

# Journal of Dinda

Data Science, Information Technology, and Data Analytics

Vol. 6 No. 2 (2026) 106 - 113

E-ISSN: 2809-8064

## Clickbait News Classification for Indonesian Headlines Using the Mamba Selective State Space Model

Ridhwan Ardiyansyah<sup>1</sup>, Teny Handhayani<sup>2\*</sup>, Manatap Dolok Lauro<sup>3</sup>

<sup>1,2\*,3</sup>Dept. of Informatics Engineering, Universitas Tarumanagara, Jakarta, Indonesia

<sup>1</sup>ridhwan.535220114@stu.untar.ac.id, <sup>2\*</sup>tenyh@fti.untar.ac.id, <sup>3</sup>manataps@fti.untar.ac.id

### Abstract

Clickbait headlines remain a persistent problem in Indonesian digital journalism, where sensational or intentionally misleading titles conceal essential information to inflate page traffic, steadily eroding public trust in online media and weakening readers' information literacy. This study introduces a real-time, web-based application for clickbait detection system that adopts the Mamba architecture, a Selective State Space Model, as a single end-to-end method for judging the consistency between an Indonesian news headline and its body. In contrast to conventional Transformer encoders, whose self-attention grows quadratically with sequence length, Mamba carries out sequential modelling in linear time, and its selective scanning mechanism adaptively keeps semantically relevant tokens while discarding uninformative ones, allowing long-range dependencies between title and content to be captured efficiently. The text is tokenized with the pretrained IndoBERTweet model, whose vocabulary mirrors the informal Indonesian online ecosystem of slang and abbreviations, and punctuation is intentionally retained as an additional linguistic cue. The model was trained and validated on the labelled CLICK-ID corpus of headline-content pairs and assessed through three-fold cross-validation together with a confusion matrix spanning accuracy, precision, recall, and F1-score. Averaged over the three folds, the Mamba classifier obtained an accuracy of 73.43%, precision of 71.31%, recall of 61.15%, and F1-score of 65.65%, with the best result on the second fold, indicating stable generalization to unseen news data. The system is deployed as an interactive web application that classifies a news item directly from its URL, showing that a single state-space sequence model is a viable and computationally efficient option for Indonesian clickbait detection.

Keywords: clickbait detection, deep learning, indoBERTweet, mamba, selective state space model

© 2026 Journal of DINDA

### 1. Introduction

The shift of news consumption from printed media toward digital platforms has transformed the way society obtains information, making online news portals a principal channel for daily public knowledge. Because reader traffic is tightly coupled to advertising income, this transition has sharpened the competition among media outlets to attract visitors. Social medias are one of the main ways for the Indonesian to find news [1]. To secure clicks, many newsrooms turn to clickbait, namely headlines deliberately written to be sensational, ambiguous, or provocative, or that withhold essential information so as to exploit the reader's curiosity [2]. Such headlines may raise engagement in the short term, but they steadily lower journalistic quality, hasten the spread of misinformation, and leave readers dissatisfied when the article does not fulfil the promise of its title [3].

The magnitude of this issue is far from trivial. With social media now serving as the dominant gateway

through which Indonesians encounter news, a single deceptive headline can spread within minutes and influence public opinion well before the corresponding article is actually read. Human moderation cannot match this pace, which underscores the need for an automated, language-aware detection approach able to flag misleading headlines at the very moment they are consumed.

Framed this way, the problem imposes three concrete requirements on any practical detector. It must operate in real time, so the computational cost of inference has to stay low even when a full article body is processed together with its headline. It must cope with the informal register of Indonesian online text, which calls for a vocabulary learned from that register rather than from formal writing. Finally, it must reason over the relationship between a short title and a long body, which places a premium on architectures able to model long-

Received: 26-07-2026 | Revised: 01-09-2026 | Published: 10-09-2026

range dependencies efficiently. These requirements jointly shape the design choices made in this study.

Within the Indonesian setting, the task is made harder by the informal and fast-changing register of online media, where slang, abbreviations, and emphatic punctuation appear frequently. A number of studies have tackled Indonesian clickbait detection through machine learning and deep learning [4]. Earlier solutions depended on recurrent architectures such as Long Short-Term Memory networks [5], while subsequent research has fine-tuned Transformer-based encoders, among them BERT [6], RoBERTa [7], and DistilBERT [8], together with IndoBERT variants adapted to Indonesian [9]. Although these models attain competitive accuracy, their underlying self-attention grows quadratically with sequence length, which constrains efficiency once an entire article body must be encoded jointly with its headline. This paper also detects clickbait news not only from the headline but also by considering the content of the news.

More recently, State Space Models have arisen as a strong substitute for Transformers in sequence modelling [10]. Within this family, Mamba contributes a selective scanning mechanism that lets the model retain or discard information depending on the input, attaining linear-time complexity while still capturing long-range dependencies [11]. State-space formulations have likewise proven effective for classifying long documents [12]. This efficiency is especially appealing for clickbait detection, in which the discriminative signal lies in the semantic gap between a brief, manipulative headline and a comparatively long article body.

The contributions of this paper are: Introduces Mamba as an end-to-end model for Indonesian clickbait detection by measuring headline-content consistency. Uses IndoBERTweet tokenization while retaining punctuation to capture informal language and clickbait cues. Implements the model as a web-based system that detects clickbait directly from news URLs.

Beyond its technical merit, this work is related to several United Nations Sustainable Development Goals (SDGs). By helping readers verify the credibility of a headline before consuming an article, the system supports SDG 4 (Quality Education) through improved information literacy, SDG 9 (Industry, Innovation, and Infrastructure) through the application of modern artificial intelligence, and SDG 16 (Peace, Justice, and Strong Institutions) by curbing the circulation of misleading information. Accordingly, the objective of this study is to design and implement a real-time, web-based Indonesian clickbait detector that employs the Mamba Selective State Space Model as a single classification method, combined with IndoBERTweet tokenization, and to evaluate its effectiveness through cross-validation.

The remainder of this paper is organized as follows. Section II reviews prior work on clickbait detection, from feature-based classifiers to Transformer encoders and recent state-space formulations. Section III describes the dataset, the preprocessing and tokenization stages, the Mamba architecture, and the evaluation protocol. Section IV reports the cross-validation results, discusses their implications, and sets out the limitations of the study, while Section V presents the conclusions and directions for future research.

## 2. Related Work

Work on clickbait detection has progressed in step with developments in natural language processing. The earliest systems treated the problem as ordinary text classification, feeding manually engineered lexical and syntactic features into classical machine-learning classifiers, a line of work surveyed thoroughly for the Indonesian context in [3]. While such feature-based pipelines are interpretable, they struggle to model the semantic discrepancy between a headline and its body that often betrays clickbait.

Deep learning later became the prevailing paradigm. Recurrent networks, most notably Long Short-Term Memory architectures, were applied to Indonesian clickbait headlines with promising outcomes [5], and broader deep-learning formulations were investigated as well [4]. The introduction of Transformer encoders advanced the field further still: fine-tuned BERT [6], RoBERTa [7], and DistilBERT [8], along with IndoBERT variants adapted to Indonesian [9], consistently raised accuracy by leveraging contextual representations learned from large corpora. These improvements, however, are accompanied by the quadratic cost inherent to self-attention.

Most recently, State Space Models have been put forward as an efficient counterpart to attention-based architectures for sequence modelling [10], with proven success on long-document classification [12]. Mamba extends the state-space formulation with a selective mechanism that reaches linear-time complexity while preserving the capacity to model long-range dependencies [11]. To the best of our knowledge, employing a Mamba-based model as a self-contained classifier for Indonesian clickbait detection, combined with IndoBERTweet tokenization and delivered as a real-time web service, has not been reported before; this is precisely the gap that the present study addresses.

Two design implications follow from this review. First, the linguistic register of Indonesian online media argues for a tokenizer whose vocabulary was learned from informal text rather than from formal corpora, which motivates the adoption of IndoBERTweet in this study. Second, because the discriminative signal for clickbait resides in the relationship between a short headline and

a much longer body, the chosen architecture must remain tractable on extended sequences, which is exactly the property that distinguishes the selective state-space formulation from attention-based encoders.

### 3. Methods

The proposed system is organized as a sequential pipeline that converts raw Indonesian news text into a binary clickbait decision. As outlined in Figure 1, the overall research methodology consists of data acquisition, preprocessing and tokenization, embedding, sequential modelling with Mamba, classification, and evaluation.

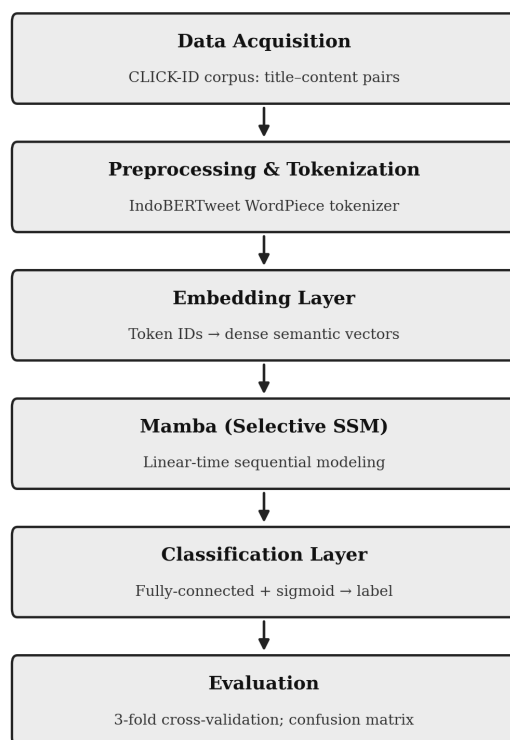


Figure 1. Overview of the research methodology pipeline for the Mamba-based clickbait detection system.

#### 3.1. Dataset and Preprocessing

This study utilizes a secondary dataset known as CLICK-ID—a specialized corpus for clickbait detection in Indonesian-language news—combined with data independently scraped from various national online news portals. The dataset comprises a total of 14,701 pairs of headlines and news content, annotated by linguists using a binary classification (1 for clickbait, 0 for non-clickbait). The class distribution within the dataset consists of 41.59% clickbait and 58.41% factual content. Model development relied on the labelled CLICK-ID corpus, which provides thousands of Indonesian news samples, each pairing a headline with its article body and a binary clickbait label. For every

entry, the headline and body were joined into one input string so that the model could learn the correspondence between them. Tokenization was carried out with the pretrained IndoBERTweet model through the WordPiece subword algorithm [13], [14]. Since IndoBERTweet was trained on informal Indonesian social-media text, its vocabulary already reflects the slang and abbreviations that often surface in clickbait titles. Departing from common preprocessing convention, punctuation marks were deliberately kept, because exclamation marks and ellipses frequently convey the sensational tone typical of clickbait.

The decision to model the headline and the body as a single sequence follows directly from the nature of the problem. A clickbait title is deceptive not in isolation but in relation to the article it introduces; the misleading effect arises from the distance between what the title promises and what the body actually delivers. A classifier that inspects the headline alone is therefore restricted to surface cues such as word choice and punctuation, whereas the joined representation places the title and its article within the same context and allows the sequence model to assess their semantic agreement directly.

Every token sequence is normalized to a fixed length of 256 positions. Sequences below this threshold are zero-padded, and an attention mask is generated that assigns one to authentic tokens and zero to padding, ensuring that padded positions are disregarded during computation. The IndoBERTweet vocabulary used for this mapping comprises 32,032 sub-word units, and each token identifier is transformed into a 256-dimensional embedding vector before entering the sequential model. The binary clickbait label of each record is represented as an integer target, producing a supervised configuration well suited to end-to-end training.

As an illustration of the tokenization step, consider a sensational Indonesian headline fragment such as “Bikin Melongo”. The WordPiece algorithm breaks such expressions down into known sub-word units, so that rare or informal terms are represented as combinations of frequent fragments instead of being dropped as out-of-vocabulary tokens. This sub-word strategy is indispensable for Indonesian online text, in which neologisms, abbreviations, and code-mixed expressions are common, and it ensures that the model still receives a meaningful representation even for words it has not encountered before.

#### 3.2. Mamba Selective State Space Model

Following tokenization, an embedding layer maps each token identifier to a dense vector, yielding a matrix that retains the semantic relationships among words [15]. The embedded sequence is then passed to the Mamba block, which is founded on the Selective State Space

Model. A continuous state-space system is discretized and rendered input-dependent, so that the parameters dictating how information is kept or discarded change from one token to the next. This selective behaviour allows the model to foreground discriminative cues, such as a mismatch between an exaggerated title and an unremarkable body, while overlooking uninformative tokens. The recurrence at the heart of the state-space formulation can be written as in (1),

$$ht = Aht - I + Bxt, \quad yt = Cht \quad (1)$$

where  $x(t)$  denotes the input token representation at step  $t$ ,  $h(t)$  the hidden state,  $y(t)$  the output, and  $A$ ,  $B$ , and  $C$  the input-dependent state-space parameters. Since this formulation runs in linear time with respect to sequence length, it handles the combined headline-body input far more efficiently than quadratic self-attention. The implementation builds on the official mamba-ssm library [11].

In architectural terms, the Mamba block handles the embedded sequence through several coordinated stages. An input projection first widens each token representation, after which a lightweight causal convolution extracts local patterns among adjacent subwords. The selective state-space scan subsequently models the sequence at a global level, relying on input-dependent gates to determine how much of each token's information should flow into the hidden state. A residual connection together with layer normalization stabilizes training, and the concluding hidden state is pooled into a single vector that encapsulates the relationship between the headline and the body.

The formal underpinning of this behaviour lies in the discretization of (1): a step-size parameter converts the continuous system into a discrete recurrence, and the matrices governing the state transition are expressed as functions of the present input. This dependence on the input is the core of the selective mechanism: rather than imposing the same fixed dynamics on every token, the model adjusts its memory position by position, strengthening the contribution of informative tokens while suppressing that of filler words. The hidden states are produced sequentially but still lend themselves to hardware-level parallelization, which keeps training efficient.

### 3.3. Classification and Evaluation

The sequence representation generated by the Mamba block is forwarded to a fully connected layer and then a sigmoid activation, which outputs the probability that the input is clickbait. To gauge generalization under a limited sample size, the system was evaluated with three-fold cross-validation [16]. At each iteration the dataset is split into three folds of equal size; two folds serve for training and the remaining one for testing, and

the cycle repeats until every fold has acted once as the test set. A fixed random seed was applied to ensure reproducibility and a fair comparison across folds. Performance was quantified through a confusion matrix, from which accuracy, precision, recall, and F1-score were obtained, and the reported values correspond to the average across the three folds.

The four-evaluation metrics follow their standard definitions, derived from the confusion-matrix counts of true positives, true negatives, false positives, and false negatives. Accuracy is the share of correctly classified headlines across all samples; precision is the proportion of headlines predicted as clickbait that are truly clickbait; recall is the proportion of actual clickbait headlines the model manages to identify; and the F1-score, as the harmonic mean of precision and recall, provides a single balanced indicator that is particularly informative when the two classes are unevenly distributed.

Cross-validation was preferred over a single hold-out split for two reasons. Every sample takes part in both training and testing across the procedure, so the available data are used to their full extent, and the averaging of scores over three independent partitions reduces the chance that one fortunate or unfortunate split distorts the assessment of the model. Both considerations are especially relevant when the sample size is limited relative to the capacity of a deep sequential model.

### 3.4. Training Configuration

The complete network, encompassing the embedding layer, the Mamba block, and the classification head, was implemented in PyTorch and trained end-to-end using the standard Adam optimizer, requiring no specialized setup beyond that of an ordinary neural network. The main hyperparameters, namely the learning rate, the batch size, and the number of training epochs, were chosen with the aid of the cross-validation procedure so that the selected values reflected consistent behaviour across folds rather than one fortunate split. Because Mamba's selective scan remains parallelizable on modern graphics hardware despite each step being input-dependent, both training and inference were performed efficiently on a GPU.

### 3.5. System Implementation

The trained model was embedded in an interactive web application so that the detector could be put to practical use, and its overall structure is summarized in Figure 2. On the server side, the Mamba classifier is exposed through a lightweight service that accepts either a news URL or raw text; when a URL is provided, the article title and body are fetched automatically via web scraping, concatenated, tokenized, and submitted to the model, which returns a clickbait probability alongside a

binary label in real time. On the client side, the outcome is shown through a straightforward interface that demands no technical expertise from the reader, consistent with the aim of fostering everyday information literacy. The overall design adheres to established machine-learning operations practices so that training, inference, and deployment remain reproducible [17].

Keeping the inference logic on the server side also simplifies maintenance. Because the model weights reside in a single central service, an improved version of the classifier can replace the current one without requiring any change on the client, and every user immediately benefits from the update. This separation between the presentation layer and the prediction service mirrors the operational practices referred to above and keeps the application straightforward to monitor and extend.

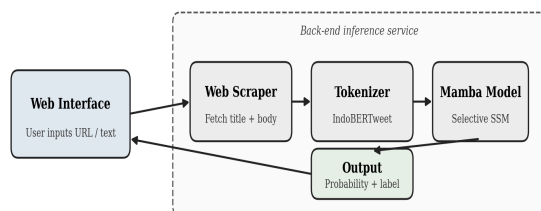


Figure 2. High-level architecture of the clickbait detection web service deployment.

#### 4. Results and Discussion

Table 1 presents the performance of the Mamba classifier on each of the three folds along with the averaged figures. The model reached a mean accuracy of 73.43% and behaved consistently from one fold to another, which confirms that the state-space approach generalizes dependably to Indonesian news data not seen during training.

| Fold    | Accuracy (%) | Precision (%) | Recall (%) | F1-Score (%) |
|---------|--------------|---------------|------------|--------------|
| 1       | 72.09        | 66.55         | 64.36      | 65.44        |
| 2       | 74.76        | 72.04         | 63.14      | 67.30        |
| 3       | 73.45        | 75.34         | 55.94      | 64.21        |
| Average | 73.43        | 71.31         | 61.15      | 65.65        |

The strongest configuration appeared on the second fold, attaining an accuracy of 74.76% and an F1-score of 67.30%. The comparatively high precision, reaching 75.34% on the third fold, indicates that headlines flagged as clickbait are correct in most instances. The lower mean recall of 61.15%, by contrast, shows that a fraction of genuine clickbait headlines still goes undetected, especially those worded in a formal yet subtly manipulative way. These per-fold scores are also presented graphically in Figure 3.

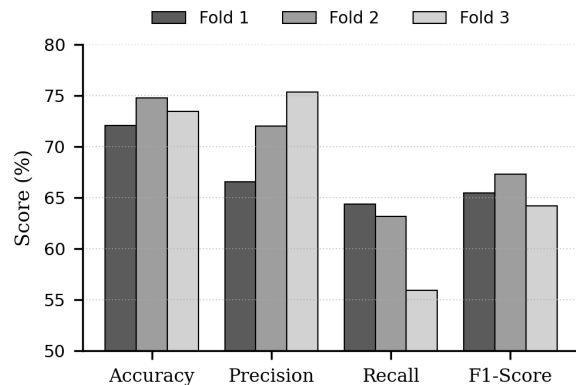


Figure 3. Accuracy, precision, recall, and F1-score of the Mamba classifier for each fold.

The variation between folds also deserves attention. Accuracy remains within a narrow band, from 72.09% on the first fold to 74.76% on the second, a spread of less than three percentage points that points to stable behaviour under different data partitions. Recall fluctuates more widely, falling to 55.94% on the third fold even as precision on that same fold rises to 75.34%. This inverse movement suggests that the composition of each partition shifts the effective decision boundary: when the test fold contains a larger share of subtly worded clickbait, the model becomes more conservative, trading missed detections for fewer false alarms.

A more detailed look at the confusion matrix accounts for this pattern. The model correctly distinguished most factual headlines from manipulative ones, producing high counts of true negatives and true positives, yet it also generated a non-negligible number of false negatives, namely authentic clickbait headlines judged to be legitimate. These errors clustered around titles that avoid sensational punctuation and instead depend on subtle semantic exaggeration, which is exactly the category that most strains lexical cues.

Functional testing supplemented the quantitative assessment. Black-box testing verified that every feature of the system, from URL-based scraping to real-time inference, operated in accordance with its specification, while a User Acceptance Test produced a uniformly valid result, signalling that end users regarded the web application as responsive and easy to use. The intentional preservation of punctuation proved advantageous, as it made the model attentive to the emphatic punctuation characteristic of sensational titles.

The User Acceptance Test involved a panel of general readers who interacted with the live web application. Their responses were consistently favourable in terms of responsiveness and ease of use, and each functional requirement was judged valid, confirming that the system is operable by non-technical users. Qualitative observations during the test nonetheless revealed two

practical shortcomings: the URL-based scraper occasionally could not retrieve content from sites with non-standard markup, and a handful of grammatically formal but semantically deceptive headlines were misclassified, echoing the recall gap seen in the quantitative results.

Placed against the wider literature on Indonesian clickbait, these results are competitive with earlier single-model studies built on recurrent or Transformer encoders [5]-[9], while adding the benefit of linear-time inference. The relatively lower recall aligns with a challenge documented throughout the literature [3], namely identifying headlines whose manipulation is semantic rather than purely lexical.

For a reader-facing tool, the model's precision-oriented behaviour is advantageous, since a false alarm that wrongly brands a legitimate headline as clickbait undermines user trust more than the occasional missed borderline case. For this reason, the web deployment displays the predicted probability together with the label, enabling readers to treat borderline scores with suitable caution.

The efficiency of the approach is a central consideration. Conventional Transformer encoders bear a cost that grows quadratically with the number of tokens, which becomes onerous once a complete article body is appended to its headline. The state-space formulation used here, in contrast, scales linearly with sequence length, so the combined headline-body input can be processed with markedly lower computational overhead while still modelling the long-range semantic gap that separates clickbait from factual reporting. This characteristic makes the method well suited to real-time deployment, where a continuous stream of news items must be screened.

Table 2 positions the proposed model among the main families of architectures that have been applied to clickbait detection. The comparison brings out the key trade-off that motivates this work: although recurrent and Transformer encoders have each been used successfully, only the selective state-space formulation pairs linear-time processing with the ability to capture the long-range headline-body relationship, and it is this property that the present study exploits.

It should be stressed that Table 2 contrasts architectural properties rather than reported scores. The cited studies were conducted on different datasets, splits, and preprocessing pipelines, so their published figures are not directly commensurable with the results in Table I; a strictly controlled benchmark of Mamba against Transformer baselines under identical experimental conditions is therefore left to future work.

Taken together, the findings support the claim that Mamba can act as a self-contained classifier for

Indonesian clickbait detection. Its linear-time complexity provides a tangible efficiency advantage over Transformer encoders, while IndoBERTtweet tokenization equips the model to handle the informal linguistic register of Indonesian online media. The main limitation concerns recall, which suggests that the decision boundary could be refined to capture more disguised forms of clickbait.

Table 2. Comparison of Sequence Architectures for Clickbait Detection

| Approach                                | Representative work                           | Complexity |
|-----------------------------------------|-----------------------------------------------|------------|
| Recurrent (LSTM)                        | Siregar et al. [5]                            | Linear     |
| Transformer (BERT, RoBERTa, DistilBERT) | Putri & Pratomo [6];<br>Sirusstara et al. [7] | Quadratic  |
| Indonesian Transformer (IndoBERT)       | Kurniawan et al. [9]                          | Quadratic  |
| Selective SSM (Mamba)                   | This work                                     | Linear     |

These observations also carry practical implications for deployment at scale. Because real-time clickbait screening must keep pace with the continuous flow of online publishing, the linear-time inference of the state-space model translates directly into lower latency and reduced hardware demands compared with attention-based encoders of comparable capacity. This efficiency, together with a tokenizer attuned to informal Indonesian, makes the approach a pragmatic foundation for browser extensions, editorial dashboards, or platform-level moderation tools intended to support healthier information consumption.

Because precision and recall react differently to the decision threshold imposed on the sigmoid output, the operating point can be adjusted to match the intended application. A reader-facing tool may prefer higher precision to limit false alarms, whereas an editorial auditing tool that screens large volumes of headlines might favour higher recall so as to surface as many suspect titles as possible for human review. Reporting the predicted probability instead of the binary label alone lets downstream users choose a threshold aligned with their tolerance for each type of error.

#### 4.1. Limitations

Several limitations qualify these findings. First, the mean recall of 61.15% indicates that the classifier still misses a meaningful share of clickbait, particularly headlines that are grammatically restrained yet semantically deceptive; lifting recall without sacrificing precision is therefore the most urgent improvement. Second, although the dataset is sizeable, it originates from a finite set of sources and may not represent the full stylistic range of Indonesian online media, which can restrict generalization to newly emerging clickbait patterns. Third, because the URL-based scraper relies on

the structure of each target website, pages with unconventional markup occasionally block automatic content retrieval and necessitate manual text entry. Addressing these issues through richer and better-balanced data, more robust scraping, and calibration of the decision threshold would strengthen the system further.

## 5. Conclusion

In conclusion, this study designed and implemented a real-time, web-based clickbait detection system for Indonesian news headlines that uses the Mamba Selective State Space Model as its sole classification method. The goal of showing that a linear-time state-space model can effectively classify Indonesian clickbait was met: under three-fold cross-validation the model achieved an average accuracy of 73.43% and an F1-score of 65.65%, while remaining computationally efficient and accessible through a URL-based web interface.

Beyond the quantitative metrics, functional evaluation confirmed the practical viability of the system. Black-box testing validated every specified feature, from URL-based content retrieval to real-time inference, and the User Acceptance Test returned uniformly valid judgements, indicating that the application can be operated comfortably by general readers without a technical background.

Future work should concentrate on improving recall so that more subtly phrased clickbait is captured, for instance by enlarging and balancing the training corpus, examining multimodal features that combine textual and visual cues, and benchmarking Mamba against state-of-the-art Transformer baselines under identical conditions. Such directions would further reinforce the reliability of automated clickbait detection and its contribution to a healthier digital information ecosystem.

## References

- [1] F. L. Gaol, A. Maulana, and T. Matsuo, "News consumption patterns on Twitter: fragmentation study on the online news media network," *Heliyon*, vol. 6, no. 10, p. e05169, Oct. 2020, doi: 10.1016/j.heliyon.2020.e05169.
- [2] A. K. Jung, S. Stieglitz, T. Kissmer, M. Mirbabaie, and T. Kroll, "Click me. . .! The influence of clickbait on user engagement in social media and the role of digital nudging," *PLoS One*, vol. 17, no. 6, June, pp. 1–22, 2022, doi: 10.1371/journal.pone.0266743.
- [3] N. A. Zuhroh and N. A. Rakhmawati, "Clickbait detection: A literature review of the methods used," *Register: Jurnal Ilmiah Teknologi Sistem Informasi*, vol. 6, no. 1, pp. 1–10, 2020, doi: 10.26594/register.v6i1.1561.
- [4] A. Chowanda, Nadia, and L. M. M. Kolbe, "Identifying clickbait in online news using deep learning," *Bulletin of Electrical Engineering and Informatics*, vol. 12, no. 3, pp. 1755–1761, 2023, doi: 10.11591/eei.v12i3.4444.
- [5] B. Siregar, I. Habibie, E. B. Nababan, and Fahmi, "Identification of Indonesian clickbait news headlines with long short-term memory recurrent neural network algorithm," *J. Phys. Conf. Ser.*, vol. 1882, no. 1, 2021, doi: 10.1088/1742-6596/1882/1/012129.
- [6] Diyah Utami Kusumaning Putri and Dinar Nugroho Pratomo, "Clickbait Detection of Indonesian News Headlines using Fine-Tune Bidirectional Encoder Representations from Transformers (BERT)," *Inform : Jurnal Ilmiah Bidang Teknologi Informasi dan Komunikasi*, vol. 7, no. 2, pp. 162–168, 2022, doi: 10.25139/inform.v7i2.4686.
- [7] J. Sirusstara, N. Alexander, A. Alfarisy, S. Achmad, and R. Sutoyo, "Clickbait Headline Detection in Indonesian News Sites using Robustly Optimized BERT Pre-training Approach (RoBERTa)," in *2022 3rd International Conference on Artificial Intelligence and Data Sciences (AiDAS)*, 2022, pp. 1–6. doi: 10.1109/AiDAS56890.2022.9918678.
- [8] R. T. Dewangga and B. Prasetyo, "Comparative Analysis of BERT Model and DistilBERT Model for Enhanced Clickbait Headline Structure Detection in Indonesian Online News," *JUITA: Jurnal Informatika*, vol. 13, no. 3, pp. 235–244, 2025, doi: 10.30595/juita.v13i3.26479.
- [9] S. Kurniawan, A. S. Pramayoga, Y. F. Ashari, and Muhammad Afrizal Amrustian, "Development and Evaluation of an IndoBERT-Based NLP Model for Automated Clickbait Detection," *Advance Sustainable Science Engineering and Technology*, vol. 8, no. 1, p. 02601021, Jan. 2026, doi: 10.26877/asset.v8i1.2637.
- [10] X. Wang *et al.*, "State Space Model for New-Generation Network Alternative to Transformers: A Survey," Apr. 2024, [Online]. Available: <http://arxiv.org/abs/2404.09516>
- [11] A. Gu and T. Dao, "Mamba: Linear-Time Sequence Modeling with Selective State Spaces," May 2024, [Online]. Available: <http://arxiv.org/abs/2312.00752>

- [12] B. Song, Y. Xu, P. Liang, and Y. Wu, “State Space Models Based Efficient Long Documents Classification,” *Journal of Intelligent Learning Systems and Applications*, vol. 16, no. 03, pp. 143–154, 2024, doi: 10.4236/jilsa.2024.163009.
- [13] F. Koto, J. H. Lau, and T. Baldwin, “INDOBERTWEET: A Pretrained Language Model for Indonesian Twitter with Effective Domain-Specific Vocabulary Initialization,” *EMNLP 2021 - 2021 Conference on Empirical Methods in Natural Language Processing, Proceedings*, pp. 10660–10668, 2021, doi: 10.18653/v1/2021.emnlp-main.833.
- [14] X. Song, A. Salcianu, Y. Song, D. Dopson, and D. Zhou, “Fast WordPiece Tokenization,” *EMNLP 2021 - 2021 Conference on Empirical Methods in Natural Language Processing, Proceedings*, no. Section 3, pp. 2089–2103, 2021, doi: 10.18653/v1/2021.emnlp-main.160.
- [15] O. Galal, A. H. Abdel-Gawad, and M. Farouk, “Rethinking of BERT sentence embedding for text classification,” *Neural Comput. Appl.*, vol. 36, no. 32, pp. 20245–20258, 2024, doi: 10.1007/s00521-024-10212-3.
- [16] A. Vabalas, E. Gowen, E. Poliakoff, and A. J. Casson, “Machine learning algorithm validation with a limited sample size,” *PLoS One*, vol. 14, no. 11, pp. 1–20, 2019, doi: 10.1371/journal.pone.0224365.
- [17] D. Kreuzberger, N. Kuhl, and S. Hirschl, “Machine Learning Operations (MLOps): Overview, Definition, and Architecture,” *IEEE Access*, vol. 11, no. February, pp. 31866–31879, 2023, doi: 10.1109/ACCESS.2023.3262138.