

Klasifikasi *Named Entity Recognition* Pada Cerita Pendek Bahasa Indonesia Menggunakan Model *Indobert*

DOI: <http://dx.doi.org/10.35889/progresif.v22i3.3731>

Creative Commons License 4.0 (CC BY –NC)



Nazifa Samsurizal^{1*}, Hendri Ahmadian², Nurrizqa³

Teknologi Informasi, UIN Ar-Raniry Banda Aceh, Aceh, Indonesia

e-mail Corresponding Author: nazifasamsurizal18@gmail.com

Abstract

Entity recognition in Indonesian short stories requires a model capable of understanding narrative language context effectively. This study aimed to implement the IndoBERT model for the Named Entity Recognition (NER) task to identify Person, Location, and Organization entities. The dataset was obtained from the Majalah Bobo e-book and organized using the BIO (Begin, Inside, Outside) format. The research included pre-processing, tokenization, label encoding, token alignment, and fine-tuning using the AutoModelForTokenClassification model. Evaluation was conducted using precision, recall, and F1-score metrics. The results showed that IndoBERT achieved F1-scores of 0.95 for Person, 0.81 for Organization, and 0.75 for Location, with a weighted average F1-score of 0.90. These results indicated that IndoBERT performed entity recognition effectively on Indonesian short stories.

Keywords: *Named Entity Recognition; IndoBERT; Indonesian short stories; Natural Language Processing; Transformer.*

Abstrak

Pengenalan entitas pada cerita pendek bahasa Indonesia memerlukan model yang mampu memahami konteks bahasa naratif secara efektif. Penelitian ini bertujuan menerapkan model IndoBERT pada tugas *Named Entity Recognition* (NER) untuk mengenali entitas Person, Location, dan Organization. *Dataset* diperoleh dari e-book Majalah Bobo dan disusun menggunakan format BIO (Begin, Inside, Outside). Penelitian dilakukan melalui tahapan *pre-processing*, tokenisasi, *label encoding*, *token alignment*, dan *fine-tuning* menggunakan model AutoModelForTokenClassification. Evaluasi dilakukan menggunakan metrik *precision*, *recall*, dan *F1-score*. Hasil pengujian menunjukkan bahwa model IndoBERT memperoleh nilai *F1-score* sebesar 0.95 pada label Person, 0.81 pada Organization, dan 0.75 pada Location, dengan *weighted average F1-score* sebesar 0.90. Hasil tersebut menunjukkan bahwa IndoBERT mampu melakukan pengenalan entitas dengan baik pada cerita pendek bahasa Indonesia.

Kata kunci: *Named Entity Recognition; IndoBERT; cerita pendek bahasa Indonesia; Natural Language Processing; Transformer.*

1. Pendahuluan

Teks merupakan salah satu media utama dalam menyampaikan informasi dalam berbagai bentuk komunikasi, baik dalam dokumen resmi, berita, maupun karya sastra seperti cerita pendek. Salah satu informasi penting yang sering muncul dalam sebuah teks adalah entitas bernama, seperti nama orang, lokasi, organisasi, maupun keterangan waktu tertentu. Dalam kajian *Natural Language Processing* (NLP), proses untuk mengenali dan mengklasifikasikan entitas tersebut dikenal dengan istilah *Named Entity Recognition* (NER) [1]. Tugas NER bertujuan untuk mengidentifikasi kata atau frasa dalam suatu teks yang merepresentasikan entitas tertentu sehingga informasi yang awalnya tidak terstruktur dapat diubah menjadi data yang lebih terorganisir dan mudah dianalisis oleh sistem komputer [2]. Teknologi ini banyak dimanfaatkan

dalam berbagai aplikasi seperti sistem pencarian informasi, analisis teks otomatis, sistem tanya jawab, serta pengolahan dokumen digital [1].

Meskipun demikian, penerapan NER pada teks naratif masih memiliki tantangan tersendiri karena struktur bahasa yang digunakan cenderung lebih kompleks dibandingkan teks formal. Berbeda dengan teks formal yang umumnya memiliki struktur bahasa yang jelas dan eksplisit, teks naratif seperti cerita pendek sering kali menggunakan gaya bahasa yang lebih bebas, kontekstual, dan bersifat implisit. Penggunaan julukan, metafora, atau penyebutan tokoh yang tidak selalu secara langsung menyebutkan nama sebenarnya menyebabkan proses ekstraksi informasi dari teks naratif menjadi lebih kompleks dibandingkan dengan teks formal [2]. Pada teks naratif, penyebutan entitas sering kali muncul dalam bentuk yang lebih variatif, misalnya melalui panggilan tokoh, kata ganti, atau deskripsi tidak langsung. Hal tersebut menimbulkan tantangan tersendiri bagi sistem NER karena model harus mampu memahami konteks kalimat secara lebih mendalam untuk mengenali entitas yang dimaksud [3].

Penelitian mengenai NER pada bahasa Indonesia telah berkembang cukup pesat, namun sebagian besar penelitian masih berfokus pada teks formal seperti berita, dokumen resmi, maupun media sosial. Anwar et al. menerapkan metode *TextRank* dan *Named Entity Recognition* (NER) dengan pembobotan TF-IDF untuk melakukan ekstraksi kata kunci pada media berita daring, serta memperoleh nilai F-measure sebesar 0,648 [4]. Penelitian lain oleh Agustini et al. menggunakan augmentasi data berbasis GPT-4 dengan model BiLSTM-CRF dan menghasilkan peningkatan performa hingga *Macro F1-score* 0,75 dan *Micro F1-score* 0,95 [5].

Selain itu, penelitian Muhammad dan Hasibuan melakukan *fine-tuning* IndoBERT menggunakan *dataset* NERGRIT dan memperoleh peningkatan *F1-score* dari 72,38% menjadi 83,67% [2]. Penelitian lain oleh Muchtar et al. menunjukkan bahwa IndoBERT mampu memberikan performa yang stabil dan konsisten dalam mengenali entitas PER, ORG, LOC, DATE, dan CRD pada transkrip video politik berbahasa Indonesia [6]. Dari beberapa penelitian tersebut, sebagian besar *dataset* yang digunakan masih berasal dari teks formal sehingga kemampuan model dalam mengenali entitas pada teks sastra atau cerita pendek belum banyak dianalisis.

Dalam beberapa tahun terakhir, pendekatan berbasis *deep learning*, khususnya model berbasis arsitektur *Transformer*, menunjukkan performa yang baik dalam berbagai tugas NLP. Salah satu model yang banyak digunakan adalah BERT (*Bidirectional Encoder Representations from Transformers*) yang memiliki kemampuan memahami konteks kata secara dua arah dalam sebuah kalimat [7]. Pada bahasa Indonesia, pengembangan model BERT diwujudkan melalui IndoBERT yang dilatih menggunakan korpus berbahasa Indonesia dalam skala besar. Proses pelatihan tersebut memungkinkan model memahami karakteristik linguistik bahasa Indonesia secara lebih baik dibandingkan model multibahasa umum [8].

Selain itu, performa model NLP bahasa Indonesia juga dievaluasi melalui benchmark IndoLEM yang menyediakan standar pengujian berbagai tugas NLP termasuk *Named Entity Recognition* [9]. Dengan kemampuan representasi kontekstual yang dimiliki IndoBERT, model ini dinilai relevan untuk menangani kompleksitas bahasa pada teks cerita pendek yang memiliki struktur kalimat lebih variatif dan kontekstual [10].

Berdasarkan permasalahan dan celah penelitian tersebut, penelitian ini bertujuan untuk mengimplementasikan dan mengevaluasi kemampuan model IndoBERT dalam melakukan *Named Entity Recognition* pada cerita pendek bahasa Indonesia. Penelitian ini berfokus pada identifikasi entitas bernama seperti nama tokoh, lokasi, dan organisasi dalam teks naratif, serta menganalisis tantangan linguistik yang muncul akibat kompleksitas gaya bahasa dan struktur kalimat pada cerita pendek. Berbeda dengan penelitian sebelumnya yang umumnya menggunakan *dataset* teks formal, penelitian ini menggunakan cerita pendek bahasa Indonesia sebagai domain teks naratif untuk evaluasi model IndoBERT pada tugas NER.

2. Metodologi

Pendekatan kuantitatif dengan metode eksperimen diterapkan dalam penelitian ini untuk mengevaluasi kemampuan model IndoBERT dalam melakukan *Named Entity Recognition* (NER) pada cerita pendek bahasa Indonesia [11]. Penelitian dilakukan menggunakan model IndoBERT berbasis arsitektur *Transformer* untuk mengenali entitas bernama pada teks naratif bahasa Indonesia. Proses penelitian meliputi beberapa tahapan pengumpulan data hingga evaluasi model dalam mengenali entitas *Person*, *Location*, dan *Organization*. Alur serta tahapan penelitian yang diterapkan dalam penelitian ini disajikan pada Gambar 1.



Gambar 1. Alur Tahapan Penelitian

Penelitian ini dilakukan melalui beberapa tahapan, yaitu pengumpulan data, *pre-processing* data, *fine-tuning* model IndoBERT, dan evaluasi model. Setiap tahapan dilakukan secara berurutan untuk menghasilkan model *Named Entity Recognition* (NER) yang mampu mengenali entitas pada cerita pendek bahasa Indonesia.

2.1 Dataset dan Pembagian Data

Dataset yang digunakan pada penelitian ini berupa cerita pendek bahasa Indonesia yang diperoleh dari e-book Majalah Bobo. *Dataset* disusun menggunakan format BIO (*Begin*, *Inside*, *Outside*) dengan label entitas PER (*Person*), LOC (*Location*), dan ORG (*Organization*) [12]. Data kemudian disimpan dalam format CSV untuk memudahkan proses pengolahan menggunakan Python dan *library* Transformers. Format BIO digunakan untuk membedakan token awal entitas, lanjutan entitas, dan token di luar entitas.

Dalam proses pelatihan dan evaluasi model, *dataset* terlebih dahulu dipisahkan ke dalam data *training*, *validation*, dan *testing*. Pembagian tersebut bertujuan untuk melatih model IndoBERT sekaligus mengevaluasi kemampuannya dalam melakukan prediksi terhadap data yang belum pernah digunakan selama proses pelatihan. Data *validation* digunakan untuk memantau performa model selama proses *fine-tuning*, sedangkan data *testing* digunakan untuk mengukur kemampuan model pada data baru yang belum pernah diproses sebelumnya.

2.2 Pre-processing Data

Sebelum memasuki tahap pelatihan model IndoBERT, *dataset* terlebih dahulu melalui proses *pre-processing* untuk menyiapkan data agar sesuai dengan kebutuhan model. *Dataset* diproses dengan memisahkan setiap kalimat menjadi token beserta label entitasnya [13]. Selanjutnya dilakukan tokenisasi menggunakan tokenizer IndoBERT untuk mengubah teks menjadi representasi token berbasis *Transformer*.

Selain tokenisasi, dilakukan proses *label encoding* untuk mengubah label entitas menjadi bentuk numerik. Setelah itu dilakukan *token alignment* untuk menyesuaikan label entitas dengan hasil tokenisasi subword IndoBERT sehingga setiap token tetap memiliki label yang sesuai. Tahap ini dilakukan agar model IndoBERT dapat memahami hubungan antara token dan label entitas secara lebih akurat selama proses pelatihan. Selain itu, penggunaan tokenisasi berbasis subword membantu model dalam mengenali kata yang tidak umum maupun variasi penulisan kata pada cerita pendek bahasa Indonesia. Hasil dari tahap *pre-processing* berupa *dataset* terstruktur yang siap digunakan pada proses *fine-tuning* dan evaluasi model *Named Entity Recognition*.

2.3 Fine-tuning Model IndoBERT

Tahap berikutnya adalah proses *fine-tuning* menggunakan model IndoBERT yang berbasis arsitektur *Transformer*. Model tersebut dipilih karena telah dilatih menggunakan korpus bahasa Indonesia sehingga memiliki kemampuan yang lebih baik dalam memahami konteks bahasa Indonesia [6]. Pada tahap ini, *dataset* hasil *pre-processing* digunakan untuk melatih model menggunakan *library* Transformers.

Proses pelatihan dilakukan menggunakan model *AutoModelForTokenClassification* untuk tugas *Named Entity Recognition*. Selama proses pelatihan, model mempelajari hubungan antara token dan label entitas sehingga mampu mengenali pola entitas pada teks cerita pendek bahasa Indonesia.

2.4 Evaluasi Model

Tahap terakhir adalah evaluasi model untuk mengukur performa IndoBERT dalam melakukan *Named Entity Recognition*. Evaluasi dilakukan menggunakan data *testing* dengan metrik *precision*, *recall*, dan *F1-score* [14].

Precision digunakan untuk menilai ketepatan prediksi entitas yang dihasilkan model, sedangkan *recall* menunjukkan kemampuan model dalam mengidentifikasi seluruh entitas yang relevan. Sementara itu, *F1-score* digunakan sebagai ukuran yang mengukur keseimbangan antara nilai *precision* dan *recall* [15]. Hasil evaluasi digunakan untuk mengetahui performa model IndoBERT dalam mengenali entitas pada cerita pendek bahasa Indonesia.

3. Hasil dan Pembahasan

Hasil dan pembahasan pada penelitian ini memaparkan implementasi model IndoBERT pada tugas *Named Entity Recognition* (NER) untuk cerita pendek bahasa Indonesia. Proses penelitian dilakukan menggunakan *dataset* cerita pendek yang telah diberi label entitas dengan format BIO (*Begin*, *Inside*, *Outside*). Selanjutnya, *dataset* tersebut dipisahkan ke dalam data *training*, *validation*, dan *testing*. Model IndoBERT kemudian dilakukan proses *fine-tuning* untuk mengenali entitas seperti nama orang, lokasi, dan organisasi pada teks naratif. Evaluasi model dilakukan menggunakan metrik *precision*, *recall*, dan *F1-score* untuk mengukur performa model dalam mengklasifikasikan entitas pada cerita pendek bahasa Indonesia.

3.1 Pengumpulan Data

Pengumpulan data menghasilkan sekumpulan cerita pendek bahasa Indonesia yang digunakan sebagai objek penelitian pada tugas *Named Entity Recognition* (NER). *Dataset* yang digunakan berasal dari *e-book* Majalah Bobo dan disusun dalam format BIO (*Begin*, *Inside*, *Outside*). Label entitas yang digunakan terdiri dari PER (*Person*), LOC (*Location*), dan ORG (*Organization*).

Dataset yang digunakan terdiri dari 62 cerita pendek dengan total 1000 kalimat dan 10.208 token yang telah diberi label entitas sesuai kategorinya masing-masing. Cerita pendek yang digunakan memiliki karakteristik Bahasa Indonesia semi-formal yang umum digunakan pada bacaan anak-anak, sehingga mengandung variasi penyebutan tokoh, lokasi, dan organisasi dalam konteks naratif. Panjang cerita yang digunakan bervariasi sesuai dengan alur dan jumlah dialog yang terdapat pada masing-masing cerita, dengan rata-rata sekitar 16 kalimat per cerita.

Hasil pengumpulan data berupa kumpulan token yang telah diberi label entitas dan siap digunakan pada tahap *pre-processing* serta pelatihan model IndoBERT. *Dataset* kemudian dibagi menjadi data *training* (70%), *validation* (10%), dan *testing* (20%) untuk mendukung proses pelatihan dan evaluasi model. Pembagian data dilakukan agar model dapat mempelajari pola entitas sekaligus mengevaluasi kemampuannya dalam melakukan prediksi terhadap data yang tidak digunakan pada tahap pelatihan.

3.2 Pre-processing

Melalui tahap *pre-processing*, *dataset* dipersiapkan dalam bentuk yang terstruktur sehingga dapat langsung digunakan pada proses pelatihan model IndoBERT. Setiap kalimat pada cerita pendek diproses menjadi token beserta label entitas menggunakan format BIO (*Begin*, *Inside*, *Outside*). Hasil pemrosesan menunjukkan bahwa token yang termasuk entitas nama orang, lokasi, dan organisasi telah berhasil diberi label sesuai kategori masing-masing.

Selain itu, hasil tokenisasi menggunakan tokenizer IndoBERT menghasilkan representasi token yang dapat diproses oleh model berbasis *Transformer*. Data hasil *pre-processing* selanjutnya digunakan pada tahap *fine-tuning* model IndoBERT. Contoh hasil *pre-processing* ditunjukkan pada Tabel 1.

Tabel 1. Alur Tahapan Penelitian

Token	Label
Dewi	B-PER
Starda	I-PER
segera	O
melapor	O
ke	O
kahyangan	B-LOC
.	O

Hasil *pre-processing* menunjukkan bahwa setiap token telah memiliki label entitas yang sesuai sehingga *dataset* dapat digunakan untuk proses pelatihan dan evaluasi model *Named Entity Recognition*. Label B-PER digunakan untuk menandai awal entitas nama orang, sedangkan I-PER digunakan untuk menandai lanjutan token yang masih termasuk dalam entitas yang sama. Selain itu, label B-LOC digunakan untuk menandai awal entitas lokasi, sedangkan token yang tidak termasuk ke dalam entitas tertentu diberi label O (Outside).

3.3 Hasil *fine-tuning* Model IndoBERT

Tahap *fine-tuning* dilakukan menggunakan model IndoBERT untuk melatih sistem dalam mengenali entitas pada cerita pendek bahasa Indonesia. Proses pelatihan dilakukan menggunakan data training dan data validation yang telah melalui tahap *pre-processing*.

Proses *fine-tuning* dilakukan menggunakan lingkungan komputasi Google Colab dengan spesifikasi hardware berupa CPU. Implementasi model menggunakan framework PyTorch dan *library* Transformers dari Hugging Face untuk pelatihan model IndoBERT. Selain itu, penelitian ini memanfaatkan *library* Pandas untuk pengolahan data, Segeval untuk evaluasi *Named Entity Recognition* (NER), serta Scikit-learn untuk analisis hasil klasifikasi. Pelatihan dan evaluasi model terhadap *dataset* cerita pendek bahasa Indonesia dilakukan dengan memanfaatkan lingkungan eksperimen tersebut.

Pada penelitian ini digunakan beberapa parameter pelatihan seperti *learning rate*, *batch size*, dan jumlah *epoch* untuk mengoptimalkan performa model selama proses pelatihan. Parameter yang digunakan pada proses *fine-tuning* ditunjukkan pada Tabel 2.

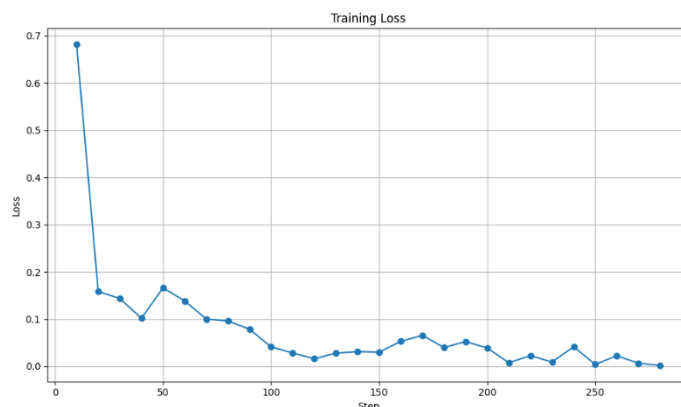
Tabel 2. Parameter *Fine-tuning* Model IndoBERT

Parameter	Nilai
<i>Learning Rate</i>	5e-5
<i>Batch Size</i>	8
<i>Epoch</i>	3
<i>Evaluation Strategy</i>	Setiap <i>Epoch</i>
<i>Save Strategy</i>	Setiap <i>Epoch</i>

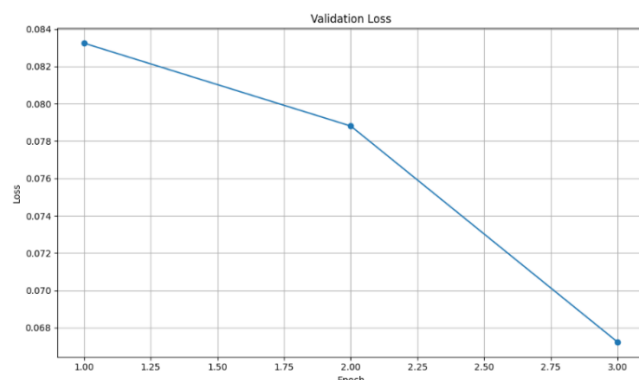
Proses pelatihan menunjukkan bahwa model IndoBERT berhasil melakukan pembelajaran dengan baik pada *dataset* cerita pendek bahasa Indonesia. Selama proses training, dilakukan pemantauan terhadap nilai *training loss* dan *validation loss* untuk melihat perkembangan performa model pada setiap *epoch*. Penggunaan evaluation strategy dan save strategy pada setiap *epoch* dilakukan untuk mengevaluasi perkembangan performa model selama proses pelatihan.

Selama proses *fine-tuning*, nilai *training loss* mengalami penurunan secara bertahap berdasarkan hasil pelatihan yang diperoleh. Selain itu, nilai *validation loss* tetap berada pada kondisi yang stabil selama proses pelatihan, yang mengindikasikan bahwa model dapat mempelajari pola entitas tanpa mengalami *overfitting* yang signifikan. Penurunan nilai *training loss* memperlihatkan bahwa model IndoBERT secara bertahap mampu mempelajari hubungan antara token dan label entitas pada *dataset* cerita pendek bahasa Indonesia selama proses *fine-tuning*.

Visualisasi perubahan nilai *training loss* dan *validation loss* ditunjukkan pada Gambar 2 dan Gambar 3.



Gambar 2. Grafik Training Loss



Gambar 3. Grafik Validation Loss

3.4 Evaluasi Model

Proses evaluasi model memanfaatkan *data testing* yang tidak digunakan selama tahap pelatihan model. Pada tahap pengujian, setiap kalimat pada data testing yang telah melalui proses *pre-processing* diproses menggunakan model IndoBERT yang telah melalui proses *fine-tuning*. Model kemudian menghasilkan prediksi label entitas untuk setiap token berdasarkan kategori PER (*Person*), LOC (*Location*), ORG (*Organization*), dan O (*Outside*). Hasil prediksi tersebut selanjutnya dibandingkan dengan label sebenarnya (*ground truth*) pada data *testing* untuk mengetahui tingkat ketepatan model dalam mengenali entitas.

Pada tahap pengujian, salah satu data pada *dataset* testing digunakan sebagai contoh untuk menunjukkan proses prediksi yang dilakukan oleh model IndoBERT terhadap setiap token. Hasil pengujian tersebut menyajikan token, label sebenarnya (*ground truth*), serta label hasil prediksi (*prediction*) yang dihasilkan oleh model. Hasil proses pengujian model ditunjukkan pada Tabel 3.

Tabel 3. Hasil Pengujian Model IndoBERT

Token	Ground Truth	Prediction
Jen	B-PER	B-PER
Chien	I-PER	I-PER
Chih	I-PER	I-PER
adalah	O	O
pedagang	O	O
yang	O	O
berasal	O	O
dari	O	O
Desa	B-LOC	B-LOC
Yutai	I-LOC	I-LOC
.	O	O

Berdasarkan Tabel 3, model IndoBERT mampu menghasilkan label entitas yang sesuai dengan label sebenarnya pada seluruh token dalam contoh data uji. Token "Jen Chien Chih" berhasil dikenali sebagai entitas *Person*, sedangkan token "Desa Yutai" dikenali sebagai entitas *Location*. Sementara token lain yang bukan merupakan entitas diberikan label O (*Outside*). Hasil tersebut menunjukkan bahwa model mampu mengenali batas awal (B-) dan token lanjutan (I-) pada setiap entitas secara tepat sehingga proses pelabelan sesuai dengan label sebenarnya (*ground truth*).

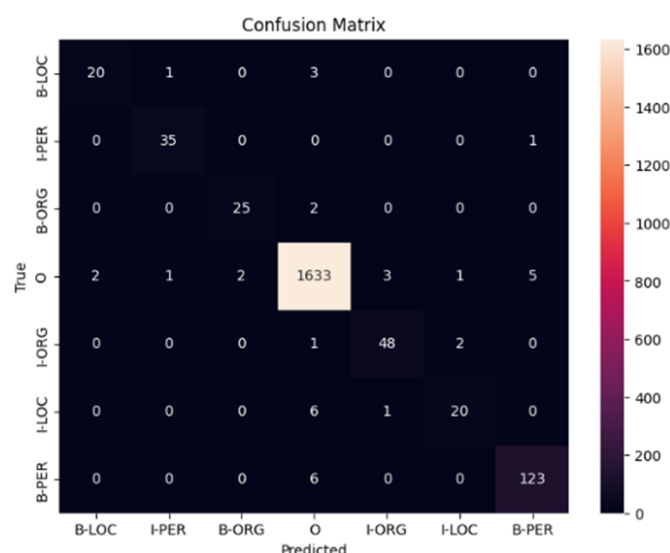
Selanjutnya, hasil evaluasi menunjukkan bahwa model IndoBERT mampu mengenali entitas dengan baik pada kategori *Person*, *Location*, dan *Organization*. Performa model dievaluasi menggunakan *classification report* yang menampilkan nilai *precision*, *recall*, dan *F1-score* pada setiap label entitas. Hasil evaluasi model ditunjukkan pada Tabel 4.

Tabel 4. Hasil Evaluasi Model IndoBERT

Label	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
LOC	0.75	0.75	0.75
ORG	0.77	0.85	0.81
PER	0.95	0.95	0.95
<i>Micro Average</i>	0.89	0.91	0.90
<i>Macro Average</i>	0.82	0.85	0.84
<i>Weighted Average</i>	0.89	0.91	0.90

Berdasarkan hasil evaluasi, model memperoleh performa terbaik pada label PER dengan nilai *F1-score* sebesar 0.95. Hasil tersebut mengindikasikan bahwa model memiliki kemampuan yang sangat baik dalam mengenali entitas nama orang pada teks cerita pendek bahasa Indonesia. Pada label ORG, model memperoleh nilai *F1-score* sebesar 0.81, sedangkan pada label LOC diperoleh nilai *F1-score* sebesar 0.75.

Selain *classification report*, pengujian juga dilakukan menggunakan *confusion matrix* untuk melihat hasil klasifikasi pada setiap label entitas. Hasil *confusion matrix* menunjukkan bahwa sebagian besar entitas berhasil diprediksi dengan benar oleh model IndoBERT.



Gambar 4. *Confusion Matrix* Model IndoBERT

Perbedaan nilai performa pada setiap label entitas dipengaruhi oleh distribusi jumlah data dan karakteristik bahasa pada cerita pendek. Label PER memperoleh performa tertinggi karena nama tokoh pada cerita pendek cenderung muncul secara konsisten dan memiliki pola penulisan yang lebih mudah dikenali oleh model. Sementara itu, label LOC dan ORG memiliki variasi konteks yang lebih beragam sehingga model memerlukan pemahaman konteks kalimat yang lebih kompleks untuk melakukan klasifikasi secara tepat.

3.5 Pembahasan

Berdasarkan hasil pengujian, model IndoBERT memperoleh nilai *weighted average F1-score* sebesar 0.90. Menurut Opitz [16], menjelaskan bahwa *precision*, *recall*, dan *F1-score* merupakan metrik yang umum digunakan untuk mengevaluasi performa model klasifikasi karena memberikan gambaran mengenai keseimbangan antara ketepatan (*precision*) dan kelengkapan (*recall*) hasil prediksi. Berdasarkan hasil evaluasi tersebut, model IndoBERT mampu mengenali entitas PER, LOC, dan ORG pada data cerita pendek bahasa Indonesia. Dengan demikian, tujuan penelitian untuk mengimplementasikan dan mengevaluasi kemampuan IndoBERT dalam mengenali ketiga jenis entitas tersebut telah tercapai.

Performa terbaik diperoleh pada label PER dengan nilai *F1-score* sebesar 0.95. Hasil tersebut mengindikasikan bahwa model IndoBERT mampu mengidentifikasi nama tokoh secara akurat, yang didukung oleh pola penulisan nama dalam cerita pendek yang relatif konsisten sehingga lebih mudah dipelajari oleh model. Sebaliknya, label LOC dan ORG memperoleh nilai *F1-score* yang lebih rendah, yaitu masing-masing sebesar 0.75 dan 0.81. Perbedaan performa tersebut dapat dipengaruhi oleh variasi konteks penggunaan entitas serta jumlah kemunculan entitas yang lebih sedikit dibandingkan entitas PER.

Hasil penelitian ini sejalan dengan konsep dasar arsitektur *Transformer* yang digunakan oleh IndoBERT. Mekanisme *self-attention* pada *Transformer* memungkinkan model memahami hubungan kontekstual antar kata dalam suatu kalimat sehingga mampu mengenali entitas secara lebih akurat [17]. Selain itu, pendekatan *bidirectional* pada BERT memungkinkan model memanfaatkan informasi dari konteks sebelum dan sesudah suatu token sehingga menghasilkan representasi bahasa yang dihasilkan lebih baik [18].

Hasil penelitian ini sejalan dengan penelitian terdahulu yang mendukung bahwa model berbasis *Transformer* memberikan performa yang baik pada tugas *Named Entity Recognition* (NER). Muhammad dan Hasibuan [2] serta Muchtar et al. [6] menunjukkan bahwa IndoBERT mampu memberikan hasil yang baik pada pengenalan entitas bahasa Indonesia. Perbedaan penelitian ini terletak pada penggunaan cerita pendek sebagai sumber data, sedangkan sebagian besar penelitian sebelumnya menggunakan teks berita, dokumen formal, atau data media sosial. Hasil yang diperoleh menunjukkan bahwa IndoBERT juga mampu bekerja dengan baik pada teks naratif yang memiliki struktur kalimat dan konteks yang lebih beragam.

Kontribusi penelitian ini adalah memberikan gambaran mengenai kemampuan IndoBERT dalam mengenali entitas pada cerita pendek bahasa Indonesia yang masih relatif jarang digunakan sebagai *dataset* NER. Hasil penelitian menunjukkan bahwa model dapat dimanfaatkan untuk mendukung proses ekstraksi informasi pada teks naratif, seperti identifikasi tokoh, lokasi, dan organisasi secara otomatis. Dengan demikian, penelitian ini dapat menjadi referensi bagi pengembangan aplikasi pemrosesan bahasa alami pada bidang pendidikan maupun analisis teks sastra [19].

Meskipun memperoleh hasil yang baik, penelitian ini masih memiliki keterbatasan, yaitu jumlah *dataset* yang digunakan relatif terbatas dan hanya mencakup tiga kategori entitas, yaitu PER, LOC, dan ORG. Selain itu, penelitian ini hanya menggunakan satu model, yaitu IndoBERT, sehingga belum dapat memberikan perbandingan dengan model *Transformer* lainnya. Oleh karena itu, penelitian selanjutnya dapat dilakukan dengan menambah jumlah *dataset*, memperluas kategori entitas, serta membandingkan performa IndoBERT dengan model lain untuk memperoleh hasil yang lebih komprehensif.

4. Simpulan

Penelitian ini berhasil menerapkan model IndoBERT untuk tugas *Named Entity Recognition* (NER) pada cerita pendek bahasa Indonesia menggunakan *dataset* yang diperoleh dari e-book Majalah Bobo. Proses penelitian dilakukan melalui tahapan *pre-processing* data menggunakan format BIO (*Begin*, *Inside*, *Outside*), tokenisasi menggunakan tokenizer IndoBERT, serta proses *fine-tuning* model berbasis arsitektur *Transformer*. Berdasarkan hasil pengujian, model IndoBERT mampu mengenali entitas *Person*, *Location*, dan *Organization* dengan performa yang baik.

Berdasarkan hasil evaluasi, label PER memperoleh nilai *F1-score* tertinggi sebesar 0.95, sedangkan label ORG dan LOC masing-masing memperoleh nilai *F1-score* sebesar 0.81 dan 0.75. Secara keseluruhan, model memperoleh nilai *weighted average F1-score* sebesar 0.90. Capaian tersebut mengindikasikan bahwa model IndoBERT memiliki kemampuan yang baik

dalam memahami konteks bahasa Indonesia pada teks naratif dan melakukan klasifikasi entitas dengan cukup baik pada cerita pendek bahasa Indonesia.

Penelitian ini juga menunjukkan bahwa pendekatan berbasis *Transformer* efektif digunakan untuk proses ekstraksi informasi pada teks berbahasa Indonesia, khususnya pada tugas NER. Meskipun demikian, performa model masih dapat ditingkatkan melalui penambahan jumlah *dataset*, variasi entitas, serta pengaturan parameter pelatihan yang lebih optimal agar hasil klasifikasi entitas menjadi lebih baik pada penelitian selanjutnya.

Daftar Referensi

- [1] M. Amien and G. F. Gunawan, "BERT dan Bahasa Indonesia: Studi tentang Efektivitas Model NLP Berbasis Transformer," *ELANG: Journal of Interdisciplinary Research*, vol. 1, no. 02, pp. 132–140, 2024.
- [2] A. F. Muhammad and M. S. Hasibuan, "View of Peningkatan Akurasi Named Entity Recognition (NER) Dengan Fine-Tuning BERT Pada Dataset Bahasa Indonesia," *CESS (Journal of Computer Engineering, System and Science)*, vol. 10, no. 02, pp. 702–713, 2025.
- [3] B. Jehangir, S. Radhakrishnan, and R. Agarwal, "A survey on Named Entity Recognition—datasets, tools, and methodologies. Natural Language Processing Journal, 3, 100017, pp. 1–12," 2023.
- [4] M. T. A. Anwar, S. H. Wijoyo, and W. H. N. Putra, "Implementasi Metode TextRank dan Named Entity Recognition Untuk Ekstraksi Kata Kunci Pada Media Online Berita," *Jurnal Sistem Informasi, Teknologi Informasi, dan Edukasi Sistem Informasi*, vol. 5, no. 1, pp. 34–41, 2024.
- [5] D. Agustini, M. I. Firdaus, M. Farida, M. E. Rosadi, H. Noor, and R. Muttaqin, "Peningkatan Kinerja Named Entity Recognition Bahasa Indonesia Melalui Augmentasi Data Berbasis Large Language Models," *Jurnal Informatika Teknologi dan Sains (Jinteks)*, vol. 7, no. 3, pp. 1361–1369, 2025.
- [6] N. U. Muchtar, D. Arifianto, and R. Umilasari, "Analisis Kinerja Transformer Untuk Named Entity Recognition (NER) Menggunakan IndoBERT Pada Transkrip Video Politik Berbahasa Indonesia," *Discovery: Jurnal Ilmu Pengetahuan*, vol. 10, no. 2, pp. 180–199, 2025.
- [7] S. Dharmawan, V. C. Mawardi, and N. J. Perdana, "Klasifikasi Ujaran Kebencian Menggunakan Metode FeedForward Neural Network (IndoBERT)," *Jurnal Ilmu Komputer dan Sistem Informasi*, vol. 11, no. 1, pp. 1–6, 2023.
- [8] M. Amien, "Sejarah dan perkembangan teknik Natural Language Processing (NLP) bahasa Indonesia: Tinjauan tentang sejarah, perkembangan teknologi, dan aplikasi NLP dalam bahasa Indonesia," *arXiv preprint arXiv:2304.02746*, pp. 1–7, 2023.
- [9] Y. O. Sihombing and N. V. Situmorang, "Prediksi Sentimen Pada Teks Media Sosial Corporate University Menggunakan RoBERTa," *Prosiding PITNAS Widyaishwara*, vol. 1, pp. 302–316, 2024.
- [10] E. Yulianti, N. Bhary, J. Abdurrohman, F. W. Dwitilas, E. Q. Nuranti, and H. S. Husin, "Named entity recognition on Indonesian legal documents: a dataset and study using transformer-based models.," *International Journal of Electrical & Computer Engineering (2088-8708)*, vol. 14, no. 5, pp. 5489–5501, 2024.
- [11] N. R. Jonathan, Syafrijon, D. Novaliendry, and K. Budayawan, "View of Ekstraksi Entitas Keterampilan pada Teks Lowongan Pekerjaan Berbahasa Indonesia Menggunakan Model IndoBERT," *Jurnal Pendidikan Tambusai*, vol. 9, no. 3, pp. 40145–40171, 2025.
- [12] W. L. Seow, I. Chaturvedi, A. Hogarth, R. Mao, and E. Cambria, "A review of named entity recognition: from learning methods to modelling paradigms and tasks," *Artif. Intell. Rev.*, vol. 58, no. 10, p. 315, 2025.
- [13] S. Khairina, N. Saffa, D. C. U. Lieharyani, and J. Hutahaeen, "Contextualized Word Embedding Untuk Ekstraksi Kutipan Berita Indonesia," *Jurnal Ilmiah Matrik*, vol. 27, no. 2, pp. 142–153, 2025.
- [14] I. I. Sholikhah, A. T. J. Harjanta, and K. Latifah, "Machine Learning Untuk Deteksi Berita Hoax Menggunakan BERT," in *Prosiding Seminar Nasional Informatika*, 2023, pp. 524–531.

-
- [15] S. C. Kusuma, T. N. Fatyanosa, and M. Data, "Fine-Tuning Bert Untuk Named Entity Recognition Pada Dokumen Klinis Gizi," *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, vol. 10, no. 2, pp. 1–10, 2026.
 - [16] J. Opitz, "A closer look at classification evaluation metrics and a critical reflection of common evaluation practice," *Trans. Assoc. Comput. Linguist.*, vol. 12, pp. 820–836, 2024.
 - [17] P. Srinivasan and R. Venkatakrisnan, "Transformer-based models for named entity recognition: A comparative study," in *2023 14th International Conference on Computing Communication and Networking Technologies (ICCCNT)*, IEEE, 2023, pp. 1–5.
 - [18] N. Patwardhan, S. Marrone, and C. Sansone, "Transformers in the real world: A survey on nlp applications," *Information*, vol. 14, no. 4, p. 242, 2023.
 - [19] I. Keraghel, S. Morbieu, and M. Nadif, "Recent advances in named entity recognition: A comprehensive survey and comparative study," *arXiv preprint arXiv:2401.10825*, pp. 1–42, 2024.