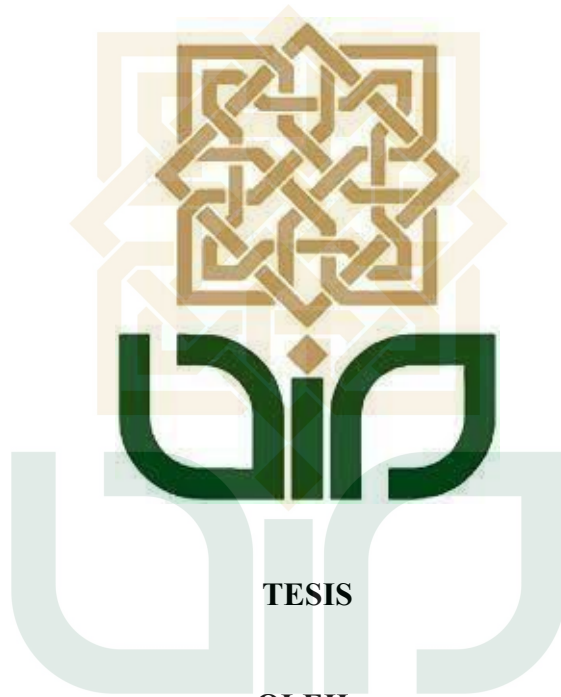


**OPTIMASI LARGE LANGUAGE MODEL: INDOBERT
DALAM KLASIFIKASI *ANAMNESIS* MEDIS BERDASARKAN
SNOMED-CT**



TESIS

OLEH:

DENY SETIAWAN

NIM: 23206051013

**PROGRAM STUDI INFORMATIKA
PROGRAM MAGISTER FAKULTAS SAINS DAN TEKNOLOGI
UIN SUNAN KALIJAGA YOGYAKARTA**

2025

PENGESAHAN TUGAS AKHIR



KEMENTERIAN AGAMA
UNIVERSITAS ISLAM NEGERI SUNAN KALIJAGA
FAKULTAS SAINS DAN TEKNOLOGI
Jl. Marsda Adisucipto Telp. (0274) 540971 Fax. (0274) 519739 Yogyakarta 55281

PENGESAHAN TUGAS AKHIR

Nomor : B-2468/Un.02/DST/PP.00.9/12/2025

Tugas Akhir dengan judul : Optimasi Model Natural Language Processing : IndoBERT dalam Klasifikasi Anamnesis Medis Berdasarkan SNOMED-CT

yang dipersiapkan dan disusun oleh:

Nama : DENY SETIAWAN, S.Kom.
Nomor Induk Mahasiswa : 23206051013
Telah diujikan pada : Senin, 03 November 2025
Nilai ujian Tugas Akhir : A-

dinyatakan telah diterima oleh Fakultas Sains dan Teknologi UIN Sunan Kalijaga Yogyakarta

TIM UJIAN TUGAS AKHIR



Ketua Sidang

Dr. Agus Mulyanto, S.Si., M.Kom., ASEAN Eng.
SIGNED

Valid ID: 693f856c252b2



Penguji I

Muhammad Mustakim, S.T. M.T.
SIGNED

Valid ID: 69169423ef8cd



Penguji II

Dr. Agung Fatwanto, S.Si., M.Kom.
SIGNED

Valid ID: 691689650a790



Yogyakarta, 03 November 2025
UIN Sunan Kalijaga
Dekan Fakultas Sains dan Teknologi

Prof. Dr. Dra. Hj. Khairul Wardati, M.Si.
SIGNED

Valid ID: 693fb955db006

PERNYATAAN KEASLIAN TESIS

PERNYATAAN KEASLIAN TESIS/TUGAS AKHIR

Saya yang bertanda tangan dibawah ini :

Nama : Deny Setiawan
NIM : 23206051013
Program Studi : Magister Informatika
Fakultas : Sains dan Teknologi

Menyatakan bahwa tesis saya yang berjudul "Optimasi Model Natural Language Processing : IndoBERT dalam Klasifikasi Anamnesis Medis berdasarkan SNOMED-CT" merupakan hasil penelitian saya sendiri tidak terdapat pada karya yang pernah diajukan untuk memperoleh gelar magister di suatu perguruan tinggi, dan bukan plagiasi karya orang lain kecuali yang secara tertulis diacu dalam naskah ini dan disebutkan dalam daftar pustaka.

Yogyakarta, 16 Agustus 2025

Mengetahui,

METERAI
TEMPEL
E98AMX426381674
Deny Setiawan

NIM. 23206051013

STATE ISLAMIC UNIVERSITY
SUNAN KALIJAGA
YOGYAKARTA

SURAT PERSETUJUAN TESIS

SURAT PERSETUJUAN TESIS

Hal : Persetujuan Tesis

Lamp : -

Kepada

Yth. Dekan Fakultas Sains dan Teknologi

UIN Sunan Kalijaga Yogyakarta

di Yogyakarta

Assalamualaikum Wr. Wb.

Setelah membaca, meneliti, memberikan petunjuk dan mengoreksi serta mengadakan perbaikan seperlunya, maka kami selaku pembimbing berpendapat bahwa tesis saudara :

Nama : Deny Setiawan

NIM : 23206051013

Judul Tesis : Optimasi Model Natural Language Processing : IndoBERT dalam
Klasifikasi Anamnesis Medis berdasarkan SNOMED-CT

Sudah dapat diajukan kembali kepada Program Studi Magister Informatika Fakultas Sains dan Teknologi UIN Sunan Kalijaga Yogyakarta sebagai salah satu syarat untuk memperoleh gelar Magister Strata Dua dalam Program Studi Magister Informatika.

Dengan ini kami mengharap agar tesis/tugas akhir saudara dapat segera di-*munaqasyah*-kan. Atas perhatiannya kami ucapkan terima kasih.

Wassalamualaikum Wr. Wb.

Yogyakarta, 16 Agustus 2025

Pembimbing,



Dr. Agus Mulyanto, S.Si., M.Kom., ASEAN Eng.

NIP. 19710823 199903 1 003

KATA PENGANTAR

Puji syukur kami panjatkan kehadirat Allah SWT yang telah melimpahkan Nikmat, Taufiq dan Hidayahnya kepada penulis sehingga Tugas Akhir ini dapat selesai. Tugas Akhir ini disusun sebagai salah satu persyaratan untuk menyelesaikan Pendidikan Program studi Magister Informatika UIN Sunan Kalijaga Yogyakarta, oleh karena itu penulis mengucapkan terima kasih kepada

1. Bapak Prof. Noorhaidi, S.Ag., M.A., M.Phil., Ph.D. selaku Rektor UIN Sunan Kalijaga Yogyakarta
2. Bapak Dr. Dra. Hj. Khurul Wardati, M.Si., selaku Dekan Fakultas Sains dan Teknologi UIN Sunan Kalijaga.
3. Bapak Dr. Ir. Sumarsono, S.T., M.Kom. selaku Ketua Program Studi Magister Informatika UIN Sunan Kalijaga Yogyakarta sekaligus Dosen Pembimbing Akademik yang sudah banyak membantu penulis mulai dari awal penulis menempuh pendidikan magister hingga sekarang
4. Bapak Dr. Agus Mulyanto, S.Si., M.Kom., ASEAN Eng. selaku Dosen Pembimbing Tesis yang selalu mencurahkan waktu, nasehat, dan pikiran serta kemudahan dalam Bimbingan.
5. Bapak Dr. Muhammad Mustakim, S.T. M.T. selaku pemimpin CV Adhijasa Informatika yang telah memberikan bantuan beasiswa pendidikan serta dukungan moral dan materiil sehingga penulis dapat melanjutkan dan menyelesaikan studi dengan lancar.
6. .Kedua orang tua saya Bapak Sajari dan Ibu Sulastri yang selalu memberikan dukungan, motivasi, dan doa yang menyertai dari awal

perkuliahan hingga saat ini. Saya sangat berterima kasih sebanyak banyaknya atas dukungan dan doa dari orang tua saya.

7. Kakak saya Fandi Ahmad yang selalu memberikan semangat dan dukungannya, saya sangat berterima kasih kepada keluarga saya yang paling saya sayangi.
8. Bapak, Ibu, Kakak, Sahabat Magister Informatika Angkatan 2023 yang selalu memberikan semangat dan bantuan untuk segera menyelesaikan Tugas Akhir ini. Penulis menyadari penelitian ini belum sempurna.

Oleh karena itu, penulis mengharapkan kritik dan saran yang membangun dari berbagai pihak, semoga penelitian ini memberikan manfaat kepada pembaca dan menambah perkembangan ilmu pengetahuan yang ada.

Yogyakarta, 15 Juli 2025

Penulis,

Deny Setiawan

NIM: 23206051013

STATE ISLAMIC UNIVERSITY
SUNAN KALIJAGA
YOGYAKARTA

ABSTRAK

Klasifikasi entitas medis dari catatan anamnesis pasien penting untuk menghasilkan data klinis yang terstruktur dan sesuai dengan standar terminologi medis seperti SNOMED-CT. Namun, catatan medis dalam bentuk teks bebas sering kali mengandung variasi bahasa yang tinggi, sehingga diperlukan pendekatan berbasis *Natural Language Processing* (NLP). Penelitian ini bertujuan mengoptimalkan model IndoBERT melalui tiga tahap *fine-tuning* untuk mengklasifikasikan entitas medis pada catatan anamnesis pasien poli saraf. Dataset terdiri dari 500 catatan anamnesis yang telah dilabeli secara manual berdasarkan kategori medis seperti gejala, bagian tubuh, riwayat medis, dan waktu.

Proses *fine-tuning* dilakukan menggunakan *learning rate* 5e-5, *batch size* 4, dan *epoch* sebanyak tiga kali dengan variasi kombinasi *hyperparameter* pada setiap versi model (V1, V2, dan V3). Hasil evaluasi menunjukkan adanya peningkatan kinerja model seiring penyesuaian *hyperparameter*. Akurasi model meningkat dari 0.856 (V1) menjadi 0.870 (V3), sedangkan *weighted F1-score* naik dari 0.844 menjadi 0.850, menunjukkan peningkatan kestabilan dan ketepatan klasifikasi entitas medis. Meskipun nilai *macro F1-score* sedikit menurun dari 0.740 menjadi 0.731, hal tersebut disebabkan oleh variasi jumlah sampel antar kelas yang masih tidak seimbang.

Secara keseluruhan, hasil penelitian ini menunjukkan bahwa *fine-tuning* bertahap pada IndoBERT efektif dalam meningkatkan performa model untuk tugas *Named Entity Recognition* (NER) di bidang medis. Model terbaik (V3) menunjukkan kemampuan yang andal dalam mengenali entitas utama pada teks anamnesis medis berbahasa Indonesia, sehingga berpotensi diterapkan untuk mendukung standarisasi dan analisis data klinis di lingkungan pelayanan kesehatan.

Kata kunci: Anamnesis; Klasifikasi; IndoBERT; NLP; SNOMED-CT

ABSTRACT

Classifying medical entities from patient anamnesis records is essential for producing structured clinical data aligned with standardized medical terminologies such as SNOMED-CT. However, the free-text format of clinical notes often contains high linguistic variability, requiring a Natural Language Processing (NLP) approach. This study aims to optimize the IndoBERT model through three stages of fine-tuning to classify medical entities in neurological outpatient anamnesis records. The dataset consists of 500 manually labeled records categorized into medical entities such as symptoms, body parts, medical history, and time expressions.

The fine-tuning process was conducted using a learning rate of $5e-5$, a batch size of 4, and three training epochs, with varying hyperparameter combinations across three model versions (V1, V2, and V3). Evaluation results show progressive performance improvement as the hyperparameters were refined. The model's accuracy increased from 0.856 (V1) to 0.870 (V3), while the weighted F1-score improved from 0.844 to 0.850, indicating greater stability and precision in medical entity classification. Although the macro F1-score slightly decreased from 0.740 to 0.731, this was primarily due to class imbalance within the dataset.

Overall, the findings demonstrate that stepwise fine-tuning of IndoBERT effectively enhances model performance for medical Named Entity Recognition (NER) tasks. The best-performing model (V3) exhibits reliable capability in identifying key medical entities within Indonesian-language anamnesis texts, suggesting its strong potential for supporting clinical data standardization and analysis in healthcare information systems.

Keywords: *Anamnesis; Classification; IndoBERT; NLP; SNOMED-CT*

DAFTAR ISI

PENGESAHAN TUGAS AKHIR.....	i
PERNYATAAN KEASLIAN TESIS	ii
SURAT PERSETUJUAN TESIS	iii
KATA PENGANTAR	iv
ABSTRAK	vi
<i>ABSTRACT</i>	vii
DAFTAR ISI	viii
DAFTAR TABEL.....	xi
DAFTAR GAMBAR.....	xii
BAB I PENDAHULUAN	1
A. Latar Belakang.....	1
B. Rumusan Masalah	5
C. Batasan Masalah	5
D. Tujuan Penelitian.....	6
E. Manfaat Penelitian.....	6
F. Keaslian Penelitian	7
BAB II TINJAUAN PUSTAKA DAN LANDASAN TEORI	8
A. Tinjauan Pustaka.....	8
B. Landasan Teori.....	12
1. Kecerdasan Buatan	12
2. <i>Mechine Learning</i> dan <i>Deep Learning</i>	13
3. <i>Natural Language Processing</i>	14
4. <i>Large Language Models</i> (LLMs)	15
5. <i>Transformers</i> Model	16
6. <i>Bidirectional Encoder Representation From Transformers</i>	17
7. IndoBERT	19
8. <i>Fine Tuning</i>	19

9.	Evaluasi Model	21
10.	Catatan Medis	22
11.	SNOMED-CT	24
BAB III METODE PENELITIAN		27
A.	Alat dan Bahan	28
1.	<i>Hardware</i>	28
2.	<i>Software</i>	29
3.	Dataset	29
B.	Tahapan Penelitian	30
1.	Pemilihan Model <i>Transformer</i>	31
2.	<i>Pre-Processing</i>	33
3.	Pembagian <i>Dataset</i>	34
4.	Pemodelan	34
5.	<i>Fine Tuning</i>	34
6.	<i>Evaluasi Matrix</i>	35
7.	Pembahasan	35
BAB IV HASIL DAN PEMBAHASAN		37
A.	Pengumpulan <i>Dataset</i>	37
B.	<i>Pre-processing</i> Data	38
C.	Pembagian Dataset	39
1.	Penyiapan Dataset	39
2.	<i>Tokenization</i> Model	40
D.	Pemodelan	41
1.	Label dan Klasifikasi Token	41
2.	Model IndoBERT	42
E.	<i>Fine Tuning</i>	43
1.	<i>Data Collator</i> dan <i>Token Classification</i>	43
2.	Evaluasi Selama Pelatihan	43
3.	<i>Logging Loss</i>	44

4. Konfigurasi <i>Training</i>	45
5. Hasil <i>Fine-Tuning</i>	47
F. Evaluasi Model	55
G. Pembahasan	65
BAB V PENUTUP	69
A. Kesimpulan	69
B. Saran	70
DAFTAR PUSTAKA	71
CONTACT PERSON / CP	76



DAFTAR TABEL

Tabel 4. 1 Tabel Evaluasi Matric Model V1	52
Tabel 4. 2 Tabel Evaluasi Matric Model V2.....	53
Tabel 4. 3 Tabel Evaluasi Matric Model V3	54
Tabel 4. 4 Hasil Classification Report Model V1	57
Tabel 4. 5 Hasil Classification Report Model V2.....	58
Tabel 4. 6 Hasil Classification Report Model V3	60



DAFTAR GAMBAR

Gambar 2. 1 Venn diagram konsep machine learning dan deep learning (Christian Janiesch)	14
Gambar 2. 2 Prosedur pre-training dan fine-tuning	18
Gambar 4. 1 Hasil Pre-processing Dataset	39
Gambar 4. 2 Kode Pembagian Dataset	40
Gambar 4. 3 Kode Model IndoBERT	40
Gambar 4. 4 Kode Label dan Klasifikasi Token	42
Gambar 4. 5 Kode Struktur Pemodelan IndoBERT untuk Token Classification	42
Gambar 4. 6 Kode Data Collator dan Token Classification	43
Gambar 4. 7 Kode Matric Evaluasi Selama Pelatihan	44
Gambar 4. 8 Kode Logging Loss	45
Gambar 4. 9 Kode Konfigurasi Training V1	46
Gambar 4. 10 Kode Konfigurasi Training V2	47
Gambar 4. 11 Kode Konfigurasi Training V3	47
Gambar 4. 12 Kode Hasil Fine-Tuning V1	48
Gambar 4. 13 Kode Hasil Fine-Tuning V2	48
Gambar 4. 14 Kode Hasil Fine-Tuning V3	49
Gambar 4. 15 Training Loss per Step Model V1	49
Gambar 4. 16 Training Loss per Step Model V2	50
Gambar 4. 17 Training Loss per Step Model V3	51
Gambar 4. 18 Kode Evaluasi Model V1	56
Gambar 4. 19 Kode Evaluasi Model V2	56
Gambar 4. 20 Kode Evaluasi Model V3	56

Gambar 4. 21 Kode Classification Report Model V1	57
Gambar 4. 22 Kode Classification Report Model V2	58
Gambar 4. 23 Kode Classification Report Model V3	59
Gambar 4. 24 Kode Confusion Matric Model V1	61
Gambar 4. 25 Hasil Confusion Matrix Model V1	61
Gambar 4. 26 Kode Confusion Matric Model V2	62
Gambar 4. 27 Hasil Confusion Matrix Model V2	63
Gambar 4. 28 Kode Confusion Matric Model V3	64
Gambar 4. 29 Hasil Confusion Matrix Model V3	64
Gambar 4. 30 Perbandingan Hasil Kinerja Model V1, V2 dan V3	67



BAB I

PENDAHULUAN

A. Latar Belakang

Kemajuan teknologi digital di abad ke-21 telah menghadirkan perubahan signifikan pada berbagai aspek kehidupan, khususnya dalam sektor kesehatan. Pada era digital ini, penggunaan *Electronic Health Records* (EHR) telah menjadi salah satu pilar modernisasi layanan kesehatan, karena mampu menyimpan riwayat medis pasien secara elektronik dan mendukung pengambilan keputusan berbasis data. Namun, sebagian besar informasi medis dalam EHR masih berbentuk teks narasi bebas yang tidak terstruktur, seperti keluhan pasien, diagnosis, hingga catatan observasi dokter. Hal ini menimbulkan tantangan dalam pemrosesan data medis, terutama ketika data tersebut perlu dianalisis, dipertukarkan, atau dimanfaatkan untuk riset lanjutan. (Soltani & Tomberg 2019)

Ketidakseragaman kata, istilah lokal, singkatan, dan ejaan tidak baku dalam catatan medis menambah kompleksitas pemrosesan. Pada titik inilah *Artificial Intelligence* (AI), khususnya cabang *Natural Language Processing* (NLP), memainkan peran penting. NLP memungkinkan komputer mengekstrak informasi bermakna dari teks tidak terstruktur, mengidentifikasi entitas medis, memahami hubungan antar konsep, hingga melakukan normalisasi istilah medis ke dalam format standar. Seiring munculnya teknologi *transformer*, model bahasa seperti BERT dan variannya berhasil membawa lompatan performa NLP di banyak domain, termasuk domain Kesehatan (Devlin et al.

2019). Studi yang dilakukan (Abdulnazar et al. 2023) menunjukkan bahwa pemanfaatan arsitektur *transformer*, seperti SapBERT, mampu meningkatkan akurasi normalisasi istilah medis ke terminologi standar seperti SNOMED CT. (Abdulnazar et al. 2023)

Di sektor kesehatan global, normalisasi konsep medis menjadi hal mendasar untuk memastikan interoperabilitas data. Salah satu terminologi standar yang banyak digunakan adalah SNOMED CT (*Systematized Nomenclature of Medicine – Clinical Terms*). SNOMED CT menyediakan kerangka hierarki konsep medis yang komprehensif dengan jutaan istilah, sinonim, dan relasi antar konsep. Implementasi SNOMED CT memungkinkan data medis dapat dipertukarkan lintas rumah sakit, negara, hingga sistem kesehatan global, tanpa kehilangan makna klinis (Chang & Sung 2024). Namun, normalisasi teks medis ke SNOMED CT bukanlah tugas yang mudah. Teks klinis sering kali mengandung variasi frasa, singkatan, dan konteks lokal yang berbeda-beda, terutama di negara dengan bahasa non-Inggris.

Berbagai studi telah mencoba menjawab tantangan tersebut melalui pengembangan model NLP domain medis. MedCAT yang dikembangkan (Kraljevic et al. 2021) adalah salah satu *toolkit* ekstraksi informasi medis *multi-domain* yang dapat menghubungkan istilah teks bebas ke SNOMED CT dengan metode pembelajaran *self-supervised*. MedCAT bahkan mampu menangani entitas medis baru yang belum ada pada data pelatihan sebelumnya (Kraljevic et al. 2021). Sementara itu, (Soltani & Tomberg 2019) dengan Text2Node mengusulkan pendekatan yang menggabungkan *embedding* kata dengan *embedding node* dari taxonomi SNOMED CT, sehingga dapat melakukan pemetaan istilah ke

konsep meskipun frasa tersebut tidak muncul pada data latih. (Soltani & Tomberg 2019)

Studi lain oleh (Kalyan & Sangeetha 2020) memanfaatkan RoBERTa untuk melakukan *embedding* serentak antara frasa input dan konsep target, lalu mengukur *cosine similarity* untuk menentukan pemetaan yang paling mendekati. Hasilnya menunjukkan peningkatan akurasi signifikan pada normalisasi istilah medis di teks buatan pengguna (Kalyan & Sangeetha 2020). (Glen et al. 2024) juga menunjukkan bahwa pemetaan otomatis catatan medis ke SNOMED CT maupun ICD pada dataset MIMIC-III mampu mencapai cakupan kode hingga 97,98% dengan menggunakan pendekatan *explainable AI*, membuka peluang integrasi NLP dengan sistem rekam medis nyata. (Glen et al. 2024)

Di Indonesia, tantangan serupa bahkan lebih kompleks karena kebanyakan catatan medis ditulis dalam bahasa Indonesia. Bahasa Indonesia memiliki struktur, kosakata, dan variasi istilah medis lokal yang berbeda dengan bahasa Inggris. Sayangnya, hampir semua model NLP domain medis yang ada saat ini dilatih menggunakan korpus berbahasa Inggris. Akibatnya, adaptasi model NLP ke bahasa Indonesia memerlukan upaya *transfer knowledge* atau pelatihan ulang dengan data lokal. Di sinilah muncul peluang penggunaan IndoBERT, model BERT khusus bahasa Indonesia, yang dikembangkan dengan korpus berbahasa Indonesia agar lebih kontekstual. Meskipun IndoBERT sudah terbukti unggul dalam tugas NLP umum seperti *sentiment analysis* atau *named entity recognition*, penelitian yang menguji kemampuannya dalam domain catatan medis terutama untuk normalisasi ke SNOMED CT masih sangat terbatas.

Normalisasi istilah medis dalam teks berbahasa Indonesia memerlukan metode khusus agar dapat memetakan frasa naratif ke konsep standar secara akurat. (Abdulnazar et al. 2023) membuktikan bahwa pendekatan SapBERT, yang berfokus pada *fine-tuning embedding*, mampu meningkatkan F1-score pada pemetaan ke SNOMED CT dari 0,322 menjadi 0,853 pada data Bahasa Inggris (Abdulnazar et al. 2023). Ini menunjukkan potensi pendekatan serupa untuk diadaptasi ke IndoBERT. Selain itu, Text2Node juga relevan sebagai inspirasi karena berhasil memetakan frasa di domain lain ke taxonomi besar meski tanpa data paralel yang banyak (Soltani & Tomberg 2019)

Dengan adanya SNOMED CT sebagai acuan normalisasi, penelitian ini diharapkan dapat berkontribusi pada upaya standarisasi data kesehatan Indonesia. Penelitian serupa oleh (Chang & Sung 2024) menekankan pentingnya integrasi SNOMED CT ke dalam *pipeline* model bahasa besar (LLMs) agar sistem EHR dapat saling terhubung di berbagai tingkatan pelayanan Kesehatan (Chang & Sung 2024). Apabila IndoBERT berhasil dioptimasi untuk tugas normalisasi, maka Indonesia akan memiliki fondasi teknologi NLP lokal yang relevan, efisien, dan mendukung interoperabilitas data secara global.

Mengingat pentingnya standarisasi ini, penelitian “Optimasi *Large Language Model*: IndoBERT dalam Klasifikasi Anamnesia Berdasarkan SNOMED CT” diharapkan dapat menjawab tantangan praktis dan metodologis di atas. Hasilnya tidak hanya memperkaya kajian NLP berbahasa Indonesia, tetapi juga membuka jalan bagi sistem EHR nasional yang terhubung dengan standar internasional. Dengan demikian, proses ekstraksi informasi dari catatan medis tidak terstruktur

dapat dilakukan secara otomatis dan konsisten, mendukung keputusan klinis yang lebih cepat, akurat, dan berdampak pada mutu pelayanan kesehatan di Indonesia.

B. Rumusan Masalah

Penelitian ini berangkat dari permasalahan belum optimalnya pemanfaatan model NLP untuk memahami dan mengklasifikasi anamnesis berbahasa Indonesia sesuai dengan terminologi medis internasional seperti SNOMED-CT. Berdasarkan uraian tersebut, penelitian ini dimaksudkan untuk menjawab pertanyaan bagaimana performa model IndoBERT dalam melakukan optimasi pemahaman dan klasifikasi anamnesis berbahasa Indonesia berdasarkan terminologi SNOMED-CT, serta sejauh mana *fine-tuning* dapat meningkatkan akurasi dan konsistensi hasil tersebut?

C. Batasan Masalah

Dalam rangka memastikan penelitian berlangsung sesuai arah dan tujuan yang diharapkan, yaitu mengoptimalkan model IndoBERT untuk proses klasifikasi anamnesis berbahasa Indonesia berdasarkan terminologi SNOMED-CT, maka penelitian ini menetapkan sejumlah batasan, yaitu:

1. Pemrograman dilakukan menggunakan bahasa Python dengan *tools* Jupyter lab
2. Model yang digunakan adalah IndoBERT base v2 (indobenchmark/indobert-base-p2) sebagai model dasar yang akan dioptimasi melalui *fine-tuning*.
3. Dataset yang digunakan merupakan kumpulan anamnesis berbahasa Indonesia, khususnya bagian keluhan/subjektif pasien, yang telah

dipersiapkan dan dianotasi secara manual sesuai standar SNOMED-CT.

4. Proses evaluasi model dilakukan menggunakan metrik berbasis akurasi, *precision*, *recall*, dan F1-score.
5. Fokus penelitian terbatas pada proses klasifikasi anamnesa terminologi medis, tidak mencakup klasifikasi *diagnosis*, prediksi penyakit, ataupun integrasi sistem elektronik rekam medis lainnya.

D. Tujuan Penelitian

Dengan merujuk pada rumusan serta batasan masalah yang telah diuraikan sebelumnya, penelitian ini diarahkan untuk mengoptimalkan kinerja model IndoBERT dalam melakukan klasifikasi anamnesis berbahasa Indonesia berdasarkan standar terminologi SNOMED-CT melalui proses *fine-tuning*. Selain itu, penelitian ini juga bertujuan untuk mengevaluasi performa model dari aspek akurasi dan konsistensi hasil klasifikasi menggunakan metrik evaluasi seperti akurasi, *precision*, *recall*, dan F1-score.

E. Manfaat Penelitian

Melalui penelitian yang dilakukan, manfaat yang diharapkan adalah sebagai berikut:

1. Memberikan gambaran terkait penerapan model IndoBERT base v2 (indobenchmark/indobert-base-p2) dalam proses klasifikasi anamnesis berbahasa Indonesia berdasarkan terminologi SNOMED-CT.
2. Mengetahui sejauh mana proses *fine-tuning* dapat meningkatkan performa model IndoBERT dalam memahami dan menyelaraskan istilah medis dengan standar SNOMED-CT.

3. Memberikan acuan awal dalam penggunaan model NLP berbasis *transformer* untuk aplikasi klasifikasi terminologi medis di Indonesia.
4. Dapat dikembangkan lebih lanjut untuk mendukung integrasi sistem rekam medis elektronik (EHR) yang mengacu pada standar internasional seperti SNOMED-CT.

F. Keaslian Penelitian

Dalam peninjauan literatur dan observasi dari beberapa penelitian terdahulu, penelitian yang secara spesifik membahas optimasi model IndoBERT base v2 (indobenchmark/indobert-base-p2) untuk tugas klasifikasi anamnesis berbahasa Indonesia berdasarkan terminologi SNOMED-CT belum pernah dilakukan.

Sebagian besar penelitian sebelumnya lebih berfokus pada penerapan model NLP untuk ekstraksi informasi medis atau klasifikasi penyakit, namun belum secara eksplisit mengkaji proses klasifikasi istilah medis terhadap SNOMED-CT menggunakan model IndoBERT. Oleh karena itu, penelitian ini memiliki kontribusi orisinal dalam mengisi celah tersebut dan dapat menjadi rujukan awal dalam pengembangan sistem NLP berbasis standar terminologi medis di Indonesia

BAB V

PENUTUP

A. Kesimpulan

Berdasarkan hasil penelitian ini, dapat disimpulkan bahwa proses fine-tuning bertahap pada model IndoBERT berhasil meningkatkan kinerja dalam melakukan klasifikasi entitas medis berdasarkan standar SNOMED-CT. Hasil evaluasi menunjukkan bahwa akurasi model meningkat dari 0,856 pada versi pertama menjadi 0,870 pada versi ketiga, menandakan adanya perbaikan kemampuan model dalam mengenali pola linguistik pada teks anamnesis medis. Nilai weighted F1-score juga mengalami peningkatan dari 0,844 menjadi 0,850, yang menunjukkan peningkatan konsistensi performa model terhadap distribusi kelas yang tidak seimbang.

Meskipun nilai macro F1-score sedikit menurun dari 0,740 menjadi 0,731, hal ini masih dapat diterima mengingat macro F1 lebih sensitif terhadap variasi antar kelas. Penurunan kecil tersebut mengindikasikan bahwa masih terdapat beberapa entitas dengan jumlah sampel terbatas yang sulit dikenali secara optimal. Namun secara keseluruhan, tren peningkatan akurasi dan stabilitas prediksi menunjukkan bahwa setiap tahap fine-tuning memberikan dampak positif terhadap kemampuan generalisasi model.

Dengan demikian, dapat disimpulkan bahwa model IndoBERT hasil fine-tuning versi ketiga (V3) memberikan performa terbaik di antara semua versi yang diuji. Model ini tidak hanya menunjukkan peningkatan akurasi tertinggi, tetapi juga mempertahankan kestabilan performa pada metrik lainnya. Hasil ini menegaskan bahwa pendekatan

fine-tuning bertahap efektif dalam mengoptimalkan kemampuan model IndoBERT untuk ekstraksi entitas medis dari teks anamnesis pasien.

B. Saran

Untuk penelitian selanjutnya, disarankan menambah jumlah data anamnesis, khususnya pada entitas yang jarang muncul, agar model lebih seimbang dalam klasifikasi. Penerapan teknik augmentasi teks juga dapat meningkatkan performa model pada label yang sulit dikenali. Eksperimen dengan versi IndoBERT yang lebih besar atau model *transformer* lain dapat menjadi alternatif untuk meningkatkan akurasi. Selain itu, integrasi model ke dalam sistem rekam medis rumah sakit perlu disertai evaluasi klinis dan pemantauan berkala untuk memastikan konsistensi dan keamanan data, terutama bila terdapat perubahan pola data *anamnesis* dari waktu ke waktu.



DAFTAR PUSTAKA

- Abdulnazar, A., Kreuzthaler, M., Roller, R. & Schulz, S., 2023, *SapBERT-Based Medical Concept Normalization Using SNOMED CT*, *Studies in Health Technology and Informatics*, vol. 302, 825–826, IOS Press BV.
- Annisa, A.Y.N., 2025, 'ANALISIS PENERAPAN ELECTRONIC MEDICAL RECORDS (RME) DALAM MENINGKATKAN KEPUASAN PASIEN DAN MUTU PELAYANAN DI RUMAH SAKIT LITERATURE REVIEW', *JURNAL KESEHATAN TAMBUSAI*, 6(1), 3598–3607.
- Anugerah Simanjuntak, Rosni Lumbantoruan, Kartika Sianipar, Rut Gultom, Mario Simaremare, Samuel Situmeang & Erwin Panggabean, 2024, 'Research and Analysis of IndoBERT Hyperparameter Tuning in Fake News Detection', *Jurnal Nasional Teknik Elektro dan Teknologi Informasi*, 13(1), 60–67.
- Bahar, A.F., 2024, *MODEL OTOMATISASI SELEKSI RESUME LAMARAN KERJA MENGGUNAKAN NATURAL LANGUAGE PROCESSING (NLP) DAN MACHINE LEARNING (ML)* – PhD thesis, UIN SUNAN KALIJAGA .
- Berges, I, Bermúdez, J., Illarramendi, A. & Berges, Idoia, no date, *Binding SNOMED CT Terms to Archetype Elements: Establishing a Baseline of Results*.
- Cahyawijaya, S., Winata, G.I., Wilie, B., Vincentio, K., Li, X., Kuncoro, A., Ruder, S., Lim, Z.Y., Bahar, S., Khodra, M.L., Purwarianti, A. & Fung, P., 2021, 'IndoNLG: Benchmark and Resources for Evaluating Indonesian Natural Language Generation'.
- Castaño, J., Gambarte, L., Otero, C. & Luna, D., 2022, *A Simple Terminology-Based Approach to Clinical Entity Recognition*.
- Chang, E. & Sung, S., 2024, *Use of SNOMED CT in Large Language Models: Scoping Review*, *JMIR Medical Informatics*, 12.

- Chang, Yupeng, Wang, X., Wang, J., Wu, Y., Yang, L., Zhu, K., Chen, H., Yi, X., Wang, C., Wang, Y., Ye, W., Zhang, Y., Chang, Yi, Yu, P.S., Yang, Q. & Xie, X., 2024, 'A Survey on Evaluation of Large Language Models', *ACM Transactions on Intelligent Systems and Technology*, 15(3).
- DeepSeek-AI, :, Bi, X., Chen, D., Chen, G., Chen, S., Dai, D., Deng, C., Ding, H., Dong, K., Du, Q., Fu, Z., Gao, H., Gao, K., Gao, W., Ge, R., Guan, K., Guo, D., Guo, J., Hao, G., Hao, Z., He, Y., Hu, W., Huang, P., Li, E., Li, G., Li, J., Li, Y., Li, Y.K., Liang, W., Lin, F., Liu, A.X., Liu, B., Liu, W., Liu, Xiaodong, Liu, Xin, Liu, Y., Lu, H., Lu, S., Luo, F., Ma, S., Nie, X., Pei, T., Piao, Y., Qiu, J., Qu, H., Ren, T., Ren, Z., Ruan, C., Sha, Z., Shao, Z., Song, J., Su, X., Sun, J., Sun, Y., Tang, M., Wang, B., Wang, P., Wang, S., Wang, Yaohui, Wang, Yongji, Wu, T., Wu, Y., Xie, X., Xie, Zhenda, Xie, Ziwei, Xiong, Y., Xu, H., Xu, R.X., Xu, Y., Yang, D., You, Y., Yu, S., Yu, X., Zhang, B., Zhang, H., Zhang, Lecong, Zhang, Liyue, Zhang, Mingchuan, Zhang, Minghua, Zhang, W., Zhang, Y., Zhao, C., Zhao, Y., Zhou, Shangyan, Zhou, Shunfeng, Zhu, Q. & Zou, Y., 2024, 'DeepSeek LLM: Scaling Open-Source Language Models with Longtermism'.
- Devlin, J., Chang, M.-W., Lee, K. & Toutanova, K., 2019, 'BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding'.
- Ehrsam, J., Gaudet-Blavignac, C., Mattei, M., Baumann, M. & Lovis, C., 2025, 'Semantics in action: a guide for representing clinical data elements with SNOMED CT', *Journal of Biomedical Semantics*, 16(1).
- Glen, J., Han, L., Rayson, P. & Nenadic, G., 2024, 'A Comparative Study on Automatic Coding of Medical Letters with Explainability'.
- Hikmah, A., Azmi, F. & Nugrahaeni, R.A., 2023, 'Implementasi Natural Language Processing Pada Chatbot Untuk Layanan Akademik', *e-Proceeding of Engineering*, 10, 371.

- Islam, K.M.S., Nipu, A.S., Wu, J. & Madiraju, P., 2025, 'LLM-based Prompt Ensemble for Reliable Medical Entity Recognition from EHRs'.
- Izza, A. Al & Lailiyah, S., 2024, 'Kajian Literatur: Gambaran Implementasi Rekam Medis Elektronik di Rumah Sakit Indonesia berdasarkan Permenkes Nomor 24 Tahun 2022 tentang Rekam Medis', *Media Gizi Kesmas*, 13(1), 549–562.
- Kalyan, K.S. & Sangeetha, S., 2020, *Medical Concept Normalization in User-Generated Texts by Learning Target Concept Embeddings, Proceedings of the 11th International Workshop on Health Text Mining and Information Analysis*, 18–23, Association for Computational Linguistics (ACL).
- Kholili, ulil, 2011, 'Pengenalan Ilmu Rekam Medis Pada Masyarakat Serta Kewajiban Tenaga Kesehatan di Rumah Sakit Introduction to Medical Records In Community Health Workers And Liabilities at hospital', *Jurnal Kesehatan Komunitas*, 1(2), 61–72.
- Koto, F., Rahimi, A., Lau, J.H. & Baldwin, T., 2020, 'IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP'.
- Kraljevic, Z., Searle, T., Shek, A., Roguski, L., Noor, K., Bean, D., Mascio, A., Zhu, L., Folarin, A.A., Roberts, A., Bendayan, R., Richardson, M.P., Stewart, R., Shah, A.D., Wong, W.K., Ibrahim, Z., Teo, J.T. & Dobson, R.J., 2021, 'Multi-domain Clinical Natural Language Processing with MedCAT: the Medical Concept Annotation Toolkit'.
- Losch, N., Plagwitz, L., Büscher, A. & Varghese, J., 2025, 'Fine-Tuning LLMs on Small Medical Datasets: Text Classification and Normalization Effectiveness on Cardiology reports and Discharge records'.

- Muhammad, T., Rahardiansyah, R., Setya Perdana, R. & Fatyanosa, T.N., 2025, *Analisis Teknik Embedding Model NV-Embed pada Large Language Models Berbasis Retrieval Augmented Generation*, vol. 9.
- Nazi, Z. Al & Peng, W., 2024, *Large Language Models in Healthcare and Medical Domain: A Review, Informatics*, 11(3).
- Nuryadi, D., Metandi, F., Alam Hadiwijaya, N., Zainul Rohman, M., Hartanto, S., Syafrizal, A., Yadie Teknologi Informasi, E., Negeri Samarinda Jl Cipto Mangunkusumo, P., Seberang, S. & Timur, K., 2025, *FINE TUNING INDOBERT UNTUK ANALISIS SENTIMEN PADA ULASAN PENGGUNA APLIKASI TIKET.COM DI GOOGLE PLAY STORE*, vol. 9.
- Park, H.A., 2024, *Why Terminology Standards Matter for Data-driven Artificial Intelligence in Healthcare, Annals of Laboratory Medicine*, 44(6), 467–471.
- Perwira, R.I., Permadi, V.A., Purnamasari, D.I. & Agusdin, R.P., 2025, 'Domain-Specific Fine-Tuning of IndoBERT for Aspect-Based Sentiment Analysis in Indonesian Travel User-Generated Content', *Journal of Information Systems Engineering and Business Intelligence*, 11(1), 30–40.
- Prasetyo, V.R., Benarkah, N. & Chrisintha, V.J., 2021, 'Implementasi Natural Language Processing Dalam Pembuatan Chatbot Pada Program Information Technology Universitas Surabaya', *Teknika*, 10(2), 114–121.
- Sahit Reddy, K., Ragavenderan, N., Vasanth, K., Naik, G.N., Prabhu, V. & Nagaraja, G.S., 2025, 'MedicalBERT: enhancing biomedical natural language processing using pretrained BERT-based model', *IAES International Journal of Artificial Intelligence*, 14(3), 2367–2378.
- Soltani, R. & Tomberg, A., 2019, 'Text2Node: a Cross-Domain System for Mapping Arbitrary Phrases to a Taxonomy'.

- Vuokko, R., Vakkuri, A. & Palojoki, S., 2023, *Systematized Nomenclature of Medicine—Clinical Terminology (SNOMED CT) Clinical Use Cases in the Context of Electronic Health Record Systems: Systematic Literature Review*, *JMIR Medical Informatics*, 11.
- Wafda, A., 2025, *Aspect-Based Sentiment Analysis terhadap Cuitan Platform X tentang Kurikulum Merdeka Menggunakan IndoBERT* – PhD thesis, Universitas Islam Indonesia .
- Wildan Amru Hidayat & Nastiti, V.R.S., 2024, 'PERBANDINGAN KINERJA PRE-TRAINED INDOBERT-BASE DAN INDOBERT-LITE PADA KLASIFIKASI SENTIMEN ULASAN TIKTOK TOKOPEDIA SELLER CENTER DENGAN MODEL INDOBERT', *JSil (Jurnal Sistem Informasi)*, 11(2), 13–20.
- Zamakhshari, F., 2024, *OPTIMASI MODEL NATURAL LANGUAGE PROCESSING UNTUK APLIKASI QUESTION-ANSWER KESEHATAN MENTAL: PERBANDINGAN MODEL INDOBERT DAN M-BERT* – PhD thesis, UIN SUNAN KALIJAGA .