

Implementasi Model *Hybrid IndoBERT - LinearSVC* untuk Deteksi Spam Judol *Obfuscated* pada Komentar YouTube

DOI: <http://dx.doi.org/10.35889/progresif.v22i3.3914>

Creative Commons License 4.0 (CC BY –NC)



Danang Budiman Hidayat^{1*}, Ahmad Abdul Chamid², Ahmad Jazuli³

Teknik Informatika, Universitas Muria Kudus, Kudus, Indonesia

*e-mail Corresponding Author: 202251235@std.umk.ac.id

Abstract

Online gambling (judol) spam comments on Indonesian YouTube are disguised using obfuscated text techniques, substituting non-standard Unicode characters designed to bypass string-matching moderation systems. Vocabulary-based Natural Language Processing (NLP) models such as IndoBERT risk representation degradation as tokenizers map non-standard characters to unknown ([UNK]) tokens. This study implemented a hybrid model using a Dual-Track Preprocessing architecture combining IndoBERT [CLS] vectors from normalized text with TF-IDF character N-Gram LinearSVC features from raw Unicode text via sparse matrix concatenation. Experiments on 6,000 YouTube judol comments show the hybrid model achieving Accuracy=0.9978, Precision=0.9956, Recall=1.0000, F1-Score=0.9978 on the test set, identical to standalone IndoBERT and outperforming standalone LinearSVC (F1=0.9944). The hybrid model preserved IndoBERT peak performance while fully closing the 4-comment obfuscated-spam gap missed by LinearSVC, with no performance penalty from fusion.

Keywords: Character N-Gram; Gambling Spam Detection; Hybrid Model; IndoBERT; LinearSVC

Abstrak

Komentar spam judi online (judol) di YouTube Indonesia disamarkan menggunakan teknik *obfuscated text*, yakni substitusi karakter Unicode non-standar yang dirancang untuk menghindari sistem moderasi berbasis pencocokan string. Model *Natural Language Processing* (NLP) berbasis kosakata seperti IndoBERT berisiko mengalami degradasi representasi karena *tokenizer* memetakan karakter non-standar ke token tidak dikenal ([UNK]). Penelitian ini mengimplementasikan model *hybrid* dengan arsitektur *Dual-Track Preprocessing* yang menggabungkan vektor [CLS] IndoBERT dari teks ternormalisasi dan fitur TF-IDF karakter N-Gram LinearSVC dari teks Unicode mentah melalui konkatenasi *sparse matrix*. Eksperimen pada 6.000 sampel komentar judol YouTube menunjukkan model *hybrid* mencapai Accuracy=0,9978, Precision=0,9956, Recall=1,0000, F1-Score=0,9978 pada *test set*, identik dengan IndoBERT *standalone* dan unggul atas LinearSVC *standalone* (F1=0,9944). Model *hybrid* mempertahankan kinerja puncak IndoBERT sekaligus menutup seluruh celah 4 komentar *obfuscated* yang terlewat oleh LinearSVC, tanpa penalti performa dari proses fusi.

Kata Kunci: Deteksi Spam; IndoBERT; Judi Online; Karakter N-Gram; LinearSVC

1. Pendahuluan

Judi online ilegal (judol) merupakan masalah hukum yang meluas di Indonesia. Berdasarkan laporan Pusat Pelaporan dan Analisis Transaksi Keuangan (PPATK), total perputaran dana judi online ilegal di Indonesia mencapai sekitar Rp 150 triliun pada tahun 2023 dengan estimasi 3,2 juta pemain aktif [1]. Platform media sosial berbasis video seperti YouTube menjadi salah satu kanal promosi yang dimanfaatkan sindikat judol melalui penyisipan komentar pada konten-konten populer.

Tantangan utama dalam mendeteksi spam judul adalah penggunaan *obfuscated text*, yakni teknik penyamaran sistematis yang menggantikan karakter Latin standar dengan karakter Unicode visual serupa dari blok berbeda, menyisipkan simbol atau angka, dan memodifikasi ejaan kata kunci. Teknik ini efektif mengelabui sistem moderasi berbasis pencocokan string [4], [8]. Tantangan serupa juga dijumpai pada deteksi ujaran kebencian berbahasa Indonesia, di mana variasi penulisan seperti pengulangan karakter (contoh: “mantap” ditulis “mantaaapppp”) menyebabkan kegagalan normalisasi pada pendekatan berbasis kosakata tetap [19].

Model *Natural Language Processing* (NLP) berbasis kosakata, termasuk model Transformer seperti IndoBERT [3] yang dibangun di atas arsitektur *Bidirectional Encoder Representations from Transformers* (BERT) [2], secara teknis berisiko mengalami degradasi representasi pada teks yang mengandung karakter Unicode non-standar semacam ini. *Tokenizer* WordPiece memetakan karakter di luar kosakata pelatihan ke token tidak dikenal ([UNK]), sehingga informasi visual dan morfologis yang terkandung dalam karakter tersamar berpotensi hilang. Inilah inti masalah yang menjadi objek kajian penelitian ini: bagaimana mendeteksi spam judul secara akurat ketika sinyal tekstualnya sengaja dirancang untuk merusak proses tokenisasi model berbasis kosakata.

Sejumlah riset terdahulu telah berupaya menyelesaikan masalah deteksi spam judul dan konten sejenis di media sosial menggunakan berbagai model deteksi spam. Perdana et al. [4] mendeteksi promosi judi online pada Twitter Indonesia menggunakan algoritma *text mining* berbasis TF-IDF dan SVM, menemukan pola leksikal yang konsisten pada konten promosi judul meskipun menggunakan variasi ejaan informal. Kamdan et al. [5] mengevaluasi IndoBERT untuk deteksi promosi judi online pada komentar YouTube dan mencapai akurasi 98,26%, membuktikan keunggulan model berbasis Transformer dibandingkan pendekatan konvensional pada teks judul YouTube berbahasa Indonesia. Manullang et al. [8] membandingkan Convolutional Neural Network (CNN), model Transformer, dan metode *machine learning* tradisional untuk klasifikasi spam judul pada komentar YouTube Indonesia, dengan hasil terbaik diperoleh model berbasis Transformer.

Nugraha et al. [7] membandingkan IndoBERT dan mBERT secara langsung untuk deteksi komentar judul pada media sosial Indonesia (Twitter dan YouTube), menemukan IndoBERT mencapai akurasi dan F1-Score 98%—unggul atas mBERT (96%)—dan menegaskan keunggulan *pre-training* monolingual untuk domain spesifik berbahasa Indonesia. Santika et al. [6] mengaplikasikan IndoBERT dengan strategi *semi-supervised learning* untuk deteksi promosi judul pada komentar YouTube, meningkatkan F1-Score dari 0,961 (IndoBERT *supervised*) menjadi 0,985 (*semi-supervised*), namun juga mencatat IndoBERT *standalone* sedikit dilampaui SVM pada kondisi data berlabel terbatas—menunjukkan keunggulan relatif antar-representasi bersifat kondisional, bukan universal. Amin et al. [10] menunjukkan efektivitas model BERT untuk deteksi spam berbahasa Indonesia secara lebih umum, dengan F1-Score yang unggul dibandingkan metode berbasis fitur konvensional. Alvin dan Winarsih [17] membandingkan IndoBERT, IndoBERTweet, dan algoritma klasik (SVM, Logistic Regression, Random Forest) pada analisis sentimen media sosial informal, menemukan IndoBERT unggul (F1=0,93) namun algoritma klasik tetap kompetitif (F1≈0,90) setelah penyeimbangan data.

Kelompok riset yang membina penelitian ini juga telah menerapkan IndoBERT dan model berbasis BERT pada berbagai tugas NLP berbahasa Indonesia lainnya. Chamid et al. [26] menerapkan IndoBERT untuk analisis opini publik pada 4.147 komentar YouTube terkait debat calon gubernur, dikombinasikan dengan Latent Dirichlet Allocation (LDA) untuk pemodelan topik—menunjukkan efektivitas IndoBERT pada data komentar YouTube berbahasa Indonesia berskala besar, domain platform yang sama dengan penelitian ini. Pada ranah *aspect-based sentiment analysis* (ABSA), Chamid et al. [27] mengusulkan pendekatan *graph-based semi-supervised deep learning* untuk mengatasi keterbatasan data berlabel, Jazuli et al. [28] mengembangkan pengindeksan semantik berbasis Transformer dengan algoritma *index generation* yang disempurnakan menggunakan BERT, dan Jazuli et al. [29] mengoptimalkan ABSA menggunakan BERT untuk analisis umpan balik mahasiswa Indonesia, mencapai F1-Score klasifikasi sentimen 0,974. Terkait kualitas anotasi data berlabel, Chamid et al. [30] mengembangkan uji konsistensi pelabelan data *multi-label* menggunakan metode Cohen Kappa, merumuskan proses pelabelan data untuk pemodelan *semi-supervised learning* [31], dan mengusulkan klasifikasi teks *multi-label* pada ulasan pengguna Indonesia menggunakan *graph neural network semi-supervised* [32].

Pendekatan berbasis fitur karakter seperti TF-IDF dengan karakter N-Gram, di sisi lain, terbukti tangguh terhadap variasi ortografis karena tidak bergantung pada tokenisasi berbasis kosakata [12]. Arifin et al. [12] menerapkan algoritma SVM dengan pembobotan TF-IDF N-Gram untuk klasifikasi teks, menunjukkan representasi ini efektif menangkap pola leksikal dalam teks Indonesia. Yazid et al. [19] menerapkan Naive Bayes dengan fitur N-Gram untuk deteksi ujaran kebencian *multi-label* berbahasa Indonesia, secara eksplisit mendokumentasikan bagaimana representasi berbasis karakter/N-Gram menangani variasi penulisan yang gagal dinormalisasi oleh pendekatan berbasis kosakata tetap—tantangan yang paralel dengan *obfuscated text* pada penelitian ini. Penerapan SVM dengan TF-IDF untuk analisis sentimen ulasan aplikasi di Google Play Store telah banyak dieksplorasi dalam konteks Indonesia, termasuk pada aplikasi layanan publik perpajakan [16], platform perpesanan [23], layanan perbankan digital [24], dan *chatbot* berbasis AI [22], dengan perbandingan fungsi *kernel* yang konsisten menunjukkan SVM sebagai pilihan kompetitif untuk klasifikasi teks berdimensi tinggi [22], [25]. Ghourabi dan Alohal [11] menggabungkan *embedding* berbasis Transformer dengan *ensemble learning* untuk klasifikasi dan deteksi pesan spam, menghasilkan peningkatan performa yang signifikan dibandingkan pendekatan individual.

Beberapa penelitian menggabungkan representasi berbasis Transformer dengan arsitektur atau fitur tambahan untuk menangani teks informal berbahasa Indonesia. Handayani dan Muljono [20] menggabungkan IndoBERTweet dengan Convolutional Neural Network (CNN) untuk klasifikasi ujaran kebencian berbahasa gaul, mencapai F1-Score 91,1%—mengungguli model berbasis *embedding* statis FastText. Ridho dan Yulianti [14] menerapkan IndoBERT dikombinasikan dengan model *machine learning* klasik (SVM, Naive Bayes, Random Forest) untuk deteksi berita hoaks berbahasa Indonesia, menemukan kombinasi tersebut menghasilkan performa yang lebih stabil dibandingkan IndoBERT *standalone*. Di luar konteks Indonesia, Khoruzhaya et al. [21] membandingkan *ensemble* model *machine learning* klasik dengan model BERT untuk klasifikasi biner laporan medis, menunjukkan kombinasi representasi vektorisasi teks yang berbeda dapat bersaing dengan model BERT murni pada domain berdimensi fitur tinggi. Al Simabua dan Alfath [9] mengoptimalkan *fine-tuning* IndoBERT-Lite untuk deteksi spam pada layanan pelanggan digital, menekankan pentingnya *deduplication* data dan *early stopping* untuk mencegah *overfitting* pada dataset spam yang mengandung template berulang. Pada tahap *preprocessing*, Ardinata et al. [18] mengembangkan metode identifikasi dan normalisasi teks *slang* menggunakan *pre-trained* FastText pada data Twitter berbahasa Indonesia, relevan dengan tantangan normalisasi teks yang dihadapi penelitian ini.

Tinjauan riset-riset terdahulu tersebut mengidentifikasi gap yang masih tersisa: penelitian-penelitian deteksi spam judul dan model deteksi spam sejenis yang ada [4]–[9], [16]–[32] secara konsisten mengevaluasi model berbasis kosakata (IndoBERT/BERT) atau model berbasis fitur karakter (TF-IDF/SVM) secara terpisah, atau menggabungkan keduanya pada input teks yang sama tanpa membedakan strategi *preprocessing* per-jalur. Belum ada penelitian yang secara eksplisit merancang *preprocessing* berbeda untuk setiap jalur representasi—teks ternormalisasi untuk model semantik dan teks Unicode mentah untuk model karakter—sekaligus menguji secara terukur apakah strategi fusi tersebut mempertahankan kinerja model semantik individual sambil memperbaiki kelemahan model karakter individual, khususnya pada subset spam yang secara aktif menggunakan *obfuscated text*.

Penelitian ini mengusulkan model *hybrid* dengan arsitektur *Dual-Track Preprocessing* untuk mengisi gap tersebut, memanfaatkan kedua pendekatan secara bersamaan: jalur IndoBERT (Track A) mengolah teks ternormalisasi untuk representasi semantik kontekstual, sementara jalur LinearSVC berbasis karakter N-Gram (Track B) mempertahankan teks Unicode mentah sebagai sinyal diskriminatif. Kedua representasi digabungkan melalui konkatenasi *sparse matrix* sebelum diklasifikasikan oleh LinearSVC akhir. Rasionalisasi konsep ini didukung oleh dua lini bukti dari riset terdahulu: pertama, representasi kontekstual BERT/IndoBERT [2], [3] terbukti unggul menangkap makna semantik pada berbagai tugas NLP Indonesia [17], [26]–[29]; kedua, representasi karakter N-Gram terbukti tangguh terhadap variasi ortografis karena tidak bergantung pada tokenisasi berbasis kosakata [12], [19]. Model *hybrid*, dengan menggabungkan kedua sinyal yang saling melengkapi tersebut melalui satu titik fusi tunggal, diharapkan dapat mewarisi kekuatan representasi semantik tanpa kehilangan sensitivitas terhadap manipulasi karakter.

Kebaruan (*state of the art*) penelitian ini terletak pada tiga kontribusi utama: (1) arsitektur *Dual-Track Preprocessing* yang secara eksplisit memisahkan aliran teks ternormalisasi dan teks

Unicode mentah—berbeda dari penelitian-penelitian sebelumnya [5]–[9] yang mengevaluasi IndoBERT atau model berbasis fitur karakter secara individual pada input yang seragam; (2) strategi fusi representasi berbasis konkatenasi *sparse matrix* antara vektor [CLS] IndoBERT dan vektor TF-IDF karakter N-Gram, bukan sekadar *ensemble* model generik; (3) evaluasi komprehensif yang ditargetkan secara khusus pada subset spam *obfuscated*, memberikan bukti kuantitatif tentang *trade-off* (atau ketiadaannya) dalam strategi fusi representasi untuk domain *obfuscated text* berbahasa Indonesia—pengujian yang belum dilakukan secara eksplisit pada literatur deteksi spam judul maupun deteksi spam berbahasa Indonesia secara umum sebelumnya.

2. Metodologi

Implementasi model hybrid IndoBERT-LinearSVC pada penelitian ini terdiri dari lima tahapan utama yang dieksekusi secara berurutan: (1) akuisisi dan pra-pemrosesan dataset, (2) pelatihan Model A (IndoBERT fine-tuned), (3) pelatihan Model B (LinearSVC karakter N-Gram), (4) fusi representasi dan pelatihan Model C (Hybrid), dan (5) evaluasi model. Kelima tahapan ini disajikan kembali pada Bab Hasil dan Pembahasan dengan fokus pada proses eksekusi dan hasil yang diperoleh.

1) Tahap 1: Akuisisi dan Pra-pemrosesan Dataset

Dataset yang digunakan berasal dari dataset publik "Dataset Komentar Judi Online di YouTube" yang dikumpulkan oleh Muhammad Yusuf Tri Daryanto (Kaggle username: kyyyyy8), terdiri dari 45.592 baris data mentah dengan tiga kolom: text, label, dan text_preprocessed. Proses pembersihan data menghapus 1.213 baris dengan nilai kosong pada kolom text_preprocessed (tersisa 44.379 baris), kemudian menghapus 10.343 baris duplikat berdasarkan kolom text (tersisa 34.036 baris bersih—76% dari data mentah).

Distribusi kelas pada data bersih (34.036 baris) menunjukkan 30.719 komentar non-spam (90,3%) dan 3.317 komentar spam (9,7%), dengan rasio ketidakseimbangan 9,26:1. Untuk menjamin keseimbangan kelas pada pelatihan, dipilih subset 6.000 sampel (3.000 spam, 3.000 non-spam) dari data bersih. Karena hanya 89 dari 3.317 spam yang bersifat *non-obfuscated*, seluruh 3.000 sampel spam pada subset sengaja diambil dari pool *obfuscated* (3.228 tersedia) agar evaluasi secara konsisten menguji *robustness* terhadap *obfuscation*—mengakibatkan 100% kelas spam pada subset penelitian bersifat *obfuscated by design*. Subset kemudian dibagi menggunakan *stratified split* rasio 70:15:15 (RANDOM_SEED=42), menghasilkan Train 4.200 (2.100/2.100), Validation 900 (450/450), dan Test 900 (450/450) (Tabel 1). Dataset menyediakan kolom text (Unicode mentah) dan text_preprocessed (ternormalisasi oleh pemilik dataset), yang menjadi input Track B dan Track A secara berurutan pada tahap-tahap berikutnya.

Tabel 1 Distribusi Dataset Penelitian (Subset Seimbang)

Partisi	Total	Spam	Non-Spam	Rasio
Train	4.200	2.100	2.100	1:1
Validation	900	450	450	1:1
Test	900	450	450	1:1
Total	6.000	3.000	3.000	1:1

Tabel 2 Konfigurasi Pelatihan Masing-Masing Model (Ringkasan Tahap 2–4)

Model	Komponen	Konfigurasi Utama	Input
Model A	IndoBERT-p2 fine-tuned	lr=2e-5, 5 <i>epoch</i> penuh, <i>early stopping</i> (patience=2, tidak terpicu)	text_preprocessed
Model B	LinearSVC + TF-IDF	Ablation N-Gram (5 konfigurasi), C <i>grid</i> manual (<i>validation set</i>)	text (raw)
Model C	Hybrid (A+B)	hstack([CLS,tfidf]), C <i>grid</i> manual (<i>validation set</i>)	Keduanya

2) Tahap 2: Pelatihan Model A — IndoBERT Fine-Tuning (Track A)

Model A menggunakan indobenchmark/indobert-base-p2 yang di-*fine-tune* untuk klasifikasi biner. Input teks ternormalisasi (*text_preprocessed*) ditokenisasi dengan IndoBERTTokenizer (*max_length=128*). Prediksi kelas dihasilkan dari vektor [CLS] pada lapisan Transformer terakhir melalui lapisan klasifikasi linear dan fungsi aktivasi *softmax*, dirumuskan pada Persamaan (1):

$$\hat{y} = \text{softmax}(W \cdot h_{\{[CLS]\}} + b) \quad (1)$$

Keterangan: $\hat{y}$ = probabilitas kelas prediksi; $h_{\{[CLS]\}}$ = hidden state token [CLS] berdimensi 768 dari lapisan Transformer terakhir; W , b = bobot dan bias lapisan klasifikasi yang dilatih.

Pelatihan menggunakan *optimizer* AdamW ($\text{lr}=2 \times 10^{-5}$, $\text{weight_decay}=0,01$) dan *linear warmup scheduler*, berjalan penuh selama 5 *epoch*. Mekanisme *early stopping* (*patience=2*) dikonfigurasi untuk mencegah *overfitting* namun dirancang berhenti hanya jika *validation* F1 tidak meningkat dalam dua *epoch* berturut-turut. *Checkpoint* dengan *validation* F1 tertinggi disimpan sebagai model final, dan vektor [CLS] (768 dimensi) dari *checkpoint* tersebut diekstraksi untuk seluruh *split* data sebagai komponen representasi Track A pada Tahap 4.

3) Tahap 3: Pelatihan Model B — LinearSVC Karakter N-Gram (Track B)

Model B menggunakan representasi *Term Frequency-Inverse Document Frequency* (TF-IDF) berbasis karakter N-Gram [12] pada teks Unicode mentah (kolom *text*), dengan *max_features=50.000* dan *min_df=2*. Berbeda dari TF-IDF berbasis kata, unit yang dihitung adalah sub-string karakter sepanjang N , sehingga tangguh terhadap substitusi karakter Unicode karena tidak bergantung pada kosakata kata utuh. Pembobotan TF-IDF dihitung melalui tiga tahap berikut.

$$TF(t,d) = f_{\{t,d\}} / \sum_{\{t'\}} f_{\{t',d\}} \quad (2)$$

Keterangan: $TF(t,d)$ = term frequency substring karakter t pada dokumen d ; $f_{\{t,d\}}$ = jumlah kemunculan t pada d ; $\sum_{\{t'\}} f_{\{t',d\}}$ = total kemunculan seluruh substring karakter pada d .

$$IDF(t) = \log(N / df(t)) \quad (3)$$

Keterangan: N = jumlah total dokumen dalam korpus; $df(t)$ = jumlah dokumen yang mengandung substring karakter t .

$$TF-IDF(t,d) = TF(t,d) \times IDF(t) \quad (4)$$

Rentang N-Gram optimal ditentukan melalui *ablation study* pada set validasi terhadap lima konfigurasi: (1,1), (2,2), (3,3), (1,3), dan (2,4) [default awal], dengan *hyperparameter* $C=1,0$ tetap untuk mengisolasi efek rentang N-Gram (Tabel 3). Hasil *ablation* disajikan pada Bab Hasil dan Pembahasan.

Tabel 3 Konfigurasi *Ablation Study* Karakter N-Gram (Set Validasi, $C=1,0$)

Konfigurasi	N-Gram Range	Jml. Fitur	Keterangan
Abl-1	(1, 1)	663	Unigram karakter
Abl-2	(2, 2)	2.899	Bigram karakter
Abl-3	(3, 3)	8.176	Trigram karakter
Abl-4	(1, 3)	11.738	Unigram–Trigram
Abl-5	(2, 4)	28.771	Bigram–4-Gram (default awal)

Catatan: "Ablation Abl-6" pada proposal penelitian (Bab III Tabel 3.8) merujuk pada evaluasi subset *spam obfuscated* pada Tahap 5, bukan konfigurasi N-Gram tambahan.

Vectorizer di-*refit* setelah konfigurasi N-Gram terbaik ditentukan, dan *hyperparameter* C pada LinearSVC di-*tuning* melalui pencarian *grid* manual pada *validation set*—dilatih sekali per kandidat pada *training set* dan dievaluasi pada *validation set* (bukan *cross-validation*, untuk menghindari *data leakage* antar-tahap)—pada kandidat $C \in \{0,01; 0,1; 1,0; 10,0\}$. LinearSVC mengklasifikasikan vektor fitur melalui fungsi keputusan linear pada Persamaan (5), dengan syarat *margin* pada Persamaan (6)–(7) mengikuti formulasi Vapnik dan Cortes [2]:

$$f(x) = w^T x + b \quad (5)$$

$$y_i(w^T x_i + b) \geq 1, \forall i \quad (6)$$

$$\|w\| \text{ diminimalkan terhadap batasan Persamaan (6)} \quad (7)$$

Keterangan: $f(x)$ = fungsi keputusan; w = vektor bobot hyperplane; b = bias; x_i = vektor fitur TF-IDF sampel ke- i ; $y_i \in \{-1, +1\}$ = label kelas sampel ke- i .

4) Tahap 4: Fusi Representasi dan Pelatihan Model C (Hybrid)

Model C menggabungkan representasi Track A dan Track B melalui konkatenasi sparse matrix, dirumuskan pada Persamaan (8):

$$v_{\text{hybrid}} = [v_{\text{CLS}}] \oplus v_{\text{tfidf}} \quad (8)$$

Keterangan: $v_{\text{CLS}} \in \mathbb{R}^{768}$ = vektor [CLS] dari Model A (dikonversi ke sparse matrix); $v_{\text{tfidf}} \in \mathbb{R}^{11738}$ = vektor TF-IDF karakter N-Gram terbaik dari Model B; $\oplus$ = operator konkatenasi horizontal (hstack); $v_{\text{hybrid}} \in \mathbb{R}^{12506}$ = vektor fitur gabungan.

Vektor gabungan v_{hybrid} digunakan sebagai fitur untuk melatih LinearSVC baru mengikuti fungsi keputusan yang sama pada Persamaan (5), dengan *hyperparameter* C di-*tuning* ulang menggunakan pencarian *grid* manual pada *validation set* dengan kandidat C yang sama seperti Tahap 3.

5) Tahap 5: Evaluasi Model

Evaluasi dilakukan pada *test set* (900 sampel) menggunakan empat metrik standar klasifikasi biner, dengan kelas Spam (label=1) sebagai kelas positif pada seluruh pelaporan metrik, dirumuskan pada Persamaan (9)–(12):

$$\text{Accuracy} = (TP + TN) / (TP + TN + FP + FN) \quad (9)$$

$$\text{Precision} = TP / (TP + FP) \quad (10)$$

$$\text{Recall} = TP / (TP + FN) \quad (11)$$

$$F1\text{-Score} = 2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall}) \quad (12)$$

Keterangan: TP = True Positive (spam terdeteksi benar); TN = True Negative (non-spam terdeteksi benar); FP = False Positive (non-spam salah terdeteksi spam); FN = False Negative (spam lolos terdeteksi).

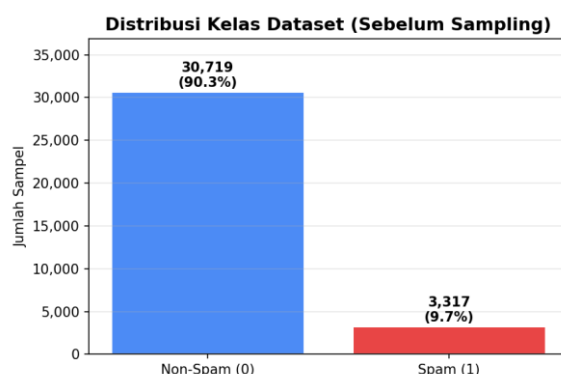
Evaluasi khusus juga dilakukan pada 450 komentar spam dalam *test set*, selain evaluasi global pada seluruh *test set* (evaluasi ini diberi label Ablation Abl-6 sesuai proposal penelitian), untuk mengukur kemampuan deteksi teks yang secara aktif memanipulasi karakter. Karena seluruh 450 spam pada *test set* berasal dari pool *obfuscated* (Tahap 1), evaluasi Abl-6 secara matematis identik dengan breakdown kelas spam pada evaluasi global—konsekuensi dari desain sampling subset, bukan evaluasi pada data yang independen.

3. Hasil Dan Pembahasan

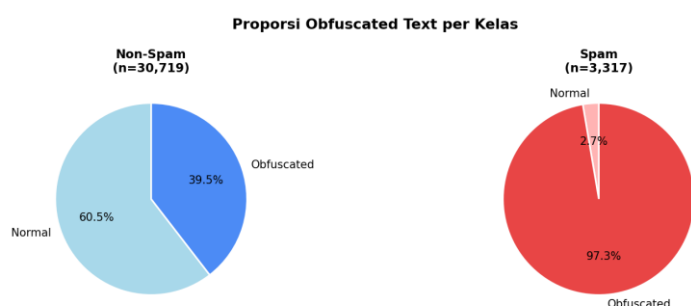
1) Hasil Tahap 1: Dataset Setelah Pra-pemrosesan

Proses akuisisi dan pembersihan pada Tahap 1 menghasilkan 34.036 baris data bersih dari 45.592 baris mentah. Gambar 1 menunjukkan distribusi kelas pada data bersih tersebut: 30.719 non-spam (90,3%) dan 3.317 spam (9,7%), rasio 9,26:1. Analisis *obfuscation* (Gambar 2) menunjukkan 97,3% komentar spam (3.228 dari 3.317) mengandung karakter Unicode non-standar (rata-rata 9,34 karakter tersamar/komentar), dibandingkan 39,5% pada komentar non-spam (rata-rata 1,06 karakter tersamar/komentar). *Obfuscation* bukan indikator biner yang

eksklusif untuk spam, namun densitas kemunculannya per komentar jauh lebih tinggi pada kelas spam—sejalan dengan pola manipulasi karakter sistematis yang didokumentasikan pada penelitian deteksi judul sebelumnya [4], [5].



Gambar 1 Distribusi kelas pada dataset bersih (34.036 sampel, sebelum seleksi subset): 30.719 non-spam (90,3%) dan 3.317 spam (9,7%), rasio 9,26:1.

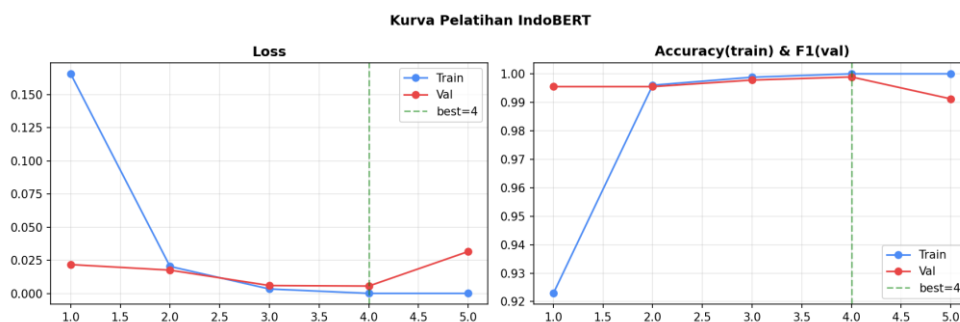


Gambar 2 Proporsi komentar mengandung karakter *obfuscated* per kelas pada dataset bersih: 97,3% spam vs 39,5% non-spam.

Hasil seleksi subset seimbang dan pembagian data mengikuti prosedur Tahap 1 menghasilkan distribusi persis sesuai rancangan (Tabel 1): Train 4.200, Validation 900, dan Test 900, masing-masing dengan rasio 1:1 antara kelas spam dan non-spam.

2) Hasil Tahap 2: Pelatihan dan Evaluasi Model A (IndoBERT)

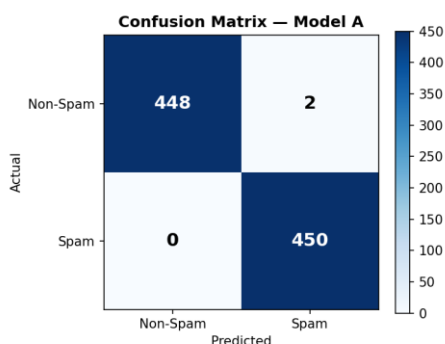
Model A (indobenchmark/indobert-base-p2) dilatih penuh selama 5 *epoch* mengikuti prosedur Tahap 2. *Checkpoint* terbaik dipilih pada *epoch* ke-4 berdasarkan *validation* F1 tertinggi (0,9989); mekanisme *early stopping* (patience=2) tidak pernah benar-benar terpicu karena tidak ada dua *epoch* berturut-turut tanpa peningkatan (Gambar 3). Pada *test set*, *checkpoint* ini menghasilkan Accuracy=0,9978, Precision=0,9956, Recall=1,0000, F1-Score=0,9978 (Tabel 4, Gambar 4)—Recall sempurna dengan FN=0, berarti seluruh 450 komentar spam *obfuscated* dalam *test set* berhasil terdeteksi oleh IndoBERT *standalone* tanpa satu pun yang lolos.



Gambar 3 Kurva *training loss* dan *validation loss* IndoBERT selama 5 *epoch* pelatihan (Tahap 2).

Tabel 4 Evaluasi Model A pada *Test Set*

Model	Accuracy	Precision	Recall	F1-Score
IndoBERT Fine-Tuned (A)	0,9978	0,9956	1,0000	0,9978

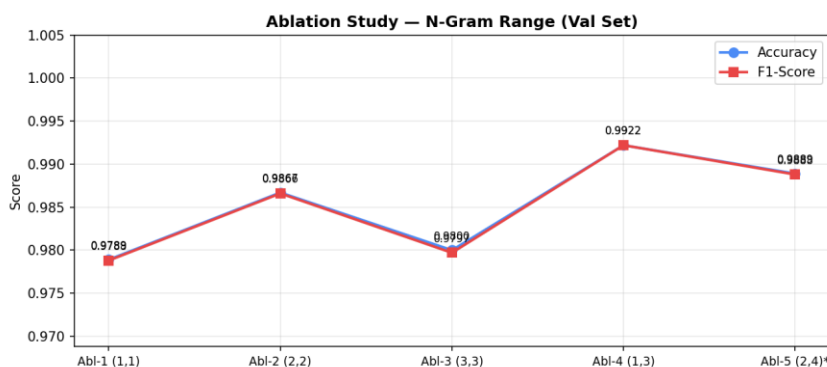
**Gambar 4** Confusion matrix Model A. TP=450, TN=448, FP=2, FN=0.

3) Hasil Tahap 3: Ablation, Tuning, dan Evaluasi Model B (LinearSVC)

Hasil *ablation study* mengikuti prosedur Tahap 3 (Tabel 5) menunjukkan rentang N-Gram (1,3) menghasilkan F1-Score tertinggi (0,9922) pada set validasi, mengungguli konfigurasi default awal (2,4) yang hanya mencapai F1=0,9888 (Gambar 5). Rentang unigram-trigram lebih efektif dibanding rentang yang lebih sempit seperti (3,3) atau lebih lebar seperti (2,4), mengindikasikan kombinasi unit karakter pendek hingga menengah lebih diskriminatif untuk pola substitusi Unicode pada dataset ini [12].

Tabel 5 Hasil *Ablation Study* N-Gram pada Set Validasi

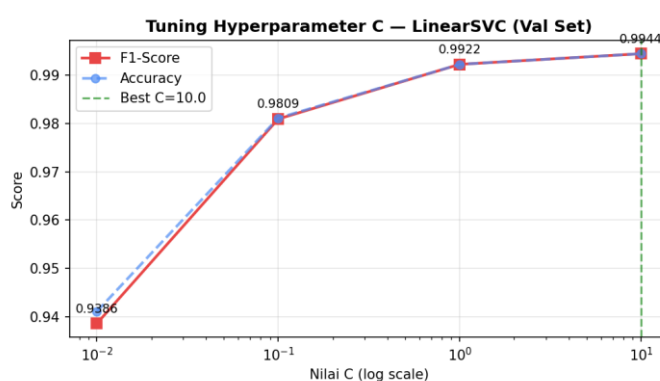
Konfigurasi	N-Gram	Accuracy	Precision	Recall	F1-Score
Abl-1	(1,1)	0,9789	0,9843	0,9733	0,9788
Abl-2	(2,2)	0,9867	0,9910	0,9822	0,9866
Abl-3	(3,3)	0,9800	0,9954	0,9644	0,9797
Abl-4*	(1,3)	0,9922	0,9978	0,9867	0,9922 ✓
Abl-5	(2,4)	0,9889	0,9977	0,9800	0,9888

**Gambar 5** Perbandingan F1-Score *ablation study* konfigurasi N-Gram (set validasi).

Tuning hyperparameter C mengikuti prosedur Tahap 3 dengan kandidat {0,01; 0,1; 1,0; 10,0} menggunakan N-Gram terbaik (1,3) (Tabel 6, Gambar 6). C=10,0 menghasilkan F1-Score tertinggi (0,9944) dan dipilih sebagai nilai final.

Tabel 6 Hasil *Tuning Hyperparameter C* pada *Validation Set* (Model B)

Param. C	Accuracy	Precision	Recall	F1-Score	Keterangan
0,01	0,9411	0,9806	0,9000	0,9386	<i>Underfitting</i>
0,1	0,9811	0,9909	0,9711	0,9809	
1,0	0,9922	0,9978	0,9867	0,9922	Default awal
10,0*	0,9944	0,9978	0,9911	0,9944	Terbaik ✓

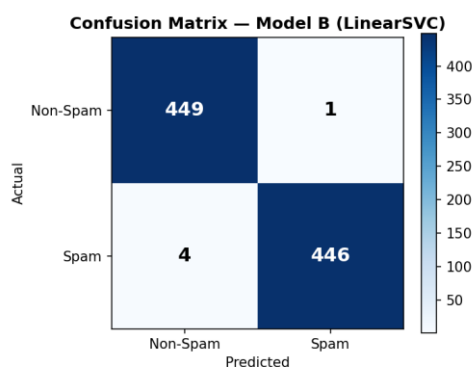


Gambar 6 Kurva F1-Score pada *validation set* terhadap kandidat C untuk LinearSVC (Model B).

Model B dengan konfigurasi final (N-Gram (1,3), C=10,0) menghasilkan F1-Score 0,9944 pada *test set* (Tabel 7, Gambar 7). Terdapat 4 *False Negative* (FN=4), seluruhnya merupakan komentar spam *obfuscated* yang pola karakternya tidak sepenuhnya tertangkap oleh representasi N-Gram (1,3) meskipun C sudah optimal.

Tabel 7 Evaluasi Model B pada *Test Set*

Model	Accuracy	Precision	Recall	F1-Score
LinearSVC Char N-Gram (B)	0,9944	0,9978	0,9911	0,9944



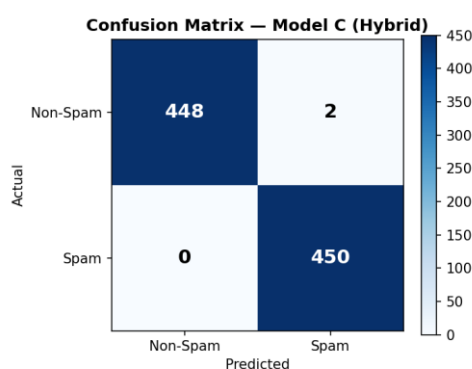
Gambar 7 Confusion matrix Model B. TP=446, TN=449, FP=1, FN=4.

4) Hasil Tahap 4: Fusi dan Evaluasi Model C (Hybrid)

Model C dibangun melalui konkatenasi *sparse matrix* vektor [CLS] IndoBERT (768 dimensi) dengan vektor TF-IDF N-Gram (1,3) Model B (11.738 dimensi) mengikuti prosedur Tahap 4, menghasilkan 12.506 dimensi fitur gabungan. *Tuning* C pada *validation set* (kandidat sama: {0,01; 0,1; 1,0; 10,0}) menghasilkan skor identik ($F1=0,9989$) untuk seluruh kandidat, sehingga $C=0,01$ dipilih mengikuti aturan *tie-break* (C terkecil). Pada *test set*, Model C mencapai $Accuracy=0,9978$, $Precision=0,9956$, $Recall=1,0000$, $F1-Score=0,9978$ (Tabel 8, Gambar 8)—secara numerik identik dengan Model A pada seluruh metrik dan pola error ($FP=2$, $FN=0$ pada keduanya).

Tabel 8 Evaluasi Model C pada *Test Set*

Model	Accuracy	Precision	Recall	F1-Score
Hybrid IndoBERT-LinearSVC (C)	0,9978	0,9956	1,0000	0,9978



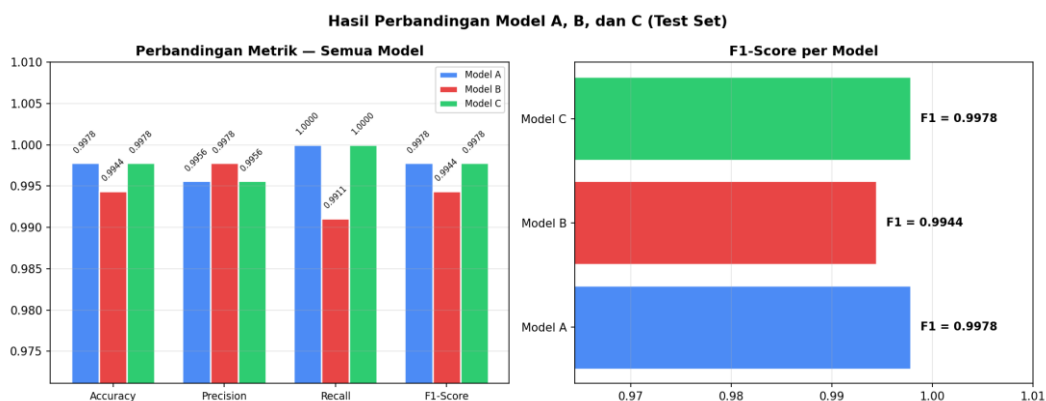
Gambar 8 Confusion matrix Model C. $TP=450$, $TN=448$, $FP=2$, $FN=0$.

5) Hasil Tahap 5: Evaluasi Perbandingan dan Pembahasan

Tabel 9 merangkum perbandingan performa ketiga model pada *test set* ($N=900$) mengikuti prosedur evaluasi pada Tahap 5; Gambar 9 memvisualisasikannya. Karena seluruh 450 sampel kelas spam pada *test set* berasal dari pool *obfuscated text* (Tahap 1), kolom FN pada Tabel 9 secara langsung merepresentasikan hasil evaluasi Ablation Abl-6 (evaluasi subset *obfuscated* sesuai proposal): tidak diperlukan tabel terpisah karena tidak terdapat subset spam *non-obfuscated* dalam *test set* ini untuk dibandingkan.

Tabel 9 Perbandingan Performa Model A, B, dan C pada *Test Set* ($N=900$)

Model	Accuracy	Precision	Recall	F1-Score	FP	FN
Model A – IndoBERT	0,9978	0,9956	1,0000	0,9978	2	0
Model B – LinearSVC	0,9944	0,9978	0,9911	0,9944	1	4
Model C – Hybrid	0,9978	0,9956	1,0000	0,9978	2	0



Gambar 9 Perbandingan Accuracy, Precision, Recall, dan F1-Score ketiga model pada *test set*.

Temuan utama penelitian ini adalah Model C (*Hybrid*) mencocokkan secara eksak kinerja Model A (IndoBERT *standalone*) pada seluruh metrik—Accuracy, Precision, Recall, dan F1-Score keduanya identik (0,9978/0,9956/1,0000/0,9978), termasuk pola error yang sama persis (FP=2, FN=0). Ini berarti proses fusi representasi pada Tahap 4 tidak kehilangan sedikit pun kinerja IndoBERT, sekaligus sepenuhnya menutup kelemahan Model B (LinearSVC) yang melewatkan 4 dari 450 komentar spam *obfuscated* (FN=4).

Temuan ini perlu diinterpretasikan secara proporsional. Penelitian ini tidak membuktikan bahwa arsitektur *hybrid* melampaui IndoBERT *standalone*, maupun bahwa fusi representasi mencapai sesuatu yang tidak dapat dicapai oleh salah satu jalur saja—pada *test set* ini, IndoBERT *standalone* semata sudah cukup untuk mendeteksi seluruh spam *obfuscated* tanpa kesalahan. Kontribusi empiris yang dapat dipertanggungjawabkan dari eksperimen ini adalah bahwa fusi representasi karakter N-Gram tidak memberikan penalti performa terhadap jalur semantik IndoBERT yang sudah sangat kuat, sekaligus terbukti mampu memperbaiki performa jalur karakter N-Gram *standalone* hingga menyamai performa terbaik yang tersedia dalam eksperimen ini—pola yang konsisten dengan temuan Khoruzhaya et al. [21] dan Ridho dan Yulianti [14] bahwa fusi representasi cenderung menstabilkan performa dibanding model tunggal, meskipun tidak selalu meningkatkan puncak performa itu sendiri.

Beberapa keterbatasan penting perlu digarisbawahi. Pertama, proses pembersihan data pada Tahap 1 menghapus 1.213 baris *null* dan 10.343 baris duplikat dari 45.592 baris mentah (24% dari total data mentah), sebuah tahap yang berdampak signifikan pada ukuran dan distribusi dataset akhir. Kedua, seluruh 450 sampel spam pada *test set* sengaja diambil dari pool *obfuscated* (100% *by design*), sehingga penelitian ini tidak memiliki data pembanding *non-obfuscated* spam dalam evaluasinya sendiri—klaim tentang ketahanan relatif terhadap teks *obfuscated* dibanding *non-obfuscated* tidak dapat diuji langsung dari desain eksperimen ini. Ketiga, kolom *text_preprocessed* yang menjadi input Track A dibuat oleh pemilik dataset, bukan oleh peneliti, sehingga kualitas normalisasi berada di luar kendali metodologis penelitian. Keempat, subset penelitian berdistribusi seimbang 50:50, sedangkan distribusi dataset bersih adalah 9,26:1, sehingga metrik yang dilaporkan tidak langsung mencerminkan ekspektasi performa pada kondisi produksi. Kelima, evaluasi terbatas pada satu dataset tunggal dari satu platform.

4. Simpulan

Penelitian ini mengimplementasikan model *hybrid* IndoBERT-LinearSVC dengan arsitektur *Dual-Track Preprocessing* untuk deteksi spam judi online (judol) *obfuscated* pada komentar YouTube berbahasa Indonesia, melalui lima tahapan: akuisisi dan pra-pemrosesan dataset, pelatihan Model A (IndoBERT), pelatihan Model B (LinearSVC karakter N-Gram), fusi representasi menjadi Model C (*Hybrid*), dan evaluasi. Arsitektur ini memisahkan aliran teks ternormalisasi untuk representasi semantik IndoBERT [3] dan teks Unicode mentah untuk representasi karakter N-Gram LinearSVC [12], kemudian menggabungkan keduanya melalui konkatenasi *sparse matrix*.

Eksperimen pada subset 6.000 sampel menunjukkan Model C (*Hybrid*) mencapai F1-Score=0,9978 pada *test set*—secara numerik identik dengan Model A (IndoBERT *standalone*) pada seluruh metrik (Accuracy, Precision, Recall, F1-Score) dan pola error (FP=2, FN=0 pada keduanya), serta unggul atas Model B (LinearSVC *standalone*, F1=0,9944, FN=4). Temuan ini

menunjukkan bahwa fusi representasi karakter N-Gram ke dalam representasi semantik IndoBERT tidak mengorbankan performa dan berhasil menutup celah yang dimiliki LinearSVC *standalone* pada deteksi spam *obfuscated*—namun penelitian ini tidak membuktikan keunggulan *hybrid* di atas IndoBERT *standalone*, karena pada *test set* ini IndoBERT saja telah mencapai kinerja puncak yang identik.

Penelitian lanjutan disarankan untuk mengevaluasi arsitektur ini pada skenario di mana IndoBERT *standalone* berpotensi tidak mencapai kinerja sempurna—misalnya pada data dengan distribusi kelas asli (9,26:1), pada subset spam yang mencampur komentar *obfuscated* dan *non-obfuscated*, atau pada platform dan domain bahasa lain—untuk menguji apakah keunggulan fusi representasi baru terlihat pada kondisi yang lebih menantang dibanding *test set* penelitian ini. Pengembangan *pipeline preprocessing* Unicode mandiri sebagai Track A, alih-alih bergantung pada kolom yang disediakan pemilik dataset, juga disarankan untuk meningkatkan kendali metodologi penelitian selanjutnya.

5. Ucapan Terima Kasih

Penulis mengucapkan terima kasih kepada Universitas Muria Kudus atas dukungan fasilitas penelitian, kepada Muhammad Yusuf Tri Daryanto (kyyyy8) atas ketersediaan dataset komentar judul YouTube secara publik, serta kepada tim pengembang IndoBERT atas tersedianya model pra-latih Bahasa Indonesia.

Daftar Referensi

- [1] Pusat Pelaporan dan Analisis Transaksi Keuangan (PPATK), "Laporan Tahunan PPATK 2023," PPATK, Jakarta, 2023. [Online]. Tersedia: <https://www.ppatk.go.id>
- [2] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: *Pre-training of Deep Bidirectional Transformers for Language Understanding*," in *Proc. 2019 Conf. North American Chapter ACL: Human Language Technologies (NAACL-HLT)*, Minneapolis, MN, 2019, pp. 4171-4186, doi: 10.18653/v1/N19-1423.
- [3] B. Wilie et al., "IndoNLU: Benchmark and Resources for Evaluating Indonesian Natural Language Understanding," in *Proc. 1st Conf. Asia-Pacific Chapter ACL 8th Int. Joint Conf. NLP (AACL-IJCNLP 2020)*, Suzhou, China, Dec. 2020, pp. 843-857. [Online]. Available: <https://aclanthology.org/2020.aacl-main.85>
- [4] R. B. Perdana, Ardin, I. Budi, A. B. Santoso, A. Ramadiah, and P. K. Putra, "Detecting Online Gambling Promotions on Indonesian Twitter Using Text Mining Algorithm," *Int. J. Adv. Comput. Sci. Appl. (IJACSA)*, vol. 15, no. 8, pp. 942-949, 2024, doi: 10.14569/IJACSA.2024.0150893.
- [5] K. Kamdan, M. P. Anugrah, M. J. Almutaali, R. Ramdani, and I. L. Kharisma, "Performance Analysis of IndoBERT for Detection of Online Gambling Promotion in YouTube Comments," in *Proc. 7th Int. Global Conf. Ser. ICT Integration Technical Education & Smart Society*, MDPI, Sep. 2025, p. 66, doi: 10.3390/engproc2025107066.
- [6] S. P. Santika, M. A. Saputra, A. Zahra, and D. Suhartono, "Online Gambling Promotion Detection in Indonesian YouTube Comments Using Semi-Supervised IndoBERT Classification," *Procedia Comput. Sci.*, vol. 269, pp. 1269-1278, 2025, doi: 10.1016/j.procs.2025.09.068.
- [7] S. A. Nugraha, C. Lestari, K. B. Sanjaya, R. A. Naya, and J. Jolie, "Comparative Analysis of IndoBERT and mBERT for Online Gambling Comment Detection in Indonesian Social Media," *J. Tek. Inform. (Jutif)*, vol. 7, no. 2, pp. 1931-1943, Apr. 2026, doi: 10.52436/1.jutif.2026.7.2.5677.
- [8] M. C. T. Manullang, A. Z. Rakhman, H. Tantriawan, and A. Setiawan, "Comparative Analysis of CNN, Transformers, and Traditional ML for Classifying Online Gambling Spam Comments in Indonesian," *J. Appl. Inform. Comput. (JAIC)*, vol. 9, no. 3, pp. 592-602, 2025, doi: 10.30871/jaic.v9i3.9468.
- [9] F. M. Al Simabua and L. Alfat, "Optimization of IndoBERT-Lite Fine-Tuning for Spam Detection in Digital Customer Services," *Sistemasi: J. Sist. Inf.*, vol. 15, no. 5, pp. 1886-1899, 2026, doi: 10.32520/stmsi.v15i5.6398.
- [10] M. B. M. Amin, G. Hakim, M. T. Maulana, M. F. Alwan, H. S. Anggraheni, M. J. Naufal, and N. Yudistira, "Deteksi Spam Berbahasa Indonesia berbasis Teks menggunakan Model

- Bert," *J. Teknol. Inf. Dan Ilmu Komput. (JTIK)*, vol. 11, no. 6, pp. 1291-1302, Des. 2024, doi: 10.25126/jtiik.2024118121.
- [11] A. Ghourabi and M. Alohal, "Enhancing Spam Message Classification and Detection Using Transformer-Based Embedding and Ensemble Learning," *Sensors*, vol. 23, no. 8, pp. 1-17, 2023, doi: 10.3390/s23083861.
- [12] N. Arifin, U. Enri, and N. Sulistiyowati, "Penerapan Algoritma Support Vector Machine (SVM) dengan TF-IDF N-Gram untuk Text Classification," *STRING (Satuan Tulisan Ris. dan Inov. Teknol.)*, vol. 6, no. 2, p. 129, 2021, doi: 10.30998/string.v6i2.10133.
- [13] A. Gasparetto, M. Marcuzzo, A. Zangari, and A. Albarelli, "A Survey on Text Classification Algorithms: From Text to Predictions," *Information*, vol. 13, no. 2, p. 83, Feb. 2022, doi: 10.3390/INFO13020083.
- [14] M. Y. Ridho and E. Yulianti, "From Text to Truth: Leveraging IndoBERT and Machine Learning Models for Hoax Detection in Indonesian News," *J. Ilm. Tek. Elektro Komput. dan Inform. (JITEKI)*, vol. 10, no. 3, pp. 544-555, 2024, doi: 10.26555/jiteki.v10i3.29450.
- [15] H. Murfi, S. T. Gowandi, G. Ardaneswari, and S. Nurrohmah, "BERT-based Combination of Convolutional and Recurrent Neural Network for Indonesian Sentiment Analysis," *Appl. Soft Comput.*, vol. 151, p. 111112, Jan. 2024, doi: 10.1016/j.asoc.2023.111112.
- [16] R. I. Firdaus, A. A. Chamid, and A. Jazuli, "Penerapan Algoritma Support Vector Machine dalam Analisis Sentimen terhadap Ulasan Pengguna pada Aplikasi NewSakpole," *Jutisi: J. Ilm. Tek. Inform. dan Sist. Inf.*, vol. 15, no. 2, pp. 525-536, Apr. 2026, doi: 10.35889/jutisi.v15i2.3478.
- [17] F. Alvin and N. A. S. Winarsih, "Perbandingan Kinerja Model IndoBERT, IndoBERTweet, dan Algoritma Klasik pada Analisis Sentimen Isu Indonesia Gelap," *Build. Informatics Technol. Sci. (BITS)*, vol. 7, no. 3, pp. 1601-1613, Dec. 2025, doi: 10.47065/bits.v7i3.8636.
- [18] P. M. S. Ardinata, A. A. J. Permana, and I. N. S. W. Wijaya, "Identifikasi dan Normalisasi Teks Slang dengan FastText pada Twitter dalam Bahasa Indonesia," *J. Pendidik. Teknol. dan Kejuru.*, vol. 21, no. 1, pp. 34-44, 2024, doi: 10.23887/jptkundiksha.v21i1.66381.
- [19] R. M. Yazid, F. R. Umbara, and P. N. Sabrina, "Deteksi Ujaran Kebencian dengan Metode Klasifikasi Naïve Bayes dan Metode N-Gram pada Dataset Multi-Label Twitter Berbahasa Indonesia," *Informatics and Digital Expert (INDEX)*, vol. 4, no. 2, pp. 46-52, 2023, doi: 10.36423/index.v4i2.894.
- [20] M. V. S. Handayani and Muljono, "Arsitektur Hibrida IndoBERTweet-CNN untuk Klasifikasi Ujaran Kebencian Berbahasa Gaul di Media Sosial," *Infotekmesin*, vol. 17, no. 1, pp. 39-47, Jan. 2026, doi: 10.35970/infotekmesin.v17i1.3026.
- [21] A. N. Khoruzhaya, D. V. Kozlov, K. M. Arzamasov, and E. I. Kremneva, "Comparison of an Ensemble of Machine Learning Models and the BERT Language Model for Analysis of Text Descriptions of Brain CT Reports to Determine the Presence of Intracranial Hemorrhage," *Sovremennye Tehnologii v Medicine*, vol. 16, no. 1, p. 27, 2024, doi: 10.17691/stm2024.16.1.03.
- [22] M. Putri, M. Afdal, R. Novita, and M. Mustakim, "Perbandingan Evaluasi Kernel Support Vector Machine dalam Analisis Sentimen Chatbot AI pada Ulasan Google Play Store," *J. Teknol. Sist. Inf. dan Apl. (JTSI)*, vol. 7, no. 3, pp. 1236-1245, Jul. 2024, doi: 10.32493/jtsi.v7i3.41354.
- [23] N. A. Nevrada and M. A. Syaputra, "Sentiment Analysis of Telegram App Reviews on Google Play Store Using the Support Vector Machine (SVM) Algorithm," *J. Appl. Inform. Comput. (JAIC)*, vol. 9, no. 1, pp. 96-105, Jan. 2025, doi: 10.30871/jaic.v9i1.8851.
- [24] S. A. S. Mola, D. L. B. Baun, I. O. Nunes, and M. M. A. R. Sani, "Analisis Sentimen Aplikasi Halo BCA di Google Play Store Menggunakan Metode Naive Bayes, Support Vector Machine, dan Random Forest," *HOAQ: High Educ. Organ. Arch. Qual. – J. Teknol. Inf.*, vol. 15, no. 2, pp. 69-79, Dec. 2024, doi: 10.52972/hoaq.vol15no2.p69-79.
- [25] P. Arsi and R. Waluyo, "Analisis Sentimen Wacana Pemindahan Ibu Kota Indonesia Menggunakan Algoritma Support Vector Machine (SVM)," *J. Teknol. Inf. dan Ilmu Komput. (JTIK)*, vol. 8, no. 1, p. 147, 2021, doi: 10.25126/jtiik.0813944.
- [26] A. A. Chamid, R. Nindyasari, N. Azizah, and A. Hariyadi, "Analysis of Public Opinion on the Governor Candidate Debate Using LDA and IndoBERT," *Kinetik: Game Technol. Inf. Syst. Comput. Netw. Comput. Electron. Control*, vol. 10, no. 3, pp. 295-306, Aug. 2025, doi: 10.22219/kinetik.v10i3.2221.

- [27] A. A. Chamid, Widowati, and R. Kusumaningrum, "Graph-Based Semi-Supervised Deep Learning for Indonesian Aspect-Based Sentiment Analysis," *Big Data Cogn. Comput.*, vol. 7, no. 1, p. 5, Mar. 2023, doi: 10.3390/bdcc7010005.
- [28] A. Jazuli, Widowati, A. A. Chamid, and R. Kusumaningrum, "Transformer-Based Semantic Indexing for Aspect-Based Sentiment Analysis Using an Enhanced Index Generation Algorithm with BERT," *Int. J. Adv. Technol. Eng. Explor. (IJATEE)*, vol. 12, no. 127, pp. 907-926, 2025, doi: 10.19101/IJATEE.2024.111102114.
- [29] A. Jazuli, Widowati, and R. Kusumaningrum, "Optimizing Aspect-Based Sentiment Analysis Using BERT for Comprehensive Analysis of Indonesian Student Feedback," *Appl. Sci.*, vol. 15, no. 1, pp. 1-28, 2025, doi: 10.3390/app15010172.
- [30] A. A. Chamid, Widowati, and R. Kusumaningrum, "Labeling Consistency Test of Multi-Label Data for Aspect and Sentiment Classification Using the Cohen Kappa Method," *Ingénierie des Systèmes d'Information*, vol. 29, no. 1, pp. 161-167, 2024, doi: 10.18280/isi.290118.
- [31] A. A. Chamid, W. Widowati, and R. Kusumaningrum, "Text Data Labeling Process for Semi-Supervised Learning Modeling," in *12th Int. Semin. New Paradigm Innov. Nat. Sci. Its Appl. (12th ISNPINSA), AIP Conf. Proc.*, vol. 3165, p. 030011, 2024, doi: 10.1063/5.0216320.
- [32] A. A. Chamid, Widowati, and R. Kusumaningrum, "Multi-Label Text Classification on Indonesian User Reviews Using Semi-Supervised Graph Neural Networks," *ICIC Express Lett.*, vol. 17, no. 10, pp. 1075-1084, 2023, doi: 10.24507/icicel.17.10.1075.