

International Doctoral Colloquium 2026

A Credibility Model Based on Influencer Sentiment Analysis and User Reviews on Digital Investment Platforms

1st Irma Dwi Kusuma,^{1*} 2nd Dwi Suhartini², 3rd Wiwik Handayani³, 4th Catur Suratnoajii⁴
Doctoral Program in Management Science, Faculty of Economics and Business Universitas Pembangunan Nasional
Veteran Jawa Timur, Surabaya, Indonesia

*Corresponding author email address: irmadwikusuma@gmail.com

Abstract. This study develops a digital credibility model in the context of Indonesian retail investment by integrating two sources that shape public perception, namely social media influencers and collective reviewers on digital platforms. Unlike conventional approaches that position influencers as the dominant source of credibility and reviewers as mere electronic word of mouth, this study reconstructs both as equal and complementary credibility entities. The study employs a quantitative approach based on sentiment analysis of 57,808 Google Play Store review data points and TikTok comments from three major investment platforms in Indonesia — Stockbit, Bibit, and Bareksa — covering the period 2025 to 2026. Data were analyzed using a fine-tuned IndoBERT model, BiLSTM with an attention mechanism, and four classical machine learning models with class imbalance handling techniques based on SMOTE and class weighting. Indonesia's digital investment ecosystem during this period was also characterized by dual validation from the Indonesia Stock Exchange (IDX) through JATS system connectivity and KSEI recording, as well as recognition by TradingView as a global analysis platform that listed all three applications among its Preferred Brokers/Platforms in Southeast Asia. On the marketing side, all three platforms leverage a diverse spectrum of influencers, ranging from macro-influencer entertainment figures such as Raditya Dika, Vidi Aldiano, and Maudy Ayunda, to niche TikTok influencers targeting specific communities. Evaluation was conducted using F1-macro, precision, recall, and ROC-AUC metrics. Results indicate that influencer sentiment tends to build emotional and aspirational perceptions, while collective reviewers form rational evaluations based on real experiences — including responses to specific features such as Bareksa Emas and Robo Advisor. The integration of both sources produces a dynamic Dual-Source Digital Credibility Model (DSDCM). These findings extend the Source Credibility Theory framework to a platform-based digital ecosystem and offer practical implications for investment platform management in designing marketing communication strategies aligned with actual user experiences.

Keywords: digital credibility, sentiment analysis, influencer, collective reviewer, IndoBERT, retail investment, Source Credibility Theory, dual-source model

I. Introduction

Indonesia's retail investment ecosystem has undergone significant transformation throughout the 2025–2026 period, marked by the exponential growth of digital investment platforms and increasing social media engagement in shaping the investment decisions of the general public (OJK, 2025). Three leading investment platforms, namely Stockbit, Bibit, and Bareksa, have successfully reached millions of retail users through digital marketing strategies that intensively involve TikTok influencers, while simultaneously facing direct user feedback through Google Play Store reviews.

The credibility of these three platforms does not rest solely on the size of their user communities; it is also supported by two pillars of external legitimacy. First, official recognition from the Indonesia Stock Exchange (IDX): all three applications process transactions directly connected to the JATS (Jakarta Automated Trading System), so that every traded security is legally recorded at KSEI (Kustodian Sentral Efek Indonesia) (BEI, 2025). Second, global standard validation through TradingView, the world's leading technical analysis platform, which lists all three applications among its Preferred Brokers/Platforms in Southeast Asia owing to the stability of the real-time data feeds they provide to Indonesian retail investors (TradingView, 2025). This combination of local-IDX and global-TradingView validation forms an institutional trust foundation that reinforces the effectiveness of influencer-based marketing strategies.

This phenomenon raises a fundamental question in the field of digital marketing management: how do two informationally distinct sources, namely *influencers* who are aspirational in nature and

reviewers who are evaluative, jointly shape the credibility of an investment platform in the public's eyes? This question has not been comprehensively addressed in the existing literature, as the majority of prior studies tend to examine influencers and user reviews separately (Hovland et al., 1953; Ohanian, 1990; Mudambi & Schuff, 2010).

The influencer strategies employed by all three platforms are also multi-layered. At the macro-influencer level, entertainment figures such as Raditya Dika, Vidi Aldiano, and Maudy Ayunda serve as psychological bridges to dismantle the stigma that investing is complicated, by conveying messages of simplicity through lifestyle content on Instagram, TikTok, and YouTube (OJK, 2025). At the niche influencer level, the platforms select figures whose values are congruent with their target segments, such as Islamic educators for Muslim women's communities (Stockbit), Gen Z couples for young investors (Bibit), and religious motivators for worship-oriented aspirations (Bareksa). This dual strategy produces a complex sentiment landscape that warrants empirical investigation.

Source Credibility Theory, developed by Hovland, Janis, and Kelley (1953), traditionally defines communication credibility along two primary dimensions: *expertise* and *trustworthiness*. In the context of modern social media, Ohanian (1990) added the dimension of *attractiveness*. This theoretical framework was constructed within a one-way communication context and does not yet accommodate the dynamics of a digital ecosystem in which two credibility sources operate simultaneously with different orientations.

This study addresses that gap through three main contributions. Conceptually, the study proposes a Dual-Source Digital Credibility Model (DSDCM) that positions influencers and collective reviewers as equal credibility entities within an ecosystem also shaped by IDX institutional legitimacy and TradingView global standards. Methodologically, the study applies deep learning-based sentiment analysis (IndoBERT and BiLSTM) to a large-volume dataset in the context of Indonesian retail investment, including keyword-level analysis to identify specific topics such as product features and spiritual narratives. Practically, the study identifies the gap between influencer promotional narratives and actual user experiences, including responses to specific asset features such as Bareksa Emas, as an indicator of platform credibility.

II. Literature Review

Source Credibility Theory and Its Evolution

Hovland et al. (1953) stated that the effectiveness of a communication message is determined by the receiver's perception of the source's credibility. Credibility is conceptualized through two dimensions: *expertise* as the source's capacity to provide valid information, and *trustworthiness* as the source's motivation to provide honest information. Ohanian (1990) added the dimension of *attractiveness* comprising the components of familiar, likable, classy, sexy, and elegant for the context of celebrity endorsers. This elaboration is relevant in understanding why non-expert influencers in the field of investment, such as entertainment celebrities Raditya Dika, Vidi Aldiano, and Maudy Ayunda, are still capable of stimulating public interest in investing through their social capital and personal appeal (Ohanian, 1990; Kamins, 1990).

In the digital context, Fogg (2003) introduced the concept of *web credibility*, underscoring that credibility in a digital ecosystem is constructive in nature, built from the interaction between content, design, and context. This concept is relevant given that users' assessments of investment platforms do not depend solely on official platform statements, but also on the collective experiences of users documented in digital reviews. On the other hand, the integration of all three platforms into the IDX ecosystem through JATS and KSEI provides a layer of *institutional credibility* that reinforces the overall digital credibility construction (BEI, 2025).

The Influencer Spectrum: From Macro-Influencers to Niche Influencers

Influencer marketing has evolved into a dominant marketing strategy in the financial and investment industry (Lou & Yuan, 2019). In the context of Indonesian retail investment during the

2025–2026 period, the influencer spectrum utilized by investment platforms can be categorized into two tiers that serve distinct yet complementary functions.

1. At the macro-influencer tier, public entertainment figures such as Raditya Dika, Vidi Aldiano, and Maudy Ayunda function as stigma-change agents (*destigmatization agents*). The core message communicated focuses on the ease of starting to invest and the security guaranteed by IDX and OJK regulations. The content produced is generally of an aspirational lifestyle type, in which investment is subtly integrated as part of achieving financial goals without abandoning one's primary career. A potential bias at this level is the tendency to emphasize the "guaranteed profit" aspect, causing the inherently volatile nature of market risk to be overlooked.
2. At the niche influencer tier, platforms select figures with a high degree of value congruence with specific audience segments (Kamins, 1990). Unlike conventional celebrities, social media influencers build credibility through the perception of a *parasocial relationship*, a one-directional bond that feels personal between the influencer and their followers (Horton & Wohl, 1956).

Lou and Yuan (2019) found that influencer credibility in a purchase context consists of four dimensions: *trustworthiness*, *expertise*, *similarity*, and *attractiveness*. The dimension of *similarity*, referring to the perceived likeness between the influencer and the audience, proved to be the most powerful factor in driving behavioral intention. This finding explains the selection of niche influencers based on the following value congruence criteria:

- a. Female Islamic educators for Muslim women's communities (Stockbit)
- b. Gen Z couples for young investors (Bibit)
- c. Religious motivators for worship-oriented aspirations (Bareksa)

The effectiveness of influencers in building long-term credibility depends on the degree of *congruence* between the influencer's persona and the values of the product being promoted (Kamins, 1990). A mismatch between the influencer's image and the actual product reality, which can be identified through user reviews, has the potential to create a *credibility backfire* effect that is detrimental to the platform.

Electronic Word of Mouth and Collective Reviewers

Hennig-Thurau, Gwinner, Walsh, and Gremler (2004) defined Electronic Word of Mouth (eWOM) as positive or negative statements made by potential, actual, or former consumers about a product or company, made available to a multitude of people and institutions via the internet. Reviews on digital application platforms represent one of the most structured and verified forms of eWOM, as they are directly linked to actual usage experiences.

Mudambi and Schuff (2010) found that reviews with high diagnostic content, that is, those providing specific information about product attributes, have a greater impact on consumer decisions compared to reviews that are general in nature. In the context of investment platforms, reviews that discuss transaction execution speed, ease of fund withdrawal, the quality of specific features such as Bareksa Emas and Robo Advisor, or the quality of customer service function as *diagnostic reviews* that contribute to shaping the expectations of new users. Review distribution patterns, particularly the *U-shaped distribution* phenomenon characterized by a dominance of one-star and five-star reviews, reflect polarized satisfaction dynamics that are often associated with specific technical experiences rather than holistic evaluations of the platform (Shim & Lee, 2021). This pattern was identified in the Stockbit and Bareksa review data in the present study.

Sentiment Analysis in the Context of Digital Finance

The application of natural language processing (NLP)-based sentiment analysis in the financial domain has expanded rapidly alongside the availability of transformer-based language models (Devlin et al., 2019). For the Indonesian language context, the IndoBERT model developed based on the BERT architecture and trained on a large-scale Indonesian language corpus has proven to provide superior performance compared to bag-of-words approaches in sentiment classification tasks (Wilie et al., 2020).

Koto et al. (2020) identified specific challenges in Indonesian sentiment analysis, including the use of slang, informal abbreviations, and code-mixing between Indonesian and English. These linguistic characteristics are particularly prominent in the application reviews and TikTok comments that serve as data sources in this study. Keyword frequency analysis (*wordcloud analysis*) provides a qualitative layer of understanding regarding the most relevant topics for each sentiment segment, complementing the quantitative classification produced by transformer-based models (Koto et al., 2020).

III. Research Method

Research Design and Approach

This study employs a quantitative approach with a cross-sectional comparative study design. The unit of analysis is the individual sentiment in each review or comment, while the unit of observation consists of three investment platforms (Stockbit, Bibit, Bareksa) across two data source channels, namely Google Play Store and TikTok.

Data Collection

Data were collected from two primary sources covering the period January 2025 to February 2026. The first source consists of Google Play Store reviews collected using a web scraping technique based on the Google Play Store API. A total of 74,252 reviews were successfully collected from three platforms, with the following breakdown: Stockbit (17,056 reviews), Bibit (48,705 reviews), and Bareksa (8,491 reviews). The second source consists of TikTok comments from three niche influencers selected purposively based on content relevance, value congruence with the platform's audience segment, and interaction volume. A total of 975 comments were collected from the *bundatradersahamsyariah* video for Stockbit (201 comments), the *shafyrahadian* video for Bibit (715 comments), and the *zahidsamosir* video for Bareksa (59 comments). This selection was conducted at the niche influencer tier, while macro-influencer campaigns (such as those involving Raditya Dika, Vidi Aldiano, and Maudy Ayunda) were not included in the scope of this study's text analysis due to limited accessibility of comment data at that level.

Table I presents a summary of the final dataset distribution following deduplication and data cleaning.

TABLE I RESEARCH DATASET DISTRIBUTION

Platform	Play Store	TikTok	Total
Stockbit	17,056	201	17,257
Bibit	48,705	715	49,420
Bareksa	8,491	59	8,550
Total	74,252	975	75,227

After deduplication: 57,808 samples

Data Labeling

Play Store data were labeled based on star rating conversion: ratings of 4 and 5 stars as Positive, 3 stars as Neutral, and 1 and 2 stars as Negative. TikTok data were labeled using the VADER Sentiment Intensity Analyzer with automatic translation into English, subsequently verified manually on a random 10% sample to ensure label validity (inter-rater agreement $\kappa = 0.78$, classified as substantial agreement).

Text Preprocessing

Given the linguistic characteristics of the data, which encompass informal Indonesian, slang, abbreviations, and code-mixing, the preprocessing pipeline was designed in four stages.

1. The first stage is normalization, applying a normalization dictionary containing 120 entries of informal Indonesian slang and abbreviations (e.g., *gak* becomes *tidak*, *bgt* becomes *banget*, *app* becomes *aplikasi*).
2. The second stage is negation handling, employing a token-merging mechanism that combines negation tokens with the following word to prevent loss of semantic context (e.g., *tidak bagus* becomes *tidak_bagus*). Negation words such as *tidak*, *bukan*, and *jangan* are excluded from the stopword removal process.
3. The third stage is standard cleaning, which includes case folding, removal of URLs, mentions, hashtags, and special characters, as well as whitespace normalization.
4. The fourth stage is emoji and punctuation feature extraction. Emojis are converted into separate numerical features, namely the count of positive emojis, negative emojis, total emojis, exclamation marks, question marks, and the ratio of uppercase letters, to be used as additional features in classical models.

Keyword Frequency Analysis (Wordcloud Analysis)

As a complement to quantitative sentiment classification, keyword frequency analysis (*wordcloud analysis*) was conducted on each sentiment class (Positive, Neutral, Negative) for each platform and channel. This analysis provides qualitative understanding of the specific topics driving each sentiment class. Words extracted after normalization and stopword removal were visualized using frequency-weighted wordcloud representations. Dominant keywords such as "*kak*", "*rdpu*", "*uang*", and "*umroh*" in TikTok comments, as well as feature-related words such as "*emas*", "*pencairan*", and "*robo*" in Play Store reviews, were used to enrich the interpretation of quantitative findings.

Model Architecture

1. *IndoBERT Fine-tuned*: The primary model in this study is IndoBERT (Wilie et al., 2020) (*indobenchmark/indobert-base-p1*) fine-tuned for three-class sentiment classification. The additional architecture includes a custom classification head with two fully connected layers (768 to 256 to 3), dropout 0.3, and ReLU activation. This model has a total of 124,638,979 parameters, all of which are trainable. The fine-tuning process uses a differential learning rate (BERT layers: 2×10^{-5} , classification layers: 2×10^{-4}) with a linear warmup scheduler over 4 epochs. Data were split into three partitions: training data (41,621 samples), validation data (4,625 samples), and test data (11,562 samples).
2. *BiLSTM with Attention*: As a deep learning comparison model, a Bidirectional Long Short-Term Memory (BiLSTM) with an attention mechanism was used (Bahdanau et al., 2015). The architecture consists of an embedding layer (dimension 256), two hidden BiLSTM layers (dimension 256 per direction), an attention mechanism for context aggregation, and a classification layer (512 to 128 to 3). The vocabulary size of this model is 10,628 words. The model was trained for 10 epochs using the AdamW optimizer.
3. *Classical Machine Learning Models*: As baselines, four classical machine learning models were trained using TF-IDF representations with an n-gram range of (1,3) and 50,000 feature dimensions, supplemented by 6 emoji features: Logistic Regression, Linear SVC, SGD Classifier, and Random Forest. Emoji and punctuation features were combined with the TF-IDF matrix using sparse matrix concatenation operations.

Class Imbalance Handling

Class distribution imbalance in the training data (46,246 samples) was addressed through two complementary strategies. SMOTE (Synthetic Minority Over-sampling Technique) (Chawla et al., 2002) was applied to the training data to oversample the minority class, resulting in 97,137 balanced training samples with an even distribution of 32,379 samples per class. In addition, proportional class weighting (Negative: 1.46, Neutral: 4.62, Positive: 0.48) was applied to the loss function for all models.

Model Evaluation

All models were evaluated using stratified 5-fold cross-validation on the training data and final evaluation on the test data (11,562 samples, 20% of total data). Evaluation metrics include: F1-macro

(primary metric, treating all classes equally), precision-macro, recall-macro, accuracy, and ROC-AUC per class. The selection of F1-macro as the primary metric is based on the importance of balanced performance across all three sentiment classes, given that the minority class (Neutral) carries significant interpretive relevance in the context of credibility gap analysis.

Credibility Gap Analysis

The Net Sentiment Score (NSS) was calculated for each platform-source combination using the following formula:

$$NSS_{p,s} = \frac{N_{positive} - N_{negative}}{N_{total}} \times 100 \quad (1)$$

Where $N_{positive}$, $N_{negative}$, and N_{total} represent the number of positive, negative, and total sentiments on platform *p* from source *s*, respectively. The Credibility Gap is then formulated as:

$$CG_p = NSS_{p,TikTok} - NSS_{p,Playstore} \quad (2)$$

A positive CG value indicates that sentiment from the influencer channel is more optimistic than actual user experiences, suggesting the presence of promotional bias.

IV. Results and Discussion

Results

After deduplication, the final dataset comprises 57,808 samples. The label distribution prior to deduplication shows: Positive 57,212, Negative 13,460, and Neutral 4,555 out of 75,227 raw data points. After deduplication, the distribution changed to Positive 40,474 (70.0%), Negative 13,163 (22.8%), and Neutral 4,171 (7.2%). In the test data (11,562 samples), the class distribution is: Positive 8,095, Negative 2,633, and Neutral 834. This distribution is consistent with the general pattern of digital platform reviews, in which highly satisfied and highly dissatisfied users are more motivated to leave reviews compared to users with average experiences.

Table II presents the sentiment distribution per platform based on Play Store star ratings and TikTok labels, along with their respective NSS values.

TABLE II SENTIMENT DISTRIBUTION AND NSS PER PLATFORM

Platform	Samples	Positive%	Neutral%	Negative%	NSS
<i>Google Play Store</i>					
Stockbit	13,454	61.9	6.8	31.3	30.6
Bibit	36,726	75.0	6.1	19.0	56.0
Bareksa	6,678	62.3	9.4	28.2	34.1
<i>TikTok Niche Influencer</i>					
Stockbit	199	39.7	45.2	15.1	24.6
Bibit	693	47.6	42.3	10.1	37.5
Bareksa	58	55.2	37.9	6.9	48.3

A notable finding from Table II is the proportion of Neutral sentiment on TikTok, which is considerably higher than on Play Store, ranging from 37.9% to 45.2% on the TikTok channel compared to only 6.1–9.4% on Play Store. This can be explained by the nature of TikTok

comment interactions, which frequently take the form of informative questions, empathetic support statements, or non-evaluative responses, in contrast to Play Store reviews that explicitly solicit ratings.

Based on formula (2), the Credibility Gap calculation results are presented in Table III.

TABLE III NET SENTIMENT SCORE AND CREDIBILITY GAP PER PLATFORM

Platform	NSS TikTok	NSS PS	CG	Interpretation
Stockbit	24.6	30.6	-6.0	Convergent*
Bibit	37.5	56.0	-18.5	Moderate
Bareksa	48.3	34.1	+14.2	Divergent

*Negative CG: Play Store sentiment is higher than TikTok

Table III yields findings that partially contradict hypothesis H2. The Play Store NSS for Stockbit and Bibit is actually higher than TikTok, resulting in negative CG values for both. For Stockbit, the NSS difference is relatively small (-6.0), reflecting the high proportion of neutral-informative comments on TikTok rather than a dominance of negative sentiment. Bibit has the largest absolute CG (-18.5), indicating that although Bibit's Play Store recorded the highest NSS (56.0) among the three platforms, its TikTok sentiment is considerably lower (37.5) due to the large proportion of neutral comments consisting of questions from novice investors. Bareksa is the only platform with a positive CG (+14.2), indicating that the spirituality and religious savings narrative constructed by influencer *zahidsamosir* successfully shaped emotional expectations that surpassed actual user technical evaluations. This finding provides partial support for H2: the hypothesis is confirmed only for Bareksa, while for Stockbit and Bibit, Play Store NSS is actually higher.

Stockbit exhibits a polarized sentiment distribution on Play Store (13,454 samples after deduplication) with Positive sentiment at 61.9%, Neutral at 6.8%, and Negative at 31.3%. The Play Store NSS of 30.6 is the lowest among the three platforms, indicating that the proportion of technical complaints on Stockbit is relatively higher. The *U-shaped distribution* pattern is clearly visible, with a dominance of five-star ratings (approximately 60%) and a surge in one-star reviews (3,601 reviews), indicating specific technical disappointments such as system disruptions following app updates, login issues, or order execution delays during active market hours. Keyword analysis (*wordcloud*) on the negative sentiment cluster reveals a dominance of phrases such as "*tidak bisa login*", "*error terus*", and "*tidak bisa transaksi*", while the positive cluster is dominated by keywords such as "*fitur lengkap*", "*komunitas bagus*", and "*data real-time*". The keywords "*kak*" and "*rdpu*" appear prominently in Neutral sentiment on the TikTok channel, reflecting the nature of comments as questions directed at the influencer. On the TikTok channel, influencer *bundatradersahamsyariah* generated 199 analyzable comments with the distribution: Neutral 45.2%, Positive 39.7%, Negative 15.1%, yielding a TikTok NSS of 24.6 and a CG of -6.0, classified as convergent.

Bibit recorded the best performance among the three platforms on the Play Store channel. Out of 36,726 samples, the sentiment distribution shows Positive 75.0%, Neutral 6.1%, and Negative 19.0%, with the highest NSS of 56.0. The strong dominance of five-star ratings (approximately 69.6% of total reviews) indicates widespread user satisfaction. Thematic analysis identifies four main factors driving positive sentiment: ease of interface, effectiveness of the Robo Advisor feature in recommending products according to risk profile, cost transparency, and diversity of investment products (mutual funds, SBN, FR bonds, stocks). Influencer *shafyrahadian* (Fyra & Dian, Gen Z couple) generated the largest response (693 comments) with the distribution: Positive 47.6%, Neutral 42.3%, Negative 10.1%, and NSS 37.5, yielding a CG of -18.5.

Bareksa has the most varied distribution among the three platforms. On the Play Store channel (6,678 samples), the distribution shows Positive 62.3%, Neutral 9.4%, and Negative 28.2%, with an NSS of 34.1. Keyword analysis identifies three distinct sentiment clusters: the five-star positive cluster dominated by words related to ease of SBN ordering and monthly coupon monitoring; the four-star positive cluster specifically reinforced by the Bareksa Emas feature praised for facilitating asset diversification (mutual funds, SBN, gold) within a single application; and the one-star negative cluster concentrated on complaints about fund withdrawal duration and less intuitive interface. Influencer

zahidsamosir generated 58 analyzable comments with the distribution: Positive 55.2%, Neutral 37.9%, Negative 6.9%, and NSS 48.3, yielding a positive CG of +14.2.

Regarding model performance, Table IV presents a comparison of all models tested on the test data (11,562 samples).

TABLE IV SENTIMENT ANALYSIS MODEL PERFORMANCE COMPARISON

Model	Accuracy	F1-mac	Precision	Recall
Logistic Regression	0.744	0.601	0.588	0.651
SGD Classifier	0.770	0.614	0.605	0.646
Random Forest	0.789	0.617	0.612	0.631
Linear SVC	0.793	0.618	0.605	0.639
BiLSTM + Attention	0.805	0.649	0.637	0.679
IndoBERT FT	0.827	0.665	0.654	0.687

IndoBERT fine-tuned achieved the best performance with an F1-macro of 0.665 and accuracy of 0.827, surpassing the best TF-IDF baseline model (Linear SVC, F1-macro 0.618) by 4.7 percentage points and the BiLSTM model (F1-macro 0.649) by 1.6 points. The fine-tuning process demonstrated consistent improvement on validation: epoch 1 yielded Val F1 0.6595, epoch 3 rose to 0.6759, and the best value was achieved at epoch 4 (Val F1 0.6759). BiLSTM achieved its best validation F1 of 0.6615 at the end of 10 training epochs.

TABLE V INDOBERT PER-CLASS PERFORMANCE ON TEST DATA

Class	Prec.	Rec.	F1	Support
Negative	0.753	0.782	0.767	2,633
Neutral	0.260	0.393	0.313	834
Positive	0.948	0.886	0.916	8,095
Macro avg	0.654	0.687	0.665	11,562
Weighted avg	0.854	0.827	0.839	11,562

The Neutral class remains the greatest challenge with an F1 of 0.313, caused by the small number of samples (834 out of 11,562 test data points) and high linguistic ambiguity. The majority of the Neutral class on the TikTok channel consists of communal expressions such as the greeting "*kak*", religious affirmations "*aamiin*", and contextual questions that are semantically neutral but difficult for the model to distinguish from positive expressions. The Positive class achieves the best performance (F1 0.916) due to the dominance of samples, while the Negative class demonstrates a good balance between precision and recall (F1 0.767).

V. Conclusion

This study successfully developed the Dual-Source Digital Credibility Model (DSDCM), which reconstructs the relationship between influencers and collective reviewers as two equal credibility entities operating through different mechanisms, with institutional legitimacy from IDX (JATS/KSEI connection) and global validation from TradingView as supporting prerequisites. This model constitutes a theoretical novelty by extending the Source Credibility Theory (Hovland et al., 1953) beyond its original one-way communication context into a platform-based digital ecosystem where two credibility sources operate simultaneously with distinct orientations. Influencers build credibility through an emotional pathway grounded in values, aspirations, and personal trust, both through celebrity macro-influencers who dismantle the stigma of investing and niche influencers who build connections with specific communities. Collective reviewers, meanwhile, build credibility through a rational pathway grounded in actual technical experiences and responses to specific product features such as Bareksa Emas, Robo Advisor, and SBN quality.

Among the three platforms studied, Bibit demonstrates the best Play Store performance with an NSS of 56.0 and a positive sentiment distribution of 75.0%, indicating the most widespread user satisfaction and consistency between promotional image and actual experience. Stockbit has the most convergent CG

(−6.0) between the two sources, while Bareksa shows the greatest divergence (+14.2) as a result of an influencer strategy based on spiritual-umroh narratives that successfully built emotional expectations surpassing actual technical experiences, which also encompasses the advantage of the Bareksa Emas feature as a driver of positive four- and five-star reviews. These findings provide partial support for H2: influencer NSS is higher only for Bareksa, while for Stockbit and Bibit, Play Store reviewers actually provide more positive evaluations. Methodologically, IndoBERT fine-tuned proved superior with an F1-macro of 0.665 and accuracy of 0.827. The Neutral class remains the primary challenge (F1 0.313), partly attributable to the linguistic uniqueness of TikTok comments in the form of communal greetings ("*kak*"), abbreviations ("*rdpu*"), and religious affirmations ("*aamiin*", "*umroh*") that characterize the Neutral sentiment in Indonesian investment data.

The theoretical implication of this study lies in the proposition that digital credibility in a platform-based ecosystem is no longer a singular construct, but a dynamic dual-source construct shaped by the simultaneous interaction of emotional influencer signals and rational reviewer evaluations, both reinforced by an institutional legitimacy foundation. The practical implication is that investment platform management needs to monitor the Credibility Gap between influencer and reviewer channels as a key performance indicator, ensuring that promotional narratives remain aligned with actual user experiences to sustain long-term public trust. Future research is recommended to test the DSDCM causally using structural equation modeling, expand the scope to other investment platforms such as Ajaib, Pluang, and Tokopedia Investasi, integrate comment data from macro-influencer campaigns for quantitative comparison across influencer tiers, analyze the temporal dynamics of the credibility gap to identify inflection points in public trust, and integrate visual content analysis from influencer videos using multimodal sentiment analysis techniques.

References

- Bei, Y. (2010). Source credibility theory. In S. H. Priest (Ed.), *Encyclopedia of science and technology communication*. SAGE.
- Bahdanau, D., Cho, K., & Bengio, Y. (2015). Neural machine translation by jointly learning to align and translate. *Proceedings of ICLR 2015*, San Diego, CA, USA.
- Bickart, B., & Schindler, R. M. (2001). Internet forums as influential sources of consumer information. *Journal of Interactive Marketing*, 15(3), 31–40.
- Bursa Efek Indonesia. (2025). *Panduan anggota bursa dan sistem perdagangan JATS*. BEI.
- Chawla, N. V., Bowyer, K. W., Hall, L. O., & Kegelmeyer, W. P. (2002). SMOTE: Synthetic minority over-sampling technique. *Journal of Artificial Intelligence Research*, 16, 321–357.
- Devlin, J., Chang, M., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of NAACL-HLT 2019*, Minneapolis, MN, USA, 4171–4186.
- Fogg, B. J. (2003). *Persuasive technology: Using computers to change what we think and do*. Morgan Kaufmann.
- Hajli, M. N. (2014). A study of the impact of social media on consumers. *International Journal of Market Research*, 56(3), 387–404.
- Hennig-Thurau, T., Gwinner, K. P., Walsh, G., & Gremler, D. D. (2004). Electronic word-of-mouth via consumer-opinion platforms: What motivates consumers to articulate themselves on the internet? *Journal of Interactive Marketing*, 18(1), 38–52.
- Horton, D., & Wohl, R. R. (1956). Mass communication and para-social interaction: Observations on intimacy at a distance. *Psychiatry*, 19(3), 215–229.
- Hovland, C. I., Janis, I. L., & Kelley, H. H. (1953). *Communication and persuasion: Psychological studies of opinion change*. Yale University Press.
- Kamins, M. A. (1990). An investigation into the 'match-up' hypothesis in celebrity advertising: When beauty may be only skin deep. *Journal of Advertising*, 19(1), 4–13.
- Koto, F., Rahimi, A., Lau, J. H., & Baldwin, T. (2020). IndoLEM and IndoBERT: A benchmark dataset and pre-trained language model for Indonesian NLP. *Proceedings of the 28th International Conference on Computational Linguistics (COLING)*, 757–770.
- Lou, C., & Yuan, S. (2019). Influencer marketing: How message value and credibility affect consumer trust of branded content on social media. *Journal of Interactive Advertising*, 19(1), 58–73.
- Mudambi, S. M., & Schuff, D. (2010). What makes a helpful online review? A study of customer reviews on Amazon.com. *MIS Quarterly*, 34(1), 185–200.
- Ohanian, R. (1990). Construction and validation of a scale to measure celebrity endorsers' perceived expertise, trustworthiness, and attractiveness. *Journal of Advertising*, 19(3), 39–52.
- Otoritas Jasa Keuangan. (2025). *Laporan perkembangan keuangan digital Indonesia 2025*. OJK.

- Senecal, S., & Nantel, J. (2004). The influence of online product recommendations on consumers' online choices. *Journal of Retailing*, 80(2), 159–169.
- Shim, K. J., & Lee, J. K. (2021). The effects of user reviews on app downloads: Moderating role of app type. *Journal of Business Research*, 130, 475–483.
- TradingView. (2025). *Preferred broker program: Asia Pacific and Southeast Asia*. TradingView Inc. <https://www.tradingview.com/broker/>
- Wilie, B., Vincentio, K., Winata, G. I., Cahyawijaya, S., Li, X., Lim, Z. Y., Soleman, S., Mahendra, R., Fung, P., Bahar, S., & Purwarianti, A. F. (2020). IndoNLU: Benchmark and resources for evaluating Indonesian natural language understanding. *Proceedings of the 1st Conference of the Asia-Pacific Chapter of ACL*, 843–857.
- Zhang, W., Yoshida, T., Tang, X., & Wang, Q. (2020). Text classification based on multi-word with support vector machine. *Knowledge-Based Systems*, 21(8), 879–886.