



Tersedia online di www.journal.unipdu.ac.id
Unipdu

Halaman jurnal di www.journal.unipdu.ac.id/index.php/teknologi



Deteksi Promosi Judi Online Di Komentar Youtube Menggunakan IndoBERT *Fine-Tuning* Dengan Pendekatan *Focal Loss*

Heni Prasetyo ^a, Ulfa Kusnia ^b

^aProgram Studi Ekonomi Syariah, Universitas Islam Negeri Kiai Ageng Muhammad Besari Ponorogo, Ponorogo, Indonesia

^bProgram Studi Informatika, Universitas Islam Darul 'Ulum Lamongan, Lamongan, Indonesia

email: ^a* heni.prasetyo@uinponorogo.ac.id

*Korespondensi

Dikirim 22 Mei 2026; Direvisi 09 Juni 2026; Diterima 11 Juni 2026; Diterbitkan 27 Juni 2026

Abstrak

Pesatnya perkembangan teknologi informasi menjadikan promosi lebih efektif, efisien dan memiliki jangkauan yang lebih luas. Salah satu strategi promosi yang diterapkan adalah dengan melakukan promosi melalui kolom komentar YouTube. Hal tersebut kemudian banyak disalahgunakan oleh pihak tertentu untuk mempromosikan konten ilegal, termasuk promosi judi online. Meskipun YouTube telah menerapkan algoritma deteksi spam dan pelanggaran kebijakan komunitas, banyak komentar promosi judi yang masih lolos dari sistem. Oleh karena itu, untuk membantu dalam mendukung pemberantasan judi online, penelitian ini mengembangkan model kecerdasan buatan yang mampu mengklasifikasi dan mendeteksi promosi judi online secara otomatis. Metode yang digunakan dalam penelitian ini dilakukan dengan pendekatan Natural Language Processing menggunakan fine-tuning model IndoBERT yang merupakan pretrained language model berbasis arsitektur BERT (Bidirectional Encoder Representations from Transformers) dan disesuaikan untuk Bahasa Indonesia. IndoBERT dilatih secara khusus menggunakan data dari Wikipedia Bahasa Indonesia, artikel berita dari Kompas, Tempo, dan Liputan6, serta korpus web Indonesia. Hasil evaluasi menunjukkan bahwa kombinasi IndoBERT dengan pendekatan Focal Loss mendapat akurasi 97%. Penelitian ini menunjukkan bahwa penggabungan model berbasis bahasa alami dengan pendekatan Focal Loss dapat memberikan hasil yang akurat dalam mendeteksi komentar promosi judi online.

Kata Kunci: Deteksi Promosi Judi, Komentar YouTube, IndoBERT, Fine Tuning, Focal Loss

Detection of Online Gambling Promotions in YouTube Comments Using IndoBERT Fine-Tuning with a Focal Loss Approach

Abstract

The rapid development of information technology has made promotion more effective, efficient, and capable of reaching a wider audience. One promotional strategy widely applied is promotion through YouTube comment sections. However, this method has been widely misused by certain parties to promote illegal content, including the promotion of online gambling. Although YouTube has implemented spam detection algorithms and community guideline violation filters, many gambling promotion comments still manage to bypass the system. Therefore, to support efforts in eradicating online gambling, this research develops an artificial intelligence model capable of automatically classifying and detecting online gambling promotions. The method used in this study employs a Natural Language Processing (NLP) approach by fine-tuning the IndoBERT model, a pretrained language model based on the BERT (Bidirectional Encoder Representations from Transformers) architecture and specifically adapted for the Indonesian language. IndoBERT was specially trained using data from Indonesian Wikipedia, news articles from Kompas, Tempo, and Liputan6, as well as the Indonesian web corpus. Evaluation results show that the combination of IndoBERT with the Focal Loss approach achieved an accuracy of 97%. This research demonstrates that integrating a natural language-based model with the Focal Loss approach can produce highly accurate results in detecting comments that promote online gambling.

Keywords: : Gambling Promotion Detection, YouTube Comments, IndoBERT, Fine Tuning, Focal Loss

Untuk mengutip artikel ini dengan APA Style:

Prasetyo, H., Kusnia, U. (2026). Deteksi Promosi Judi Online Di Komentar Youtube Menggunakan IndoBERT Fine-Tuning Dengan Pendekatan Focal Loss. TEKNOLOGI: Jurnal Ilmiah Sistem Informasi, 16(1), 80-87 : <https://doi.org/10.26594/teknologi.v16i1.6481>



© 2022 Penulis. Diterbitkan oleh Program Studi Sistem Informasi, Universitas Pesantren Tinggi Darul Ulum. Ini adalah artikel open access di bawah lisensi CC BY-NC-NA (<https://creativecommons.org/licenses/by-nc-sa/4.0/>).

1. Pendahuluan

Perjudian daring (judi online) merupakan salah satu bentuk kejahatan siber yang berkembang pesat di Indonesia. Indonesia menjadi negara tertinggi pemain judi online, tercatat pemain judi online di Indonesia sebanyak 4.000.000 orang. Pemain judi online, tidak hanya berasal usia dewasa tetapi juga anak-anak. Berdasarkan data demografi, pemain judi online usia di bawah 10 tahun mencapai 2% dari pemain, dengan total 80.000 orang. Sebaran pemain antara usia antara 10 tahun s.d. 20 tahun sebanyak 11% atau kurang lebih 440.000 orang, kemudian usia 21 sampai dengan 30 tahun 13% atau 520.000 orang. Usia 30 sampai dengan 50 tahun sebesar 40% atau 1.640.000 orang dan usia di atas 50 tahun sebanyak 34% dengan

jumlah 1.350.000 orang [1]. Pemerintah Indonesia telah membentuk satgas judi online dengan dikeluarkannya Keputusan Presiden RI Nomor 21 Tahun 2024 Tentang Satuan Tugas Pemberantasan Judi Daring. Menurut data Kementerian Komunikasi dan Digital (Komdigi) hingga April 2025, komdigi telah menangani lebih dari 1,3 konten perjudian online yang terdiri dari 1.192.000 situs judi dan 127.000 konten judi di media sosial dengan cara memblokir konten tersebut [2], namun promosi terus bermutasi ke bentuk yang lebih sulit dideteksi. Salah satu kanal promosi yang paling efektif saat ini adalah kolom komentar YouTube. Pelaku memanfaatkan anonimitas dan volume komentar yang sangat besar untuk menjangkau audiens muda yang menjadi target utama.

Komentar promosi judi online biasanya disamarkan dengan bahasa tidak langsung, penggunaan emoji, singkatan, kode tertentu (contoh: "Jepe", "Gacor", "Hoki", "Jackpot"), serta ajakan untuk bergabung melalui link website. Moderasi manual oleh uploader atau tim YouTube tidak lagi efektif karena volume komentar pada video populer dapat mencapai puluhan ribu per hari. Oleh karena itu, diperlukan sistem deteksi otomatis berbasis pembelajaran mesin yang mampu memahami konteks bahasa Indonesia sehari-hari.

Model bahasa berbasis transformer seperti BERT telah terbukti sangat baik dalam tugas klasifikasi teks. Namun, model BERT multibahasa (mBERT) kurang optimal untuk bahasa Indonesia karena hanya dilatih pada porsi kecil data bahasa Indonesia. IndoBERT, yang sepenuhnya dilatih pada korpus bahasa Indonesia, menawarkan performa yang jauh lebih baik. Tantangan utama dalam kasus ini adalah ketidakseimbangan kelas yang sangat ekstrem (rasio positif:negatif bisa mencapai 1:50 atau lebih) sehingga model cenderung bias terhadap kelas mayoritas dan menghasilkan recall rendah pada kelas promosi.

Untuk mengatasi masalah tersebut, penelitian ini menggunakan pendekatan Focal Loss yang dirancang khusus untuk menangani hard examples dan class imbalance pada tahap pelatihan. Dengan menggabungkan kemampuan representasi bahasa IndoBERT dan mekanisme Focal Loss, diharapkan diperoleh model yang memiliki precision tinggi (meminimalkan false positive) sekaligus recall yang memadai pada kelas promosi judi online.

2. State of the Art

Penelitian terkait klasifikasi teks dengan menggunakan metode IndoBERT Fine-Tuning telah banyak dilakukan sebelumnya dan menunjukkan bahwa metode tersebut mampu menghasilkan klasifikasi yang akurat. Juarto dkk. (2023) misalnya, mengklasifikasikan berita berbahasa Indonesia dengan menggunakan beberapa metode dan metode IndoBERT mencatat akurasi sebesar 94.5%, metode tersebut menghasilkan akurasi tertinggi dibandingkan metode yang lain. Studi lain oleh Kunaefi dkk. (2025) menggunakan IndoBERT dalam klasifikasi berita hoaks bahasa Indonesia dan mendapatkan akurasi sebesar 97.9% kemudian setelah menambahkan pendekatan focal loss akurasi meningkat menjadi 98.3%. Hal tersebut membuktikan bahwa penggunaan pendekatan focal loss dapat meningkatkan performa dari metode IndoBERT. Rijal dkk. (2025) juga melakukan penelitian terkait dengan berita hoaks yang tersebar di media sosial dengan menggunakan metode IndoBERT dan mendapatkan akurasi sebesar 92%, F1-Score 92%, dan ROC 98%. Tobing dkk (2025) melakukan penelitian dengan mendeteksi berita hoaks khusus kategori politik, penelitian tersebut menggunakan dataset berjumlah 20.928 artikel berita faktual dan 2.251 artikel berita hoaks. Hasil yang didapatkan dari penelitian tersebut adalah akurasi sebesar 95% dan ROC AUC 94.6%

3. Metode Penelitian

Penelitian ini dilakukan dengan menggunakan pendekatan kuantitatif, untuk melaksanakan penelitian ini dibutuhkan tahapan-tahapan yang harus dilakukan agar penelitian dapat berjalan secara terarah dan dapat menghasilkan output yang sesuai dengan tujuan. Proses dimulai dengan membuat desain penelitian, mengumpulkan data, memproses data, implementasi metode dan yang terakhir adalah melakukan evaluasi terhadap hasil yang didapatkan dari implementasi metode.

3.1. Desain Penelitian

Penelitian dimulai dengan membuat desain alur penelitian sebagai landasan dalam proses pelaksanaan penelitian.



Gambar 1. Desain penelitian

Metode yang digunakan dalam penelitian ini adalah IndoBERT, dimana IndoBERT merupakan adaptasi dari BERT untuk bahasa Indonesia yang telah dilatih secara khusus dengan menggunakan data dari Wikipedia Bahasa Indonesia, artikel berita dari Kompas, Tempo, dan Liputan6 serta korpus web Indonesia.

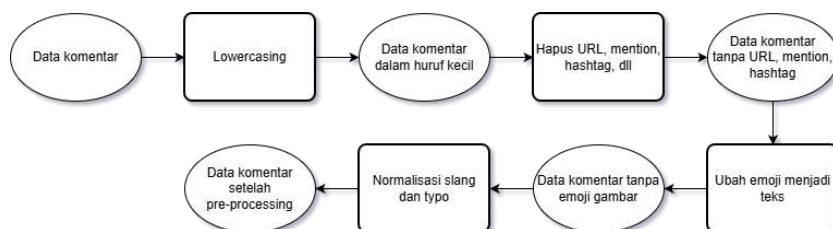
3.2. Pengumpulan Data

Data yang dibutuhkan dalam proses penelitian ini adalah data komentar dari platform Youtube. Proses pengumpulan data dalam penelitian ini dilakukan dengan mengcrawling komentar dari platform Youtube dengan menggunakan Youtube API. Untuk bisa menggunakan Youtube API, user harus mendaftar terlebih dahulu sebagai developer dan kemudian akan mendapatkan api key dan secret key.

Data komentar yang telah dikumpulkan kemudian dilabeli secara manual dengan melihat username dan komentar untuk menentukan komentar tersebut merupakan promosi judi online atau bukan. Hal tersebut bertujuan untuk mengklasifikasi komentar dan nantinya akan digunakan sebagai data training.

3.3. Text Pre-Processing

Data yang akan digunakan dalam penelitian ini adalah data berbentuk teks sehingga dibutuhkan pemrosesan data sebelum data digunakan. Pemrosesan data bertujuan untuk menghilangkan noise dan mendapatkan fitur-fitur yang dapat digunakan untuk proses deteksi.



Gambar 2. Alur pre-processing data

Gambar 2 menunjukkan tahapan dalam text pre-processing, tahapan dalam text pre-processing terdiri atas :

- Lowercasing

Tahapan pertama dalam text pre-processing data adalah lowercasing atau mengubah semua huruf menjadi huruf kecil.

Tabel 1. Lowercasing

Input	Output
Bakal gw ulang podcast ini buat naikin dopamin. Moodbooster banget 🚀 , makasih sudah pertemukan mereka ber 3. Semoga tim kreatif baca ini , buatkan mereka program . Pasti pecah	bakal gw ulang podcast ini buat naikin dopamin. moodbooster banget 🚀 , makasih sudah pertemukan mereka ber 3. semoga tim kreatif baca ini , buatkan mereka program . pasti pecah

- Hapus URL, mention, hashtag, dll

Tahap ini berfungsi untuk membersihkan noise yang ada dalam data komentar karena dalam proses NLP, atribut-atribut tersebut dianggap tidak memiliki makna atau arti.

Tabel 2. Hapus URL, mention, hashtag, dll

Input	Output
bakal gw ulang podcast ini buat naikin dopamin moodbooster banget 🚀 makasih sudah pertemukan mereka ber 3 semoga tim kreatif baca ini buatkan mereka program pasti pecah	bakal gw ulang podcast ini buat naikin dopamin moodbooster banget roket makasih sudah pertemukan mereka ber 3 semoga tim kreatif baca ini buatkan mereka program pasti pecah

- Ubah emoji menjadi teks

Komentar promosi judi online biasanya dipenuhi dengan penggunaan emoji, jika dihilangkan akan menghilangkan konteks. Oleh karena itu, emoji harus diubah menjadi text.

Tabel 3. Ubah emoji menjadi teks

<i>Input</i>	<i>Output</i>
bakal gw ulang podcast ini buat naikin dopamin moodbooster banget 🚀 makasih sudah pertemukan mereka ber 3 semoga tim kreatif baca ini buatkan mereka program pasti pecah	bakal gw ulang podcast ini buat naikin dopamin moodbooster banget roket makasih sudah pertemukan mereka ber 3 semoga tim kreatif baca ini buatkan mereka program pasti pecah

- Normalisasi slang dan typo

Stemming adalah proses konversi term / kata kebentuk dasar dengan menghilangkan imbuhan yang terdapat dalam kata. Imbuhan yang dihilangkan dalam proses stemming adalah awalan, sisipan, akhiran, dan kombinasi awalan dan akhiran.

Tabel 4. Normalisasi slang dan typo

<i>Input</i>	<i>Output</i>
bakal gw ulang podcast ini buat naikin dopamin moodbooster banget roket makasih sudah pertemukan mereka ber 3 semoga tim kreatif baca ini buatkan mereka program pasti pecah	bakal saya ulang podcast ini buat menaikkan dopamin moodbooster banget roket terimakasih sudah mempertemukan mereka ber 3 semoga tim kreatif baca ini membuatkan mereka program pasti pecah

3.4. Tokenisasi

Tahap tokenisasi dilakukan untuk mengubah data teks menjadi format numerik yang dapat diproses oleh model BERT. Tokenisasi bertugas memecah setiap kalimat atau dokumen menjadi unit-unit yang lebih kecil yang disebut token, yang dapat berupa kata utuh atau potongan sub-kata sesuai dengan aturan tokenisasi BERT. Setiap teks diawali dengan token khusus [CLS] yang berfungsi sebagai representasi keseluruhan kalimat, dan diakhiri dengan token [SEP] sebagai penanda akhir kalimat. Setiap token kemudian dikonversi menjadi representasi numerik berupa token ID, yang mencerminkan posisi token dalam kosakata BERT, dan attention mask, yang menandai token mana yang harus diproses oleh model dan mana yang harus diabaikan, seperti token hasil padding. Penggunaan attention mask sangat penting untuk memastikan bahwa model hanya fokus pada bagian input yang informatif, sehingga meningkatkan efektivitas pemahaman konteks dalam tugas-tugas pemrosesan bahasa alami. Oleh karena itu, tokenisasi memainkan peran penting dalam mentransformasikan data teks mentah menjadi format yang sesuai untuk proses komputasi dalam model BERT.

IndoBERT merupakan model BERT yang telah dilatih sebelumnya menggunakan kumpulan data Bahasa Indonesia (Indo4B) yang ekstensif dan telah disanitasi. IndoBERT menerapkan metode tokenisasi WordPiece, yang merupakan salah satu jenis tokenisasi sub-kata (sub-word tokenization). Metode ini sangat efektif dalam menangani kata-kata yang tidak ada dalam kamus pelatihan (out-of-vocabulary), termasuk kata tidak baku, kesalahan ejaan, atau istilah slang, dengan cara memecahnya menjadi unit subkata yang lebih kecil dan dikenali.

4. Implementasi Metode

Pada tahap ini, model IndoBERT digunakan sebagai arsitektur dasar dalam proses fine-tuning untuk menyelesaikan proses deteksi komentar promosi judi online. Fine-tuning merujuk pada proses penyesuaian lanjutan terhadap parameter-parameter model yang telah dilatih sebelumnya (pre-trained), agar model dapat beradaptasi secara optimal terhadap tugas klasifikasi yang lebih spesifik. Dalam penelitian ini, klasifikasi dilakukan terhadap dua kelas yaitu komentar promosi judi online dan komentar bukan promosi judi online.

Pemanfaatan IndoBERT yang telah melalui tahap pre-training pada korpus besar berbahasa Indonesia memberikan dasar representasi bahasa yang kuat. Namun, untuk mencapai performa yang lebih akurat pada tugas klasifikasi tertentu, diperlukan proses fine-tuning. Proses ini memungkinkan model untuk mempelajari pola-pola khusus dari data yang lebih sempit dan relevan dengan domain yang dituju, seperti karakteristik linguistik dalam komentar promosi judi online. Penyesuaian ini penting karena model pre-trained secara umum dilatih pada data yang luas dan bersifat umum, sehingga belum tentu optimal untuk tugas klasifikasi yang spesifik.

Fine-tuning pada model berbasis Transformer seperti BERT secara signifikan meningkatkan performa karena mampu menyesuaikan representasi internal model terhadap distribusi data dan label yang baru. Oleh karena itu, pendekatan ini dipilih untuk memperoleh hasil klasifikasi yang lebih akurat dan relevan dalam mendeteksi komentar promosi judi online berbahasa Indonesia.

Dataset yang digunakan dalam penelitian ini bersifat tidak seimbang, dengan distribusi atau split data sebanyak 80% data untuk kelas non-promosi judi online dan 20% data untuk kelas promosi judi online. Ketidakseimbangan ini berpotensi menyebabkan bias model terhadap kelas mayoritas. Untuk mengatasi hal tersebut, digunakan loss function dengan focal loss sebagai pengganti dari Cross-Entropy Loss. Fungsi ini pertama kali diperkenalkan dalam konteks deteksi objek pada model RetinaNet. Focal Loss dirancang untuk memberikan bobot lebih besar terhadap sampel yang sulit diklasifikasikan (dalam hal ini, berita hoaks yang berjumlah lebih sedikit), serta mengurangi kontribusi dari sampel yang lebih mudah dipelajari (berita non-hoaks yang lebih banyak).

$$FL(P_t) = -a_t(1 - P_t)^y \log(P_t)$$

Dimana :

P_t : probabilitas prediksi untuk kelas yang benar

a_t : balancing factor antara kelas

y : focus parameter

5. Evaluasi

Pada tahap evaluasi model, penelitian ini mengadopsi dua pendekatan utama, yaitu confusion matrix dan classification report, untuk menilai efektivitas model IndoBERT fine tuning dengan focal loss dalam melakukan deteksi terhadap komentar promosi judi online. Confusion matrix merupakan sebuah representasi tabular yang mengilustrasikan hubungan antara hasil prediksi model dan label aktual pada data uji. Matriks ini terdiri dari empat komponen utama, yaitu True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN), yang masing-masing mencerminkan jenis-jenis prediksi yang benar maupun salah oleh model terhadap dua kelas yang dianalisis. Dengan menyediakan rincian informasi mengenai distribusi kesalahan dan keberhasilan prediksi, confusion matrix memungkinkan peneliti untuk melakukan analisis menyeluruh terhadap keunggulan serta keterbatasan model dalam mendeteksi masing-masing kelas.

$$Precision = \frac{TP}{TP + FP}$$

$$Recall = \frac{TP}{TP + FN}$$

$$F1 - Score = \frac{2 * Precision * Recall}{Precision + Recall}$$

6. Hasil dan Pembahasan

Dataset yang digunakan untuk melakukan percobaan adalah data komentar yang dicrawl dengan menggunakan Youtube API pada video <https://www.youtube.com/watch?v=Nlr72bFGfts>. Dari hasil crawling didapatkan 2261 komentar dan hasil dari pelabelan didapatkan komentar bukan promosi judi online atau komentar normal sebanyak 1874 komentar dan komentar promosi judi online sebanyak 386 komentar. Proses uji coba dilakukan dengan tahapan sebagai berikut :

- Pemrosesan Data

Pemrosesan data dimulai dengan membersihkan data yang sudah di crawling. Hasil dari proses crawling dapat dilihat pada gambar berikut :

Gambar 4. Hasil crawling

Data yang akan digunakan dalam penelitian ini adalah data kolom Author dan Komentar karena selain komentar, author atau username yang digunakan juga mengandung promosi judi online. Tahapan pemrosesan data dalam penelitian ini adalah menggabungkan author dan komentar kemudian dilakukan proses lowercasing, membersihkan data dari URL, mention, hashtag, mengubah emoji menjadi teks, dan terakhir menormalisasi bahasa slang dan memperbaiki kesalahan penulisan atau typo. Setelah data diproses, data kemudian diklasifikasikan kedalam kelas 0 untuk komentar bukan promosi judi online dan kelas 1 untuk promosi judi online.

```

Hasil pemrosesan data :
      Hasil_Pemrosesan  Judol
0      ahmadriadi4, yang setuju bang mael tetap di po...  0
1      pulau777cariajadigo0gle, wkwwkwkwk pojok termin...  1
2      penontonw9p, semoga yg nonton yg belum dapat p...  0
3      akusukapula777, fajar sadboy emang gak pernah...  1
4      genevievejacksonn9j, sumpah.garuda ho kibenera...  1
...
2255   ccchanel3639, apa kabar saya yg hanya buruh pa...  0
2256   kazdo8438, hah?! apa hubungannya coba nonton p...  0
2257   hayazya, aamiin ya rabb.. kebetulan bgt lg car...  0
2258   mr.fuserblackwith, dih apansi, elu yg kemakan ...  0
2259   saifsyujaa3169, bang praz ini ketawanya memang...  0

[2260 rows x 2 columns]

```

Gambar 5. Hasil pemrosesan data

- Tokenisasi

Dalam model Indo-BERT, proses tokenisasi dilakukan dengan menggunakan jenis tokenisasi khusus yaitu WordPiece Tokenization. Wordpiece tokenization adalah teknik memecah kata menjadi sub-kata atau bahkan menjadi tunggal ketika kata tersebut tidak ditemukan dalam kosakata model.

```

--- INPUT_IDS ---
[2, 4592, 5175, 867, 30382, 30460, 34, 5629, 1719, 1019, 28, 830, 26, 14093, 6552, 500, 1719, 5641, 3, 0, 0, 0, ...]

--- TOKEN PER TOKEN + SPECIAL TOKEN ---
[CLS]
ahmad
#ria
#di
#4
,
yang
setuju
bang
na
#el
tetap
di
pojok
terminal
sama
bang
pras
[SEP]
[PAD]
[PAD]
[PAD]...

--- ATTENTION MASK ---
[1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 0, 0, 0, 0, 0, 0, ...]

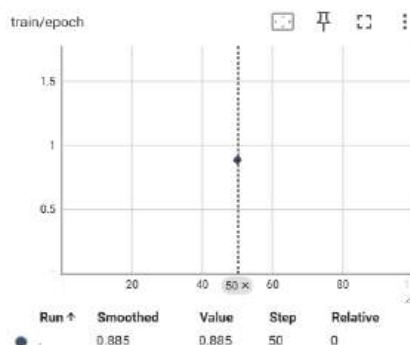
```

Gambar 6. Hasil tokenisasi

Cara kerja dari tokenisasi ini adalah kata umum dan sering muncul akan dipertahankan sebagai satu token sedangkan untuk kata yang tidak dikenal atau jarang muncul akan dipecah menjadi sub-kata. Setelah proses tokenisasi dilakukan, tahapan selanjutnya adalah mengubah token yang dihasilkan menjadi numerik termasuk token yang ditambahkan di awal dan akhir kalimat. Tujuan dari proses tokenisasi ini adalah untuk mempertahankan konteks dan membuat input standar untuk model Indo-BERT yang kemudian menghasilkan vector embedding untuk menangkap makna semantic dan hubungan kontekstual token dalam kalimat.

- Implementasi Metode

Dalam proses implementasi dataset yang telah diproses dibagi menjadi tiga dengan prosentase 70% data training, 15% data testing dan 15% data validasi. Dari hasil pembagian data didapatkan 1582 data training, 339 data testing dan 339 data validasi. Proses pelatihan data dilakukan sebanyak 3 kali dan untuk menghindari Out-of-Memory (OOM) batch_size yang akan dimuat adalah 8 sampel per step ke GPU. Di bawah ini adalah interval logging dari metrik pelatihan yang dicatat setiap 50 step.



Gambar 7. Proses training data

Output dari proses training menunjukkan model dapat mempelajari pola dalam data training dengan baik dan pelatihan model sangat efisien, hal tersebut ditunjukkan dengan kecilnya nilai training_loss 0.0248.

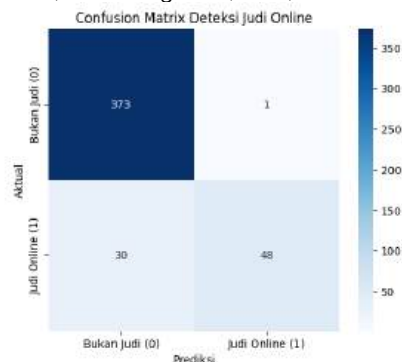
Total waktu yang dibutuhkan untuk proses pelatihan dengan menggunakan 3 epoch adalah 216.29 detik dengan kecepatan sampel 25,078 perdetik.

```
TrainOutput(
  global_step=171,
  training_loss=0.024856141959017482,
  metrics={
    'train_runtime': 216.2865,
    'train_samples_per_second': 25.078,
    'train_steps_per_second': 0.791,
    'total_flos': 1427114364272640.0,
    'train_loss': 0.024856141959017482,
    'epoch': 3.0
  }
)
```

Gambar 8. Hasil training data

- Hasil Pengujian

Setelah proses training selesai, proses terakhir dalam proses penelitian ini adalah testing. Dari hasil testing didapatkan nilai True Positive 48, True Negative, 373, False Positive 1 dan False Negatif 30.



Gambar 9. Hasil pengujian

Dari hasil tersebut maka didapatkan nilai dari precision adalah sebesar 0,98, nilai dari recall sebesar 0,62, dan nilai f1-score adalah sebesar 0,76, dan nilai akurasi adalah sebesar 0,93.

7. Kesimpulan

Berdasarkan hasil pengujian yang dilakukan terhadap data komentar sebanyak 2261 komentar dengan menggunakan Indo-BERT fine tuning dengan pendekatan focal loss didapatkan nilai precision 98%, nilai recall 62%, nilai f1-score 76% dan hasil akurasi 93%. Hasil evaluasi menunjukkan bahwa model memiliki nilai accuracy sebesar 93%, yang menandakan sebagian besar data berhasil diklasifikasikan dengan benar. Nilai precision sebesar 98% menunjukkan bahwa prediksi kelas positif yang dihasilkan model sangat akurat dengan tingkat false positive yang rendah. Namun, nilai recall sebesar 62% mengindikasikan bahwa model masih belum mampu mendeteksi seluruh data positif secara optimal. Sementara itu, nilai F1-score sebesar 76% menunjukkan bahwa keseimbangan antara precision dan recall berada pada kategori cukup baik.

8. Kontribusi Penulis

Heni Prasetyo : Data curation, Formal Analysis, Investigation, Methodology, Software, Visualization, dan Writing – original draft.

Ulfa Kusnia : Supervision, Validation, dan Writing – review & editing.

9. Declaration of Competing Interest

Penulis menyatakan tidak ada konflik kepentingan.

10. Referensi

- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In J. Burstein, C. Doran, & T. Solorio (Eds.), *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)* (pp. 4171–4186). Association for Computational Linguistics. <https://doi.org/10.18653/v1/N19-1423>
- Koto, F., Lau, J. H., & Baldwin, T. (2021). IndoBERTweet: A Pretrained Language Model for Indonesian Twitter with Effective Domain-Specific Vocabulary Initialization. <https://arxiv.org/abs/2109.04607>
- Anwari, K., & Chamidy, T. (2025). Klasterisasi Berita Berbahasa Indonesia Menggunakan Model K-Means Dan Indobert Untuk Menentukan Berita Hoaks Clustering Indonesian-Language News Using the K-Means Model and Indobert to Determine Hoax News. 17(1), 11–23. <https://doi.org/10.35891/explorit>

- Cahyo, A. N., Hermawan, R., Avin, M., Wijaya, D., & Sari, A. P. (2025). Jurnal Restikom : Riset Teknik Informatika dan Komputer Deteksi Teks Promosi Judi Online Pada Kolom Komentar YouTube dengan Metode Regex dan Fuzzy Matching. 7(2), 176–187. <https://restikom.nusaputra.ac.id>
- Jauhari, F., Riza, M., Guntara, R. G., & Nugraha, M. R. (2025). Implementasi Algoritma Naive Bayes untuk Filtrasi Spam Komentar Judi Online pada YouTube. <https://ejournal.upi.edu/index.php/IJDB>
- Jocelynne, C., Tobing, L., Lanang Wijayakusuma, I., Putu, L., & Harini, I. (2025). Detection of Political Hoax News Using Fine-Tuning IndoBERT. In Journal of Applied Informatics and Computing (JAIC) (Vol. 9, Issue 2). <http://jurnal.polibatam.ac.id/index.php/JAIC>
- Koto, F., Rahimi, A., Lau, J. H., & Baldwin, T. (2020). IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP. <http://arxiv.org/abs/2011.00677>
- Kunaefi, A., Abidin, Z., & Kusumawati, R. (2025). KLASIFIKASI BERITA HOAKS BAHASA INDONESIA MENGGUNAKAN INDOBERT FINE-TUNING DENGAN PENDEKATAN FOCAL LOSS PADA DATA TIDAK SEIMBANG. JIPI (Jurnal Ilmiah Penelitian Dan Pembelajaran Informatika), 10(2), 1706–1714. <https://doi.org/10.29100/jipi.v10i2.7811>
- Noorian, A., Harounabadi, A., & Hazratifard, M. (2024). A sequential neural recommendation system exploiting BERT and LSTM on social media posts. Complex and Intelligent Systems, 10(1), 721–744. <https://doi.org/10.1007/s40747-023-01191-4>
- Oh, H. (2021). A YouTube Spam Comments Detection Scheme Using Cascaded Ensemble Machine Learning Model. IEEE Access, 9, 144121–144128. <https://doi.org/10.1109/ACCESS.2021.3121508>
- Prisscilya, V., & Girsang, A. S. (2024). Classification of Indonesia False News Detection Using Bertopic and Indobert. Jurnal Indonesia Sosial Teknologi, 5(8). <http://jst.publikasiindonesia.id/>
- Rahmawati, A., Alamsyah, A., & Romadhony, A. (2022). Hoax News Detection Analysis using IndoBERT Deep Learning Methodology. 2022 10th International Conference on Information and Communication Technology, ICoICT 2022, 368–373.
- Ray, B., Garain, A., & Sarkar, R. (2021). An ensemble-based hotel recommender system using sentiment analysis and aspect categorization of hotel reviews. Applied Soft Computing, 98. <https://doi.org/10.1016/j.asoc.2020.106935>
- Juarto, Budi. Yulianto. (2023). Indonesian News Classification Using IndoBert. IJISAE, 11(2), 454–460. <https://orcid.org/0000-0001-9668-3819>