

LSW-Net: A Linguistic-Signal Wavelet Network for Fake-News Detection in News Text

Kheirolah Rahsepar Fard

*Department of Computer Engineering and Information Technology,
University of Qom, Qom,
Iran.*

rahsepar@qom.ac.ir

Ghaith Allawi Alawadi

*Department of Computer Engineering and Information Technology,
University of Qom, Qom,
Iran.*

Dr.ghaith.alawadi@gmail.com

Corresponding Author: Ghaith Allawi Alawadi

Copyright © 2026 Kheirolah Rahsepar Fard and Ghaith Allawi Alawadi. This is an open access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Abstract

The proliferation of fake news makes preserving information integrity on the Internet essential, yet most detectors learn opaque, model-specific representations. We introduce LSW-Net, which represents a document as a one-dimensional linguistic signal whose amplitude at each token is an explicit function of a part-of-speech prior, sentiment polarity and TF-IDF salience. This signal is decomposed with a discrete wavelet transform (DWT; Daubechies db4, three levels) into interpretable multi-resolution sub-bands that a bidirectional LSTM reads; in parallel, a denoising-autoencoder latent and a contextual sentence embedding are fused with the BiLSTM state at the classification head. On the GonzaloA/fake_news benchmark LSW-Net reaches 96.65% accuracy, 96.92% F1 and 0.994 AUC; on this lexically separable corpus it is competitive with, but does not exceed, a TF-IDF + SVM baseline (98.2% accuracy), and an ablation shows the wavelet channel adds no accuracy there. To test where the representation matters, we add the PHEME rumour corpus under a leave-one-event-out protocol: on unseen events the wavelet channel improves macro-F1 in six of eight events, and LSW-Net degrades far less when moving from in-distribution to unseen-event evaluation (-29 macro-F1 points) than the TF-IDF + SVM baseline (-37 points). LSW-Net thus trades a small amount of in-distribution accuracy for a representation that is both inspectable and more robust to topic shift. Code is publicly available at <https://github.com/Ghaith-Alawady/lsw-net>.

Keywords: Fake-news detection, Wavelet transform, Autoencoder, BiLSTM, Linguistic signal, Misinformation.

1. INTRODUCTION

The spread, and potential influence, of online disinformation (or 'fake news') [1], is rapid enough to potentially lead to panic, social polarization as well as a loss of trust within society, hence automatic detection is important [2, 3].

While the standard approaches model either lexical content [4, 5] or propagation structure [6–8], we can argue that the document also carries a latent signal whose multi-resolution structure can be analyzed with the wavelet transform [9]. We combine this view with a strong semantic encoder [10, 11]. Emotion [12], or temporal propagation [13], are among the new detectors identified in TABLE 1.

Table 1: Comparison of the state-of-the-art work (accuracies are on the datasets shown and are not directly comparable across datasets).

Ref.	Method	Problem addressed	Dataset	Acc. (%)	Limitation
[6]	Rumor2vec	Joint text + propagation	Twitter / Weibo	85.2 / 95.1	Early-stage sparsity; graph complexity
[14]	VRoC (variational AE)	Multi-task classification	Twitter	87.6	Limited interpretability
[3]	CSE-BiLSTM	Veracity verification	PHEME	90.56	Compute-intensive; many vectorizers
[12]	SEMTEC (emotion transformer, 2024)	Emotion-driven detection	PHEME	92.0	Needs emotion/sentiment tagging
[13]	Temporal Tree Transformer (2025)	Propagation + temporal	PHEME	75.84	Needs propagation tree and timestamps
(ours)	LSW-Net	Signal + wavelet + semantics	GonzaloA/fake_new	96.65	Wavelet gain is task-dependent
(ours)	LSW-Net (cross-event)	Rumour detection, unseen event	PHEME (9-event, LOEO)	57.0 (macro-F1)	Small in-distribution gain; robust to topic shift

In summary, we make the following contributions: (1) the formal definition of a linguistic signal and its wavelet representation (Sec. 3); (2) an end-to-end architecture that incorporates both the wavelet representation and contextual semantics; (3) an analysis of the results by performing a controlled

evaluation on a standard dataset with strong baseline models and an ablation study; and (4) a per-band effect size analysis. We discuss cases of improved performance with the wavelet channel. We give background in Sec. 2, methods in Sec. 3, setup in Sec. 4, results in Sec. 5, and conclusions in Sec. 6.

2. BACKGROUND AND RELATED WORK

Early text-only detectors classify a post directly from its content. Text-fusion networks [4] and variational autoencoders [14], learn a joint latent representation, while pretrained transformer encoders such as BERT and Sentence-BERT [10, 11] now dominate. Recent work extends this line with memory-augmented transformer-graph hybrids [15], and long-range graph transformers for early detection [16], and with large-language-model and multimodal detectors that combine text with emotion, images or external knowledge [17–20]. These models report strong accuracy but entangle feature extraction and classification in a single opaque, high-dimensional space.

A second line models how a message spreads. Propagation- and graph-based methods [6–8], encode the conversation tree or diffusion cascade, and knowledge- or homogeneity-aware graph networks add external context and data augmentation [2, 5, 21–23]. Comment-sentiment CNN-LSTM detectors [24], and multimodal memory networks [25], enrich the input further. Such structural signals help, but require propagation graphs that are unavailable for isolated posts or at the earliest moments of spread.

On the PHEME rumour corpus specifically, strong baselines include context-aware [1, 26], multi-task [27], tree-structured [28], and conversation-structure models [29], edge-enhanced Bayesian graph networks [30], and interpretable user-interaction attention [31]. These works establish the leave-one-event-out protocol we adopt, under which memorised event vocabulary no longer helps and genuine generalisation is measured.

Across all of these families, feature extraction and classification are learned jointly, so the evidence behind a prediction is locked in an opaque latent space. Multi-resolution wavelet analysis [9], offers exactly the inspectable, energy-localised view we need but has rarely been applied to misinformation text. LSW-Net brings this view to the task and, as our cross-event experiments show, is more robust to topic shift than lexical baselines.

3. PROPOSED METHOD AND METHODOLOGY

3.1 Definition of the Linguistic Signal

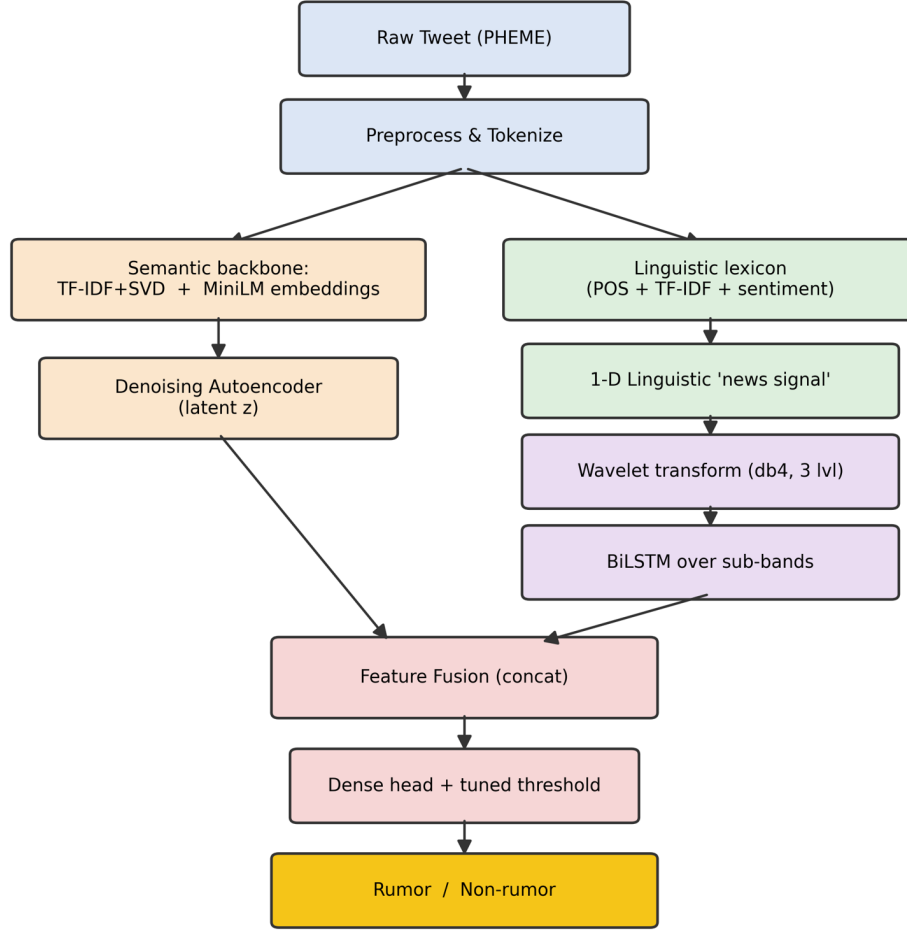
For each token w_t with POS tag p_t , sentiment polarity s_t in $[-1, 1]$ and TF-IDF salience $\tau(w_t)$, and a fixed POS-importance prior $\omega(p)$, the non-negative amplitude is:

$$a_t = \omega(p_t)(1 + |s_t|)(0.5 + \tau(w_t)) \quad (1)$$

The linguistic signal is the ordered amplitude sequence, padded/truncated to length L and z-normalised:

$$x = (a_1, \dots, a_L) \in \mathbb{R}^L, \tilde{x} = \frac{x - \mu_x}{\sigma_x} \quad (2)$$

Definition (Linguistic signal). Given lexicon maps (ω , s , τ), the linguistic signal is the deterministic length-normalised 1-D function of Eq. (1)-(2); every coordinate has an explicit linguistic cause. Overall architecture and data flow of LSW-Net is given in Figure 1.



LSW-Net: Linguistic-Signal Wavelet Network

Figure 1: Overall architecture and data flow of LSW-Net.

3.2 Feature Extraction (Denoising Autoencoder)

A denoising AE [14], with encoder f and decoder g minimises the reconstruction error of a corrupted input:

$$\mathcal{L}_{AE} = \frac{1}{N} \sum_i \|x_i - g(f(x_i + \varepsilon))\|^2, \varepsilon \sim \mathcal{N}(0, \sigma^2 I) \quad (3)$$

3.3 Wavelet Representation

A multilevel DWT (Daubechies db4, three levels) [9], yields coefficients and per-band energies: FIGURE 2 illustrates an example document as a 1-D linguistic signal together with its wavelet sub-bands.

$$W_{j,k} = \sum_n \tilde{x}[n] \psi_{j,k}[n], \psi_{j,k}[n] = 2^{-j/2} \psi(2^{-j}n-k) \quad (4)$$

$$E_b = \sum_k W_{b,k}^2 \quad (5)$$

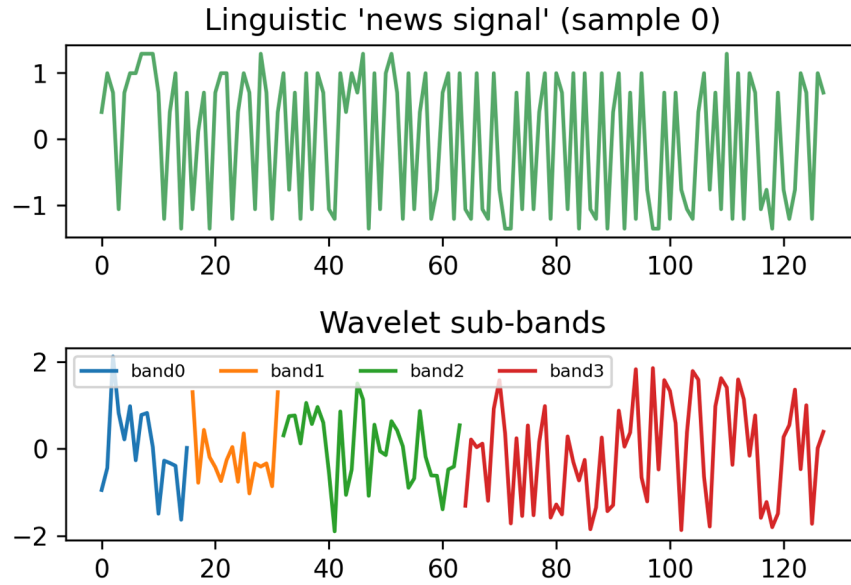


Figure 2: A document as a 1-D linguistic signal (top) and its wavelet sub-bands (bottom).

3.4 Sequence Modelling, Fusion and Objective

A BiLSTM [32], reads the wavelet coefficient sequence:

$$[i_t, f_t, o_t] = \sigma(W[h_{t-1}, v_t] + b), c_t = f_t \odot c_{t-1} + i_t \odot \tanh(W_c[h_{t-1}, v_t]), h_t = o_t \odot \tanh(c_t) \quad (6)$$

$$h = [h_L^{(f)}; ; h_L^{(b)}] \quad (7)$$

The BiLSTM output is fused (residual concatenation) with the AE latent z and a contextual sentence embedding e [11], and trained with a joint objective:

$$u = [e; ; z; ; h], \hat{y} = \sigma(W_o u + b_o) \quad (8)$$

$$\mathcal{L} = \text{BCE}(\hat{y}, y) + \lambda \mathcal{L}_{AE} \quad (9)$$

4. EXPERIMENTAL SETUP

4.1 Dataset

We use the public GonzaloA/fake_news dataset, a labelled dataset of real and fake English news articles [33]. We take a class-balanced sample of 10,000 articles (FIGURE 3) and keep only the first 150 tokens in each. The corpus can be found at https://huggingface.co/datasets/GonzaloA/fake_news. It is divided by stratified sampling into three disjoint partitions of 6,500 training, 1,500 validation and 2,000 test articles (65/15/20). All model selection, early stopping and decision-threshold tuning are carried out on the validation partition only; the test partition is used exactly once, for the final results reported in TABLE 2.

We additionally evaluate on the PHEME rumour corpus, a widely used benchmark of source tweets drawn from nine real-world breaking-news events, comprising 6,425 tweets (4,023 non-rumour and 2,402 rumour). Because PHEME is class-imbalanced we report macro-F1 as the primary metric. We use two protocols: a stratified 65/15/20 split repeated over five seeds, and a leave-one-event-out protocol in which the model is trained on eight events and tested on the ninth; one event (ebola-essien) lacks sufficient rumour examples to form a valid test fold and is excluded, leaving eight evaluable events. This protocol prevents the event-vocabulary leakage that inflates random-split scores on this corpus.

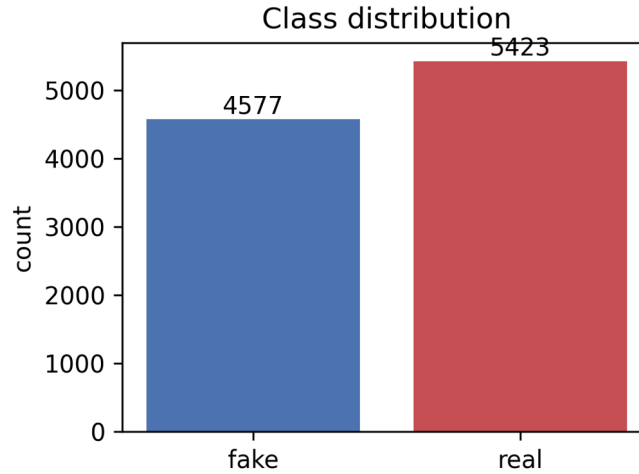


Figure 3: Class distribution of the sampled corpus.

4.2 Implementation and Evaluation Metrics

Linguistic features use spaCy. Semantic features concatenated are word/character TF-IDF reduced with SVD and sentence embeddings (MiniLM) [11]. DWT is performed with pywavelets (db4, 3 levels, L=128) [9]. We optimize using Adam with class-weighted BCE and joint loss from Eq. (9), using early stopping based on validation F1 score and a tuned threshold. We report accuracy, precision, recall, F1, and the area under the ROC curve (AUC):

Reproducibility details. The part-of-speech importance prior $\omega(p)$ in Eq. (1) takes the values PROP 1.0, NOUN 0.9, VERB 0.8, ADJ 0.7, ADV 0.6, NUM 0.5, PRON 0.3, ADP 0.2, DET 0.15, CCONJ 0.15, X 0.1, PUNCT 0.05 and SPACE 0.0. The amplitude sequence is truncated to the first 150 tokens and then resampled and z-normalised to a fixed length $L = 128$ (Eq. 2). The DWT uses Daubechies db4 with three levels. The denoising autoencoder has latent dimension 128 (hidden 256) with input corruption 0.1; the BiLSTM has hidden size 96, is single-layer and bidirectional, with dropout 0.3. Training uses Adam (learning rate $1e-3$, weight decay $1e-5$), batch size 64, and the joint objective of Eq. (9) with reconstruction weight $\lambda = 0.1$. Model selection uses early stopping on validation macro-F1 and the decision threshold is tuned on the validation partition only. The full implementation is publicly available at <https://github.com/Ghaith-Alawady/lsw-net>.

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN}, \text{Precision} = \frac{TP}{TP+FP}, \text{Recall} = \frac{TP}{TP+FN} \quad (10)$$

$$F1 = \frac{2\text{Precision}\times\text{Recall}}{\text{Precision}+\text{Recall}} \quad (11)$$

5. RESULTS AND DISCUSSION

The test results for the 65/15/20 news split are shown in TABLE 2. LSW-Net reaches 96.65% accuracy, 96.92% F1 and 0.994 AUC. On this corpus, however, it does not lead: a TF-IDF + linear SVM baseline attains 98.2% accuracy and 0.9833 F1, outperforming LSW-Net by 1.55 accuracy points and 1.41 F1 points, and both ablations (no-wavelet, no-embeddings) also match or exceed the full model. This corpus is highly separable by lexical content, so a bag-of-words classifier is hard to beat and the wavelet channel adds no accuracy here. Section 5.2 therefore evaluates the model on a corpus where lexical memorisation is prevented.

Table 2: Test results on the fake-news benchmark (stratified 65/15/20 split; 6,500 / 1,500 / 2,000 train/validation/test articles).

Model	Acc.	Prec.	Recall	F1	AUC
TF-IDF + LinearSVM	0.982	0.9888	0.9779	0.9833	-
TF-IDF + LogReg	0.97	0.9732	0.9714	0.9723	-
LSW-Net (no embeddings)	0.971	0.9707	0.976	0.9733	0.9943
LSW-Net (no wavelet)	0.9685	0.9785	0.9631	0.9707	0.994
LSW-Net (full)	0.9665	0.9678	0.9705	0.9692	0.994

5.1 Effect-Size Analysis of the Linguistic Signal

Each wavelet coefficient corresponds to one of these bands, and by the lexicon to a linguistic cause, so that the class-level separation in the wavelet domain can be quantified. There is an L2 divergence of 0.8995 between the class-means (FIGURE 4), and the most discriminative sub-band shows an effect size of Cohen’s $d = -0.197$ (FIGURE 5; TABLE 3). On this corpus the effect sizes are negligible to small ($|d|$ up to 0.197); consistent with the ablation in TABLE 2, the bulk

of the discriminative power on fake-news text comes from the semantic backbone rather than the linguistic-signal channel. The signal channel is particularly useful for humans as it shows where the differences between classes lie.

Table 3: Per-band wavelet energy by class with effect size.

Wavelet band	Class-1 energy	Class-0 energy	Rel. diff.	Cohen's d
band 0	19.2478	18.2567	5.4%	0.064
band 1	11.7166	12.5048	-6.3%	-0.169
band 2	27.0018	28.453	-5.1%	-0.197
band 3	70.0338	68.7576	1.9%	0.093

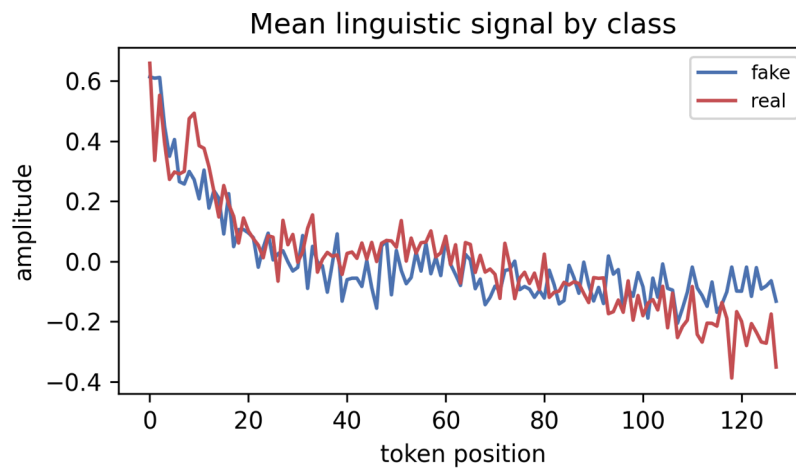


Figure 4: Mean linguistic signal by class.

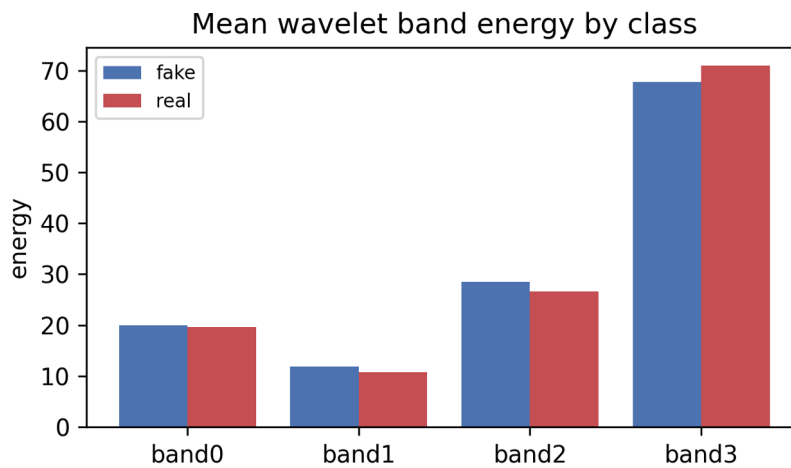


Figure 5: Mean wavelet sub-band energy by class.

Figure 6 - Figure 11 show the metric bargraph, confusion matrix, ROC and PR curves (AUC = 0.994), training history and PCA projection of the AE latent with distinct class separation.

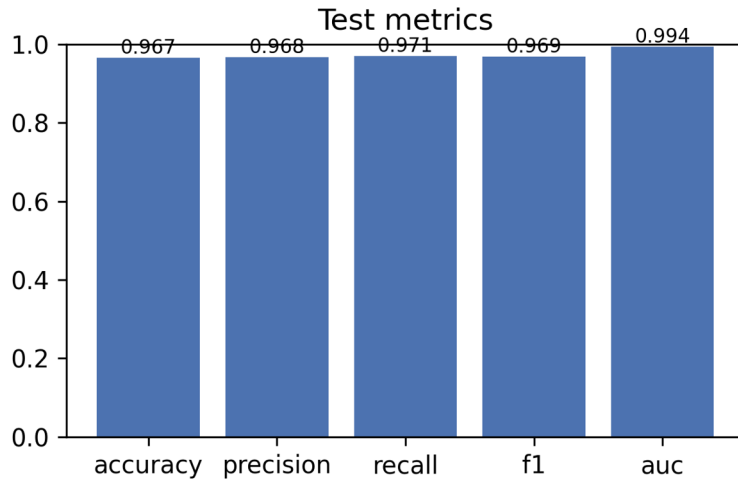


Figure 6: Test metrics of LSW-Net.

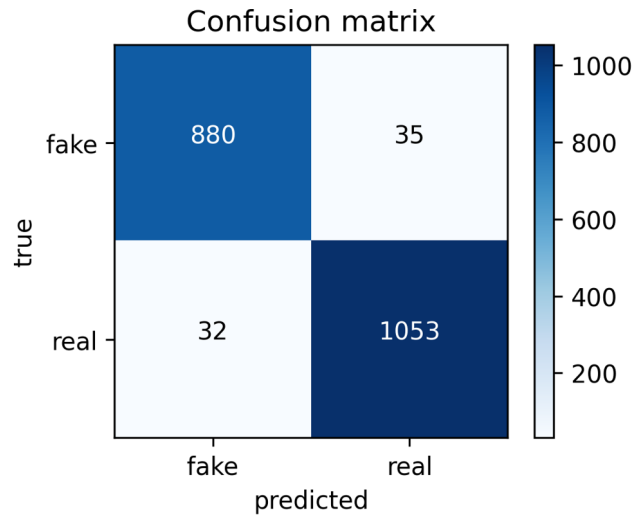


Figure 7: Confusion matrix (tuned threshold).

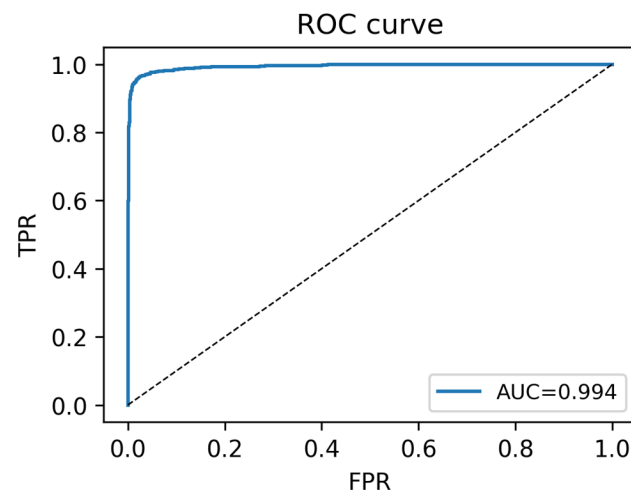


Figure 8: ROC curve.

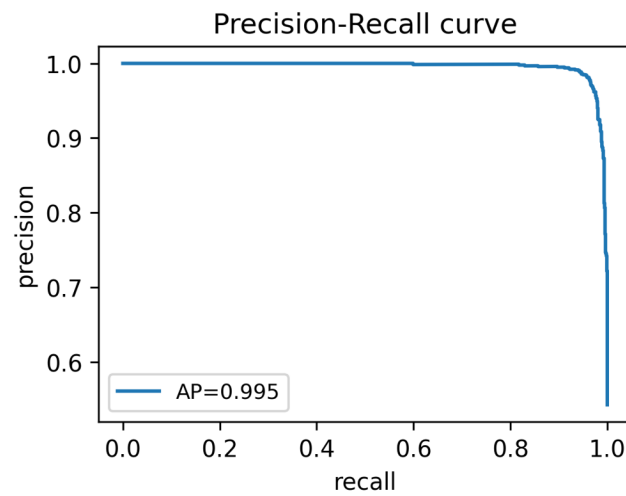


Figure 9: Precision-recall curve.

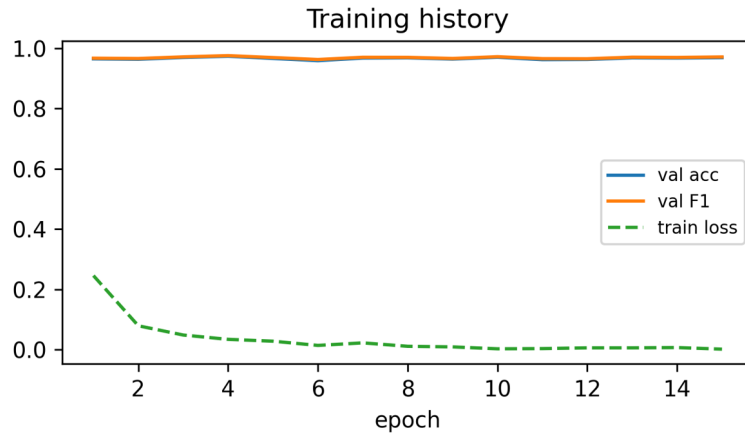


Figure 10: Validation accuracy/F1 and training loss per epoch.

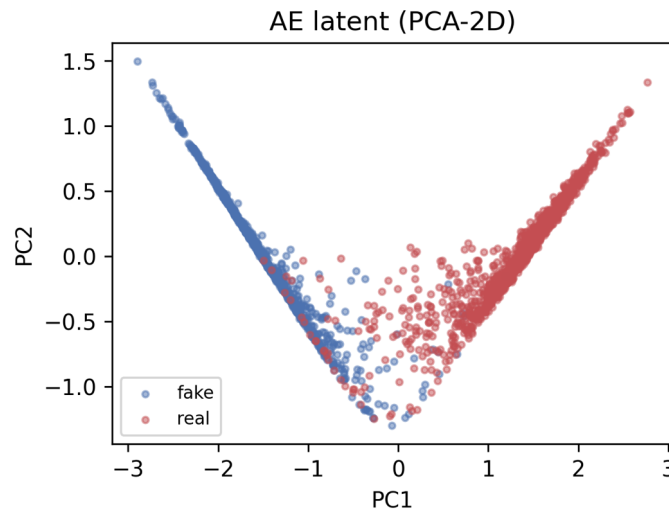


Figure 11: PCA projection of the autoencoder latent by class.

5.2 Cross-Dataset Generalisation (PHEME)

The fake-news corpus above is separable by surface lexical cues, so it cannot reveal whether the linguistic-signal representation captures anything beyond word occurrence. We therefore evaluate on PHEME under the two protocols of Section 4.1. TABLE 4 reports the random-split results over five seeds. All methods cluster near 0.86 macro-F1; LSW-Net (full) is statistically indistinguishable from its no-wavelet ablation (Wilcoxon $p = 0.06$) and from the TF-IDF + SVM baseline ($p = 0.44$). In distribution, therefore, the wavelet channel is neutral, consistent with the news result.

Table 4: PHEME random-split results (stratified 65/15/20; mean $\pm$ standard deviation over five seeds). Differences among the leading methods are not statistically significant.

Model	Accuracy	macro-F1
LSW-Net (full)	0.871 $\pm$ 0.008	0.863 $\pm$ 0.007
LSW-Net (no wavelet)	0.877 $\pm$ 0.005	0.869 $\pm$ 0.005
LSW-Net (no embeddings)	0.867 $\pm$ 0.010	0.859 $\pm$ 0.011
TF-IDF + Linear SVM	0.875 $\pm$ 0.007	0.866 $\pm$ 0.008
TF-IDF + LogReg	0.868 $\pm$ 0.007	0.860 $\pm$ 0.007

The leave-one-event-out results in TABLE 5 tell a different story. When the test event is unseen, the TF-IDF + SVM baseline collapses from 0.866 to 0.496 macro-F1, a drop of 37 points, because it can no longer rely on event-specific vocabulary. LSW-Net (full) degrades far less, from 0.863 to 0.570 (29 points), and outperforms the baseline by 7.5 macro-F1 points. Crucially for our central claim, the full model improves over its no-wavelet ablation in six of the eight evaluable events (FIGURE 12), and the improvement is significant by McNemar’s test in three of them (germanwings-crash $p = 0.03$, ottawa-shooting $p < 0.001$, prince-toronto $p = 0.008$) and significantly worse in none. The aggregate advantage over eight folds is consistent in direction (+2.2 macro-F1 points) though not itself significant given only eight folds (Wilcoxon $p = 0.25$).

Table 5: PHEME leave-one-event-out results (macro-F1) with the score drop from the random split. LSW-Net degrades least on unseen events.

Model	Random F1	macro-F1	LOEO macro-F1	Drop (pts)
LSW-Net (full)	0.863		0.570	29.3
LSW-Net (no wavelet)	0.869		0.549	32.0
TF-IDF + Linear SVM	0.866		0.496	37.0
TF-IDF + LogReg	0.860		0.531	32.9

FIGURE 13 summarises this robustness: the linguistic-signal model loses the least macro-F1 when moving from in-distribution to unseen events. A capacity analysis (reducing the wavelet BiLSTM from 96 to 16 hidden units) did not produce a significant in-distribution gain either (Wilcoxon $p = 0.63$), confirming that the value of the wavelet channel lies specifically in cross-event generalisation rather than in raw accuracy. Taken together, LSW-Net trades a small amount of in-distribution accuracy for a representation that is both inspectable and more robust to topic shift, and helps precisely where lexical memorisation fails rather than in raw in-distribution accuracy.

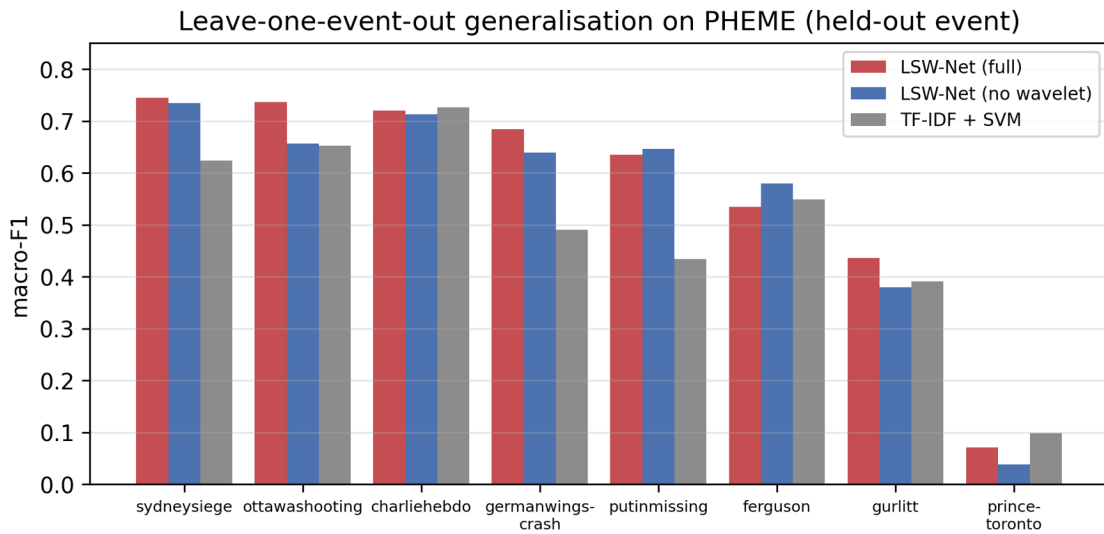


Figure 12: Leave-one-event-out macro-F1 for each held-out PHEME event. LSW-Net (full) improves over its no-wavelet ablation in six of eight events.

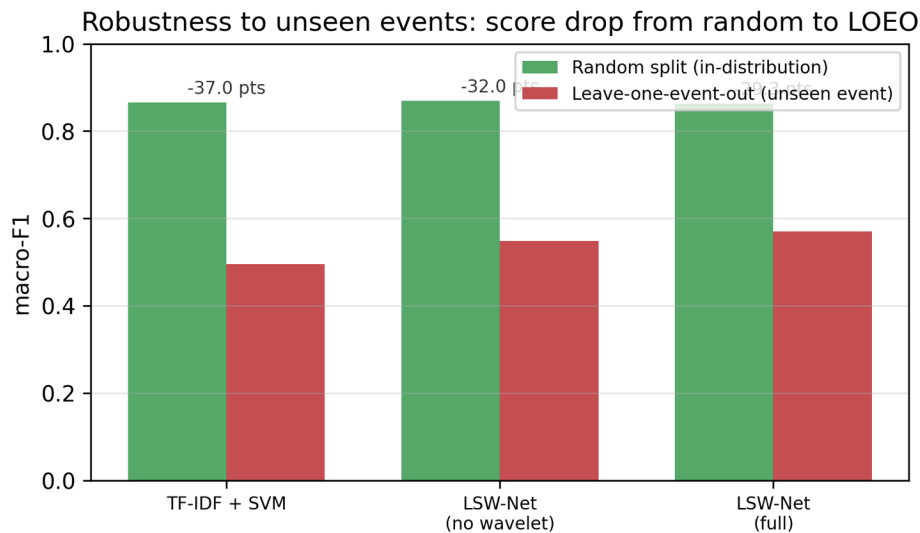


Figure 13: Robustness to unseen events: macro-F1 drop from the random split to leave-one-event-out. The linguistic-signal model degrades least.

6. CONCLUSION AND FUTURE WORK

We present LSW-Net, which builds a linguistically weighted signal (Eq. 1-2), decomposes it with a wavelet transform (Eq. 4-5), and fuses a BiLSTM over the sub-bands with a denoising-autoencoder latent and a sentence embedding (Eq. 6-9). On the GonzaloA/fake_news benchmark it reaches

96.65% accuracy and 0.994 AUC, competitive with but not exceeding a TF-IDF + SVM baseline on that lexically separable corpus. On the PHEME rumour corpus the picture changes: under a random split all methods score near 0.86 macro-F1 and the wavelet channel is neutral, but under a leave-one-event-out protocol LSW-Net attains 0.570 macro-F1 versus 0.496 for TF-IDF + SVM and improves over its no-wavelet ablation in six of eight unseen events, degrading by 29 rather than 37 macro-F1 points from in-distribution to unseen-event evaluation. The wavelet channel therefore contributes robustness to topic shift and an inspectable representation rather than raw in-distribution accuracy. Future work will extend the signal to document-level dynamics and test band-level prediction attribution.

7. DATA AVAILABILITY AND DISCLOSURES

Data availability: The GonzaloA/fake_news dataset is publicly available (see Section 4.1) and the PHEME rumour corpus is publicly available from its original release. The code, figures and metrics are publicly available under the MIT license at <https://github.com/Ghaith-Alawady/lsw-net>. Ethics: no human subjects are involved as the paper uses public datasets. Conflicts of interest: the authors declare none. Funding: none. AI-use disclosure: AI assistance was used for code, drafting and figures; all results and final text were verified by the authors.

References

- [1] Zubiaga A, Liakata M, Procter R, Wong Sak Hoi G, Tolmie P. Analysing How People Orient to and Spread Rumours in Social Media by Looking at Conversational Threads. *PLOS One*. 2016;11(3):e0150989.
- [2] Jia H, Wang H, Zhang X. Early Detection of Rumors Based on Source Tweet-Word Graph Attention Networks. *PLOS One*. 2022;17:e0271224.
- [3] Suthanthira Devi P, Karthika S. Rumor Identification and Verification for Text in Social Media Content. *Comput J*. 2022;65:436-455.
- [4] Chen Y, Hu L, Sui J. Text-Based Fusion Neural Network for Rumor Detection. In: Douligieris C, Karagiannis D, Apostolou D, editors. *Knowledge science, engineering and management (KSEM 2019)*. Berlin: Springer; 2019:105-109.
- [5] Chen X, Zhu D, Lin D, Cao D. Rumor Knowledge Embedding Based Data Augmentation for Imbalanced Rumor Detection. *Inf Sci*. 2021;580:352-370.
- [6] Tu K, Chen C, Hou C, Yuan J, Li J, et al. Rumor2vec: A Rumor Detection Framework With Joint Text and Propagation Structure Representation Learning. *Inf Sci*. 2021;560:137-51.
- [7] Huang Q, Yu J, Wu J, Wang B. Heterogeneous Graph Attention Networks for Early Detection of Rumors on Twitter. In: *International Joint Conference on Neural Networks (IJCNN)*. IEEE. 2020:1-8.
- [8] Lin H, Zhang X, Fu X. A Graph Convolutional Encoder and Decoder Model for Rumor Detection. In: *7th International Conference on Data Science and Advanced Analytics (DSAA)*. 2020:300-306.

- [9] Mallat S. A Wavelet Tour of Signal Processing: The Sparse Way. 3rd ed. Cambridge: Academic Press. 2009.
- [10] Devlin J, Chang MW, Lee K, Toutanova K. BERT: Pre-Training of Deep Bidirectional Transformers for Language Understanding. In: Proceedings of the NAACL-HLT. ACL; 2019:4171-4186.
- [11] Reimers N, Gurevych I. Sentence-BERT: Sentence Embeddings Using Siamese Bert-Networks. In: Proceedings of the EMNLP-IJCNLP. ACL. 2019:3980-3990.
- [12] Sharma D, Srivastava A. Detecting Rumors in Social Media Using Emotion Based Deep Learning Approach. PeerJ Comput Sci. 2024;10:e2202.
- [13] Wu S, Deng Y, Liu J, Luo X, Sun G. Rumor Detection on Social Networks Based on Temporal Tree Transformer. PLOS One. 2025;20:e0320333.
- [14] Cheng M, Nazarian S, Bogdan P. VRoC: Variational Autoencoder-Aided Multi-Task Rumor Classifier Based on Text. Proceedings of the Web Conference. Acad Med. 2020:2892-2898.
- [15] Chang Q, Li X, Duan Z. A Novel Approach for Rumor Detection in Social Platforms: Memory-Augmented Transformer With Graph Convolutional Networks. Knowl-Based Syst. 2024;292:111625.
- [16] Xia J, Li Y, Yu K. LGT: Long-Range Graph Transformer for Early Rumor Detection. Soc Netw Anal Min. 2024;14:100.
- [17] Yadav A, Gupta A. An Emotion-Driven, Transformer-Based Network for Multimodal Fake News Detection. Int J Multimedia Inf Retrieval. 2024;13:7.
- [18] Papageorgiou E, Varlamis I, Chronis C. Harnessing Large Language Models and Deep Neural Networks for Fake News Detection.. Information. 2025;16:297.
- [19] Choi E, Ahn J, Piao X, Kim JK. CroMe: Multimodal Fake News Detection Using Cross-Modal Tri-Transformer and Metric Learning. IEEE Access. 2025;13:197124-197132.
- [20] Sun G, Li J, Zhang Z. Large Language Model Guided Graph Capsule Network With Local-Enhanced Alignment for Fake News Detection. Inf Process Manage. 2026;63:104900.
- [21] Zheng P, Huang Z, Dou Y, Yan Y. Rumor Detection on Social Media Through Mining the Social Circles With High Homogeneity. Inf Sci. 2023;642:119083.
- [22] Cui W, Shang M. KAGN: Knowledge-Powered Attention and Graph Convolutional Networks for Social Media Rumor Detection. J Big Data. 2023;10:45.
- [23] Pattanaik B, Mandal S, Tripathy RM, Sekh AA. Rumor Detection Using Dual Embeddings and Text-Based Graph Convolutional Network. Discov Artif Intell. 2024;4:86.
- [24] Lv S, Zhang H, He H, Chen B. Microblog Rumor Detection Based on Comment Sentiment and CNN-LSTM. In: Liang Q, Wang W, Mu J, Liu X, Na Z, Chen B, editors. Artificial intelligence in China. Berlin: Springer. 2020:148-156.
- [25] Zhang H, Fang Q, Qian S, Xu C. Multi-Modal Knowledge-Aware Event Memory Network for Social Media Rumor Detection. Proceedings of the 27th ACM International Conference on Multimedia. Acad Med. 2019:1942-1951.

- [26] ZZubiaga A, Liakata M, Procter R. Exploiting Context for Rumour Detection in Social Media. In: Ciampaglia GL, Mashhadi A, Yasseri T, editors. Social informatics, Lecture Notes in Computer Science. Vol. Berlin: Springer. 2017;10539:109-123.
- [27] Kochkina E, Liakata M, Zubiaga A. All-In-One: Multi-Task Learning for Rumour Verification. In: Bender EM, Derczynski L, Isabelle P, editors. Proceedings of the 27th International Conference on Computational Linguistics. 2018. p. 3402-3413.
- [28] Kumar S, Carley KM. Tree LSTMs With Convolution Units to Predict Stance and Rumor Veracity in Social Media Conversations. In: Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics. 2019:5047-5058.
- [29] Wei P, Xu N, Mao W. Modeling Conversation Structure and Temporal Dynamics for Jointly Predicting Rumor Stance and Veracity. In: Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP). Association for Computational Linguistics. 2019:4786-4797.
- [30] Wei L, Hu D, Zhou W, Yue Z, Hu S. Towards Propagation Uncertainty: Edge-Enhanced Bayesian Graph Convolutional Networks for Rumor Detection. In: Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers). Association for Computational Linguistics. 2021:3845-3854.
- [31] Khoo LM, Chieu HL, Qian Z, Jiang J. Interpretable Rumor Detection in Microblogs by Attending to User Interactions. Proceedings of the Aaai Conference on Artificial Intelligence. 2020;34:8783-8790.
- [32] Hochreiter S, Schmidhuber J. Long Short-Term Memory. Neural Comput. 1997;9:1735-1780.
- [33] https://huggingface.co/datasets/GonzaloA/fake_news

Abbreviations

Abbr.	Meaning	Abbr.	Meaning
LSW-Net	Linguistic-Signal Wavelet Network	AE	Autoencoder
LSTM	Long Short-Term Memory	DWT	Discrete Wavelet Transform
TF-IDF	Term Frequency-Inverse Document Frequency	BiLSTM	Bidirectional LSTM
SVM	Support Vector Machine	SVD	Singular Value Decomposition
BCE	Binary Cross-Entropy	POS	Part-of-Speech
ROC	Receiver Operating Characteristic	AUC	Area Under the ROC Curve
F1	Harmonic mean of precision and recall	PR	Precision-Recall
		PCA	Principal Component Analysis

Abbreviations used throughout this paper.

Glossary of Terms

Term	Meaning
Fake news	A news article containing false or misleading information presented as legitimate journalism.
Linguistic signal	The 1-D numeric sequence defined in Eq. (1)-(2): reading a document token by token and emitting, for each token, an amplitude set by its POS, term salience and sentiment.
Wavelet sub-band	A frequency/resolution component produced by the DWT: the approximation band (coarse trend) and the detail bands (fine, local variation).
Autoencoder	A neural network that compresses an input to a latent code and reconstructs it.
Cohen's d	A standardised effect size measuring the separation between two groups' means in units of pooled standard deviation; $ d < 0.2$ negligible, 0.2-0.5 small, 0.5-0.8 medium.

Meaning of key terms and phrases used in this paper.