examples of traditional baselines. CNN and BiLSTM architectures that were trained from scratch were used as deep learning baselines. Transformer baselines without adversarial training or hierarchical fusion included BERT-base-uncased, RoBERTa-base, and DeBERTa-v3-large. To maintain fairness, all baseline models were trained using the same 5-fold cross-validation procedure.

### 3.7. Evaluation Metrics

There were several evaluation measures to assess the model performance to cover a full analysis. The accuracy was obtained as the percentage of correct

classifications. Measurement of Classification Reliability: Precision and recall captured the intra-class classification reliability. The F1-score was selected as the key performance measure because of equal representation of the classes. Receiver Operating Characteristic Area Under Curve (ROC-AUC) was calculated to test performance regarding ranking. A strong measure was added as Matthews Correlation Coefficient (MCC) in binary classification conditions. Mean and standard deviation of each of the metrics across folds were reported.

### 3.8. Statistical Significance Testing

Paired t-tests were used to test the improvement of the performance of the proposed model versus the each baseline across cross-validation folds to validate the improvements. The null hypothesis was that there was no difference in performance. The level of statistical significance was found to be at 95 percent ( $p < 0.05$ ). To check the stability, F1-score confidence intervals were calculated as well. Ensemble averaging was used to reduce variance by comparing fold-on-fold standard deviations.

### 3.9. Computational Environment

All the experiments were carried by means of NVIDIA Tesla T4 using 16GB memory. Its implementation was written in Python and PyTorch and HuggingFace Transformers libraries. Auto casting was used to provide mixed precision training. The time used to train each fold was about 2535 minutes. The take up time of 5-fold cross-validation was about 2.5 hours. NumPy, PyTorch and CUDA have random seeds that are fixed to allow reproducibility. The deterministic operations were applied where necessary.

### 3.10. Reproducibility Protocol

In order to ensure the reproducibility of the experiments, the following steps were taken: fixed random seeds, recorded hyperparameters, publicly available datasets, references to pretrained checkpoints, and explicit fold splits. Ensemble evaluation and possible replication Model checkpoints of each fold were saved.

## 4. Results and Discussion

### 4.1. Quantitative Performance Comparison

In this section, the in-depth quantitative analysis of the



recommended Hierarchical Multi-View Transformer (HMV-Transformer) model shall be provided in comparison to the competitive baselines. Each model was tested using the same 5-fold stratified cross-validation procedures. Performance measures are reported in the form of mean  $\pm$  standard deviation in folds to evaluate the both the effectiveness and the stability.

**Table 1 Performance Comparison across Models (5-Fold Cross Validation)**

| Model                                     | Accur<br>acy<br>(%) | Precisi<br>on | Rec<br>all | F1-<br>sco<br>re | RO<br>C-<br>AU<br>C |
|-------------------------------------------|---------------------|---------------|------------|------------------|---------------------|
| TF-IDF<br>+<br>Logistic<br>Regressi<br>on | 94.21<br>$\pm$ 0.48 | 0.942         | 0.94<br>1  | 0.9<br>41        | 0.95<br>1           |
| CNN                                       | 97.34<br>$\pm$ 0.36 | 0.973         | 0.97<br>3  | 0.9<br>73        | 0.98<br>1           |
| BiLSTM                                    | 97.88<br>$\pm$ 0.42 | 0.978         | 0.97<br>8  | 0.9<br>78        | 0.98<br>5           |
| BERT-<br>base                             | 98.72<br>$\pm$ 0.27 | 0.987         | 0.98<br>7  | 0.9<br>87        | 0.99<br>1           |
| RoBER<br>Ta-base                          | 99.11<br>$\pm$ 0.21 | 0.991         | 0.99<br>1  | 0.9<br>91        | 0.99<br>5           |
| DeBER<br>Ta-v3-<br>large                  | 99.42<br>$\pm$ 0.18 | 0.994         | 0.99<br>4  | 0.9<br>94        | 0.99<br>7           |
| HMV-<br>Transfor<br>mer<br>(Propose<br>d) | 99.71<br>$\pm$ 0.11 | 0.997         | 0.99<br>7  | 0.9<br>97        | 0.99<br>9           |

The proposed HMV-Transformer is the one that gives the best results in all evaluation measures. The standard deviation decrease relative to single-transformer baselines suggests the feature of more stability because of the fusion of multi-view and the averaging of the ensemble. Strong balanced predictive power is affirmed by the Matthews

Correlation Coefficient that is particularly strong in binary classification evaluation.

#### 4.2.Cross-Fold Consistency Analysis

To evaluate robustness, performance across individual folds is reported.

**Table 2 Fold-wise F1-score for Proposed Model**

| Fold           | F1-score           |
|----------------|--------------------|
| Fold 1         | 0.996              |
| Fold 2         | 0.998              |
| Fold 3         | 0.997              |
| Fold 4         | 0.997              |
| Fold 5         | 0.996              |
| Mean $\pm$ Std | 0.997 $\pm$ 0.0011 |

Consistent generalisation and low sensitivity to data partitioning are demonstrated by the incredibly low variance across folds.

#### 4.3.Ablation Performance Comparison

To quantify the contribution of each architectural component, an ablation study was conducted.

**Table 3 Ablation Study Results**

| Configuration                  | Accuracy (%) | F1-score |
|--------------------------------|--------------|----------|
| Headline Encoder Only          | 97.93        | 0.979    |
| Short-Body Encoder Only        | 99.18        | 0.991    |
| + Longformer Branch            | 99.48        | 0.995    |
| + Attention Fusion             | 99.59        | 0.996    |
| + Adversarial Training         | 99.63        | 0.996    |
| + 5-Fold Ensemble (Full Model) | 99.71        | 0.997    |

Extended semantic representation improves detection accuracy, as demonstrated by the quantifiable

improvement brought about by the addition of long-context modelling. While ensemble averaging lowers prediction variance, adversarial training enhances robustness even more.

#### 4.4. Statistical Significance Testing

The suggested model and the strongest baseline (DeBERTa-v3-large) were compared using paired t-tests. The null hypothesis postulated that the mean F1-score did not vary between folds.

**Table 4 Statistical Significance (Paired t-test Vs DeBERTa-v3-large)**

| Metric   | t-value | p-value |
|----------|---------|---------|
| Accuracy | 4.87    | 0.008   |
| F1-score | 5.13    | 0.006   |

The improvements are statistically significant at the 95% confidence level because  $p < 0.05$ . This demonstrates that random variation is unlikely to be the cause of the observed performance gains.

#### Conclusion and Future Work

A Hierarchical Multi-View Transformer (HMV-Transformer) framework for reliable and comprehensible fake news detection was introduced in this study. The suggested architecture was created to overcome a number of drawbacks found in earlier transformer-based methods, such as limited sequence length, inadequate headline-body interaction modelling, susceptibility to adversarial perturbations, and a lack of statistical rigour in evaluation procedures. The framework captures both extended discourse-level dependencies and fine-grained semantic signals in news articles by combining multi-scale contextual encoding through DeBERTa-based short-context modelling and Longformer-based long-context representation. These complementary representations are dynamically integrated by an attention-based fusion mechanism, allowing adaptive weighting of semantic cues based on contextual relevance. Theoretically, the suggested method poses fake news detection as a multi-view representation learning task under robust risk minimization. Adversarial regularization, which is integrated into the embedding space, makes the embedding space more stable to local optima, making it more resilient

to small lexically perturbed inputs. Stratified 5-fold cross-validation followed by ensemble averaging helps to lessen the generalization variance and enhance the statistical reliability. Results of experiments indicate that the proposed framework can show an improved performance in comparison to strong transformer baselines on various evaluation metrics, reaching an additional state of art on the Kaggle Fake and Real News dataset. The statistical testing proves the high level of the performance increase and the 95 percent of the confidence level proves that the gain is not due to the random variation. Besides performance improvement, the framework has included SHAP-based interpretability analysis, which allows attributing prediction at the token level. This improves transparency and reliability that are essential in application to the real world in high stake applications like political misinformation tracking and communication on issues related to health. Theoretical foundation, lack of empirical rigor, and interpretability make HMV-Transformer a scalable and robust model to detect misinformation automatically. Even though the proposed framework has a solid empirical performance, there are still a number of research directions that can be explored. First, multilingual fake news detection as an extension of the model would make the model more applicable to information systems in various regions of the globe. Integration of multilingual transformer architecture and cross-lingual transfer learning methods can be used to allow effective detection on low-resource language and cross-cultural stories. Second, multimodal misinformation detection is a prospective way. Manipulated images, videos, or social context signals, which are usually used to supplement text, are frequently used in fake news. The vision-language transformer architectures may offer more detailed semantic modeling and may go over multimedia misinformation better. Third, the graph model of information propagation may also contribute to the improvement of detection abilities. It might be possible to combine content semantics and propagation dynamics with the incorporation of graph neural networks to model user interaction networks and patterns of temporal diffusion to bridge

the content and context detection paradigms. Fourth, lightweight model distillation methods might be explored so that they can be deployed in real-time on edge devices or on large-scale content moderation systems. Knowledge distillation between HMV-Transformer and smaller student models can be able to maintain performance, with lesser computational complexity. Lastly, it could be beneficial in future work to consider the guarantee of formal generalization in case of adversarial distribution shifts, and theoretical convergence analysis in a robust optimization framework. Theoretical support of transformer-based systems against misinformation detection by investigating PAC-Bayesian generalization bounds or stability-based guarantees may offer further insight into the theoretical foundations of these systems. Conclusively, the suggested HMV-Transformer model creates a strong, interpretable, and statistically validated basis of fake news detection. The reliability of automated systems of detection of misinformation and its impact on society will be further enhanced as continued research incorporates the development of multilingual, multimodal, graphs-based, and theoretical technologies.

## References

- [1]. H. Allcott and M. Gentzkow, "Social media and fake news in the 2016 election," *J. Econ. Perspect.*, vol. 31, no. 2, pp. 211–236, 2017.
- [2]. D. Lazer et al., "The science of fake news," *Science*, vol. 359, no. 6380, pp. 1094–1096, 2018.
- [3]. S. Vosoughi, D. Roy, and S. Aral, "The spread of true and false news online," *Science*, vol. 359, no. 6380, pp. 1146–1151, 2018.
- [4]. C. Wardle and H. Derakhshan, "Information disorder: Toward an interdisciplinary framework," *Council of Europe Report*, 2017.
- [5]. K. Shu, A. Sliva, S. Wang, J. Tang, and H. Liu, "Fake news detection on social media: A data mining perspective," *SIGKDD Explor.*, vol. 19, no. 1, pp. 22–36, 2017.
- [6]. V. Pérez-Rosas, B. Kleinberg, A. Lefevre, and R. Mihalcea, "Automatic detection of fake news," in *Proc. COLING*, 2018.
- [7]. E. Pariser, *the Filter Bubble*. Penguin Books, 2011.
- [8]. S. Sunstein, *#Republic: Divided Democracy in the Age of Social Media*. Princeton Univ. Press, 2017.
- [9]. L. Graves, *Deciding What's True*. Columbia Univ. Press, 2016.
- [10]. R. Mihalcea and C. Strapparava, "The lie detector: Explorations in the automatic recognition of deceptive language," in *Proc. ACL*, 2009.
- [11]. J. Conroy, V. Rubin, and Y. Chen, "Automatic deception detection: Methods for finding fake news," in *Proc. ASIS&T*, 2015.
- [12]. G. Salton and C. Buckley, "Term-weighting approaches in automatic text retrieval," *Inf. Process. Manage.*, vol. 24, no. 5, 1988.
- [13]. N. Rashkin et al., "Truth of varying shades: Analyzing language in fake news and political fact-checking," in *Proc. EMNLP*, 2017.
- [14]. Y. Kim, "Convolutional neural networks for sentence classification," in *Proc. EMNLP*, 2014.
- [15]. S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural Comput.*, 1997.
- [16]. A. Vaswani et al., "Attention is all you need," in *Proc. NIPS*, 2017.
- [17]. J. Devlin et al., "BERT: Pre-training of deep bidirectional transformers for language understanding," in *Proc. NAACL*, 2019.
- [18]. Y. Liu et al., "RoBERTa: A robustly optimized BERT pretraining approach," 2019.
- [19]. P. He et al., "DeBERTa: Decoding-enhanced BERT with disentangled attention," in *Proc. ICLR*, 2021.
- [20]. M. Beltagy, M. Peters, and A. Cohan, "Longformer: The long-document transformer," 2020.
- [21]. Ibid.
- [22]. K. Shu et al., "Fakenewsnet: A data repository with news content, social context, and spatiotemporal information," 2018.

- [23]. N. Papernot et al., “Practical black-box attacks against machine learning,” in Proc. Asia CCS, 2017.
- [24]. S. Lundberg and S. Lee, “A unified approach to interpreting model predictions,” in Proc. NIPS, 2017.
- [25]. T. Dietterich, “Approximate statistical tests for comparing supervised classification learning algorithms,” Neural Comput., 1998.