



PROTEKSI ISI PROPOSAL

Dilarang menyalin, menyimpan, memperbanyak sebagian atau seluruh isi proposal ini dalam bentuk apapun kecuali oleh pengusul dan pengelola administrasi pengabdian kepada masyarakat

PROPOSAL PENELITIAN 2023

Rencana Pelaksanaan Penelitian: tahun 2023 s.d. tahun 2024

1. JUDUL PENELITIAN

Model Pemahaman Bahasa Indonesia Berbasis Transformers

Bidang Fokus RIRN / Bidang Unggulan Perguruan Tinggi	Tema	Topik (jika ada)	Rumpun Bidang Ilmu
Teknologi Informasi dan Komunikasi	Pengembangan sistem berbasis Kecerdasan buatan	Pengembangan aplikasi sistem cerdas	Ilmu Komputer

Kategori (Kompetitif Nasional/ Desentralisasi/ Penugasan)	Skema Penelitian	Strata (Dasar/ Terapan/ Pengembangan)	SBK (Dasar, Terapan, Pengembangan)	Target Akhir TKT	Lama Penelitian (Tahun)
Penelitian Kompetitif Nasional	Penelitian Disertasi Doktor	Riset Dasar	SBK Riset Dasar	2	2

2. IDENTITAS PENGUSUL

Nama, Peran	Perguruan Tinggi/ Institusi	Program Studi/ Bagian	Bidang Tugas	ID Sinta
TAUFIK FUADI ABIDIN Ketua Pengusul	Universitas Syiah Kuala	Kecerdasan Buatan	Sebagai ketua pelaksana penelitian (promotor) bertugas mengkoordinir tahapan-tahapan penelitian disertasi doktor, mengkoordinir penyempurnaan model pemahaman bahasa Indonesia untuk masalah klasifikasi emotion dan analisis sentimen menggunakan model transformer dan neural network dan mengujinya, mengarahkan mahasiswa S3 untuk melakukan pelatihan dan menguji model yang dikembangkan, serta membantu proses publikasi pencapaian luaran setiap tahunnya.	59393
KAHLIL Ko-Promotor	Universitas Syiah Kuala	Kecerdasan Buatan	Sebagai anggota pelaksana penelitian (co-promotor) bertugas membantu tahapan-tahapan penelitian disertasi doktor, membantu proses penyempurnaan model pemahaman	6670703

Nama, Peran	Perguruan Tinggi/ Institusi	Program Studi/ Bagian	Bidang Tugas	ID Sinta
			bahasa Indonesia untuk masalah yang dikaji, mengarahkan mahasiswa S3 untuk melakukan proses pelatihan dan pengujian model yang dikembangkan, dan membantu proses publikasi.	
HAMMAM RIZA Ko-Promotor	Universitas Syiah Kuala	Teknik Elektro	Sebagai anggota pelaksana penelitian (co-promotor) bertugas membantu tahapan-tahapan penelitian disertasi doktor, membantu proses penyempurnaan model pemahaman bahasa Indonesia untuk masalah yang dikaji, mengarahkan mahasiswa S3 untuk melakukan proses pelatihan dan pengujian model yang dikembangkan dan memastikan adanya kebaruan, serta membantu proses publikasi dan pencapaian luaran.	-
Hendri Ahmadian Mahasiswa Bimbingan	Universitas Syiah Kuala	Matematika dan Aplikasi Sains	Sebagai anggota penelitian dan mahasiswa doktor bertugas melakukan tahapan-tahapan penelitian disertasi doktor, mengembangkan model pemahaman bahasa Indonesia yang lebih baik dari model yang ada untuk masalah yang dikaji, melakukan proses pelatihan dan pengujian model yang dikembangkan, membandingkan hasil model dengan IndoNLU, dan menulis artikel ilmiah.	-

3. MITRA KERJASAMA PENELITIAN (JIKA ADA)

Pelaksanaan penelitian dapat melibatkan mitra kerjasama yaitu mitra kerjasama dalam melaksanakan penelitian, mitra sebagai calon pengguna hasil penelitian, atau mitra investor

Mitra	Nama Mitra	Dana
-------	------------	------

4. LUARAN DAN TARGET CAPAIAN

Luaran Wajib

Tahun Luaran	Jenis Luaran	Status target capaian	Keterangan
1	Artikel di Jurnal	accepted/published	Applied Computational Intelligence

			and Soft Computing, https://www.hindawi.com/journals/acisc/
2	Artikel di Jurnal	accepted/published	Journal of King Saud University - Computer and Information Sciences, https://www.sciencedirect.com/journal/journal-of-king-saud-university-computer-and-information-sciences

5. ANGGARAN

Rencana Anggaran Biaya penelitian mengacu pada PMK dan buku Panduan Penelitian dan Pengabdian kepada Masyarakat yang berlaku.

Total RAB 2 Tahun Rp. 120.000.000,00

Tahun 1 Total Rp. 60.000.000,00

Jenis Pembelian	Komponen	Item	Satuan	Vol.	Biaya Satuan	Total
Bahan	Bahan Penelitian (Habis Pakai)	Alat Tulis Kantor (ATK)	Unit	1	2.150.000	2.150.000
Bahan	Bahan Penelitian (Habis Pakai)	SSD eksternal portable 1TB 520MB/s	Unit	1	1.525.000	1.525.000
Pengumpulan Data	HR Pembantu Peneliti	Honor pengumpul data kajian (data acquisition)	OJ	180	25.000	4.500.000
Sewa Peralatan	Peralatan penelitian	Penyewaan lisensi GPU Google Colab+ Pro	Unit	6	850.000	5.100.000
Analisis Data	HR Pengolah Data	Honor pengolah data penelitian, pelatihan dan pengujian model	P (penelitian)	1	5.000.000	5.000.000
Analisis Data	Biaya konsumsi rapat	Biaya konsumsi rapat	OH	36	50.000	1.800.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	HR Sekretariat/ Administrasi Peneliti	Honor pelaksana administrasi	OB	2	300.000	600.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya konsumsi rapat	Biaya konsumsi rapat persiapan laporan dan luaran	OH	12	50.000	600.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya registrasi seminar internasional	Paket	1	2.000.000	2.000.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya penginapan mengikuti seminar internasional	Paket	1	3.920.000	3.920.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya tiket perjalanan mengikuti seminar internasional	Paket	2	4.480.000	8.960.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya uang harian mengikuti seminar internasional	Paket	2	1.590.000	3.180.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Publikasi artikel di Jurnal Internasional	Biaya APC pada jurnal internasional bereputasi	Paket	1	13.350.000	13.350.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Publikasi artikel di Jurnal Internasional	Biaya proof read artikel jurnal internasional bereputasi	Paket	1	4.700.000	4.700.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya proof read artikel camera ready prosiding	Paket	1	2.615.000	2.615.000

Tahun 2 Total Rp. 60.000.000,00

Jenis Pembelanjaan	Komponen	Item	Satuan	Vol.	Biaya Satuan	Total
Bahan	ATK	Alat Tulis Kantor (ATK)	Paket	1	2.100.000	2.100.000
Pengumpulan Data	HR Pembantu Peneliti	Honor pengumpul data kajian (data acquisition tahun ke-2)	OJ	170	25.000	4.250.000
Sewa Peralatan	Peralatan penelitian	Penyewaan lisensi GPU Google Colab+ Pro	Unit	8	850.000	6.800.000
Analisis Data	HR Pengolah Data	Honor pengolah data penelitian, pelatihan dan pengujian model	P (penelitian)	1	5.000.000	5.000.000
Analisis Data	Biaya konsumsi rapat	Biaya konsumsi rapat	OH	36	50.000	1.800.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	HR Sekretariat/ Administrasi Peneliti	Honor pelaksana administrasi	OB	2	300.000	600.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya konsumsi rapat	Biaya konsumsi rapat	OH	12	50.000	600.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya registrasi seminar internasional	Paket	1	2.000.000	2.000.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya penginapan mengikuti seminar internasional	Paket	1	1.960.000	1.960.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya tiket perjalanan mengikuti seminar internasional	Paket	1	4.480.000	4.480.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya uang harian mengikuti seminar internasional	Paket	1	1.590.000	1.590.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Biaya seminar internasional	Biaya proof read artikel camera ready prosiding	Paket	1	2.820.000	2.820.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Publikasi artikel di Jurnal Internasional	Biaya APC publikasi jurnal internasional bereputasi	Paket	1	20.000.000	20.000.000
Pelaporan, Luaran Wajib, dan Luaran Tambahan	Publikasi artikel di Jurnal Internasional	Biaya proof read artikel jurnal internasional bereputasi	Paket	1	6.000.000	6.000.000



Isian Substansi Proposal **SKEMA PENELITIAN DASAR**

Petunjuk: Pengusul hanya diperkenankan mengisi di tempat yang telah disediakan sesuai dengan petunjuk pengisian dan tidak diperkenankan melakukan modifikasi template atau penghapusan di setiap bagian.

JUDUL

Tuliskan Judul Usulan

Model Pemahaman Bahasa Indonesia Berbasis Transformers

RINGKASAN

Ringkasan penelitian tidak lebih dari 300 kata yang berisi urgensi, tujuan, dan luaran yang ditargetkan.

Penelitian ini dilakukan untuk meningkatkan kinerja model klasifikasi emosi (*emotion classification*), analisis sentimen (*sentiment analysis*), keterkaitan tekstual (*textual entailment*), dan pengenalan nama entitas (*named entity recognition*) pada model *benchmark* IndoNLU [1] menjadi lebih baik. Penelitian ini menggunakan pendekatan memodifikasi model *benchmark* IndoNLU, kemudian melakukan *fine-tuning* untuk model *benchmark* modifikasi IndoNLU dengan dataset katalog Nusa NLP Bahasa Indonesia [2]. Selanjutnya melakukan penggabungan model (*hybrid*) modifikasi IndoNLU dengan CNN/BiLSTM dan model *Generative Pre-training Transformers* versi 3 (GPT-3) pada masalah NLP model klasifikasi emosi, analisis sentimen, dan pengenalan nama entitas. Sementara untuk keterkaitan tekstual pendekatan modifikasi *benchmark* IndoNLU dengan model *graph neural network* akan digunakan. Setiap model *hybrid* akan dievaluasi dan diuji. Penelitian tahun pertama diarahkan pada modifikasi *benchmark* IndoNLU dan upaya meningkatkan kinerja model NLP *benchmark* IndoNLU untuk permasalahan klasifikasi emosi dan analisis sentimen. Luaran wajib penelitian tahun pertama akan dipublikasikan pada jurnal internasional bereputasi *Applied Computational Intelligence and Soft Computing* (Q2) dan luaran tambahan adalah publikasi artikel prosiding pada *The Pacific Rim International Conference on Artificial Intelligence* (PRICAI) 2023. Tahun kedua, penelitian difokuskan pada peningkatan kinerja model NLP *benchmark* IndoNLU pada masalah keterkaitan tekstual dan pengenalan nama entitas. Luaran wajib penelitian tahun kedua berupa publikasi pada jurnal internasional bereputasi *Journal of King Saud University - Computer and Information Sciences* (Q1).

KATA KUNCI

Kata kunci maksimal 5 kata

Pemrosesan bahasa; transformers; IndoNLU; *deep learning*

PENDAHULUAN

Penelitian Dasar merupakan riset yang memuat temuan baru atau pengembangan ilmu pengetahuan dari kegiatan riset yang terdiri dari tahapan penentuan asumsi dan dasar hukum yang akan digunakan, formulasi konsep dan/ atau aplikasi formulasi dan pembuktian konsep fungsi dan/ atau karakteristik penting secara analitis dan eksperimental.

Pendahuluan penelitian tidak lebih dari 1000 kata yang terdiri dari:

A. Latar belakang dan rumusan permasalahan yang akan diteliti

- B. Pendekatan pemecahan masalah
- C. *State of the art* dan kebaruan
- D. Peta jalan (*road map*) penelitian 5 tahun kedepan (jika dalam bentuk konsorsium harus dilengkapi dengan roadmap penelitian konsorsium)
- E. Sitasi disusun dan ditulis berdasarkan sistem nomor sesuai dengan urutan pengutipan, mengikuti format Vancouver

A. Latar belakang dan rumusan permasalahan yang akan diteliti

Natural Language Processing (NLP) adalah salah satu bidang *Artificial Intelligence* (AI) yang bertujuan untuk mengekstrak pengetahuan dari data teks dan memprosesnya untuk beragam tugas yang memberikan kesimpulan penting tentang data teks yang dianalisis. Dalam beberapa tahun terakhir ini, ada banyak penelitian tentang berbagai metode penyisipan kata dan pemodelan pada *pretrained* bahasa yang telah dilatih sebelumnya sehingga masalah NLP dari data teks tujuan untuk mencapai kinerja optimal dari setiap masalah NLP [3].

Perkembangan *deep learning* yang begitu pesat, berbagai jenis jaringan saraf (*neural network*) telah banyak digunakan untuk menyelesaikan masalah NLP, seperti *convolutional neural networks* [4]–[6], *recurrent neural networks* [7], [8], *graph-based neural networks* [9]–[11] dan *attention mechanisms* [12], [13]. Saat ini, penelitian yang terkait dengan *pretrained model* menunjukkan bahwa *pretrained model* pada data besar dapat mempelajari representasi bahasa universal yang bermanfaat untuk jenis masalah *downstream* NLP. Selain itu, *pretrained model* dapat menghindari pelatihan model baru dari awal [14]. Penelitian [15], [16] tentang ELMo (*Embeddings from Language Models*) dan BERT (*Bidirectional Encoder Representations from Transformers*) merupakan riset penting untuk mengukur *pretrained model* kontekstual bahasa. Kemudian diikuti oleh model bahasa umum seperti *General Language Understanding Evaluation* - GLUE [17], SuperGLUE [18], dan *Chinese Language Understanding and Evaluation* - CLUE [19]. Namun, penelitian tersebut hanya terbatas pada bahasa umum yang memiliki data. Sebaliknya sebagian bahasa, seperti Bahasa Indonesia, mengalami keterbatasan data dan nilai akurasi modelnya masih rendah. Meskipun sejumlah data Bahasa Indonesia tersedia, kemajuan penelitian dalam bidang NLP dalam Bahasa Indonesia sudah mulai bertambah, namun masih berjalan lambat karena kurangnya sumber data beranotasi dan belum standardisasi sumber dataset dalam Bahasa Indonesia [3].

Pengembangan model arsitektur BERT menggunakan dataset Bahasa Indonesia dimulai pada model IndoNLU (*Indonesian Natural Language Understanding*) yang terdiri dari dua belas masalah (*tasks*) [20]. Penelitian ini menggunakan *hyperparameter* model BERT [21] dan model ALBERT [22]. Hasil yang diperoleh, masalah *aspect-based sentiment analysis* dan *part of speech* mendapat nilai evaluasi masing-masing 95,69% dan 95,71%. Sementara untuk klasifikasi emosi (*emotion classification*), analisis sentimen (*sentiment analysis*), keterkaitan tekstual (*textual entailment*), dan pengenalan nama entitas (*named entity recognition*), kinerjanya perlu ditingkatkan.

Selanjutnya, penelitian [3] memperkenalkan data *benchmark* yang disebut IndoLEM (*Indonesian Language Evaluation Montage*), untuk digunakan untuk mengevaluasi masalah NLP *part-of-speech*, *named entity recognition* (NER), *dependency parsing*, *sentiment analysis*, peringkasan (*summarization*), prediksi tweet (*next tweet prediction*), dan pengaturan tweet. Untuk masalah *part-of-speech* dan *next tweet prediction*, akurasi dari masing-masing model adalah 96,8% dan 93,7%. Namun untuk NER, *dependency parsing*, *sentiment analysis*, dan *summarization*, kinerja model masih di bawah 85%.

Beberapa model NLP memiliki peranan penting dalam kajian praktis seperti klasifikasi emosi yang dapat digunakan untuk deteksi sarkasme [23] dan ekspresi emosi dalam dokumen [24], [25]. Selanjutnya, analisis sentimen metode *deep learning* [26] mulai diterapkan menggunakan dataset selain Bahasa Inggris [27]. Sementara model keterkaitan tekstual digunakan untuk menentukan kemiripan kalimat [28], [29], sementara masalah pengenalan nama entitas dapat diterapkan pada sosial media [30], [31], dan bidang kedokteran [32], [33].

Rumusan masalah dari penelitian yang diajukan ini adalah bagaimana meningkatkan kinerja model klasifikasi emosi, analisis sentimen, keterkaitan tekstual, dan pengenalan nama entitas berbasis *transformers* yang diuji menggunakan data *benchmark* IndoNLU karena kinerja model yang ada saat ini masih rendah, masing-masing sebesar 79,47%, 92,03%, 85,41%, dan 79,25% sehingga masih terbuka peluang untuk ditingkatkan kinerjanya.

B. Pendekatan pemecahan masalah

Beberapa pendekatan untuk mendapatkan hasil peningkatan kinerja pada model masalah NLP, yaitu pertama melakukan *fine-tuning* untuk model *benchmark* modifikasi IndoNLU. Selanjutnya, model modifikasi IndoNLU digabungkan dengan CNN/BiLSTM dan model GPT-3 [34] untuk masalah klasifikasi emosi [35], analisis sentimen, dan pengenalan nama entitas [36]. Khusus masalah keterkaitan tekstual, akan dilakukan menggunakan pendekatan *graph neural network* [37].

C. *State of the art* dan kebaruan

State of the art dan kebaruan dari penelitian ini dirangkum pada Tabel 1. Tabel merangkum penelitian yang telah dilakukan, urgensi, hasil yang diperoleh, dan celah (gap) penelitian yang masih dapat dikaji.

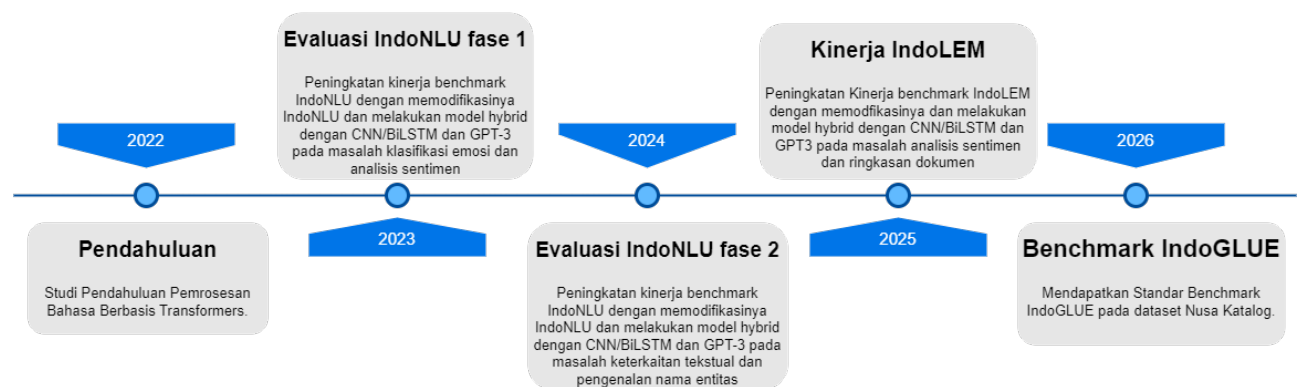
No	Penelitian	Urgensi	Hasil	Celah Penelitian Lanjutan
1	Lintang S, 2020 [38]	Riset tentang model pemahaman bahasa Indonesia, <i>pretrained</i> pemodelan Bahasa Indonesia pada masalah NLP model analisis sentiment dan <i>summarization</i> dan menggunakan dataset korpus Oscar. Selanjutnya hasil yang didapat dibandingkan dengan model <i>Multilingual BERT</i> (M-BERT)	Hasil evaluasi jenis masalah NLP pada model analisis sentimen dan <i>summarization</i> mengungguli M-BERT.	Memiliki potensi untuk melakukan penelitian lanjutan pada masalah NLP lainnya dengan dataset yang sumbernya lebih beragam.

No	Penelitian	Urgensi	Hasil	Celah Penelitian Lanjutan
2	Wilie et al., 2020 [1]	Riset tentang model pemahaman Bahasa Indonesia dengan melakukan <i>training</i> , evaluasi (<i>evaluation</i>), dan <i>benchmarking</i> yang dinamai (IndoNLU) berbasis <i>Transformers</i> pada 12 masalah NLP Bahasa Indonesia.	Pada masalah <i>aspect-based sentiment analysis</i> dan <i>part of speech</i> mendapat nilai evaluasi masing-masing 95,69% dan 95,71%. Namun untuk masalah NLP model klasifikasi emosi, analisis sentimen, keterkaitan tekstual dan pengenalan nama entitas masih rendah untuk skor hasil evaluasinya.	Pentingnya dilakukan penelitian lanjutan untuk meningkatkan kinerja pada masalah NLP model klasifikasi emosi, analisis sentiment, keterkaitan tekstual dan pengenalan nama entitas menjadi lebih baik.
3	Koto et al., 2020 [39]	Riset tentang model pemahaman Bahasa Indonesia dengan memperkenalkan dataset INDOLEM dan melakukan <i>pretrained</i> pemodelan bahasa (<i>language model</i>) berbasis <i>Transformers</i> pada dataset IndoLEM dan mengevaluasi pada 7 masalah NLP Bahasa Indonesia.	Hasil evaluasi untuk masalah NLP model <i>part-of-speech</i> dan <i>next tweet prediction</i> masing-masing 96,8% dan 93,7%. Namun untuk masalah NLP model NER, <i>dependency parsing</i> , <i>sentiment analysis</i> dan <i>summarization</i> kinerja model masih dibawah 85%.	Memiliki potensi untuk dilakukan pengkajian lanjutan untuk meningkatkan kinerja masalah NLP model NER, <i>dependency parsing</i> , <i>sentiment analysis</i> dan <i>summarization</i>
4	Usulan Penelitian 2023	Urgensi: Kinerja model pada <i>benchmark</i> IndoNLU masih terbuka peluang untuk ditingkatkan dengan menggunakan modifikasi model <i>benchmark</i> IndoNLU dan menggabungkannya dengan model yang lain (<i>hybrid</i>)		

Selain itu, kebaruan pada penelitian ini adalah peningkatan kinerja atau akurasi model untuk masalah NLP klasifikasi emosi, analisis sentimen, keterkaitan tekstual, dan model pengenalan nama entitas berbasis *transformers*. Kedua, penggabungan model modifikasi *benchmark* IndoNLU dengan beberapa model *deep learning* berbasis *neural network* untuk meningkatkan kinerja dan akurasi model.

B. Peta jalan (*road map*) penelitian 5 tahun kedepan

Penelitian ini dimulai studi pendahuluan model pemrosesan Bahasa Indonesia berbasis *transformers*. Selanjutnya, *fine-tuning benchmark* IndoNLU untuk masalah NLP klasifikasi emosi dan analisis sentimen menggunakan pendekatan metode penggabungan model CNN/BiLSTM dan GPT-3. Kemudian, masalah NLP pengenalan nama entitas dilakukan dengan pendekatan penggabungan model CNN/BiLSTM dan GPT-3, sementara untuk masalah keterkaitan tekstual, dilakukan melalui pendekatan *graph neural network*. Pada tahun kedua, evaluasi kinerja IndoLEM dilakukan untuk masalah NLP analisis sentimen dan peringkasan dokumen. Terakhir, peta jalan penelitian difokuskan pada upaya untuk mendapatkan standar *benchmark* IndoGLUE. Gambar 1 memperlihatkan peta jalan (*road map*) penelitian dalam beberapa tahun ke depan.



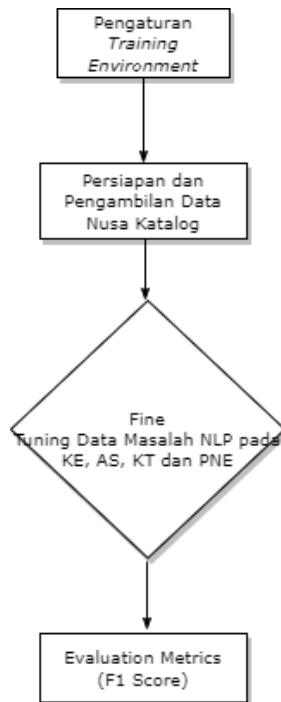
Gambar 1. Peta jalan penelitian

METODA

Metode atau cara untuk mencapai tujuan yang telah ditetapkan ditulis tidak melebihi 1000 kata. Bagian ini dapat dilengkapi dengan diagram alir penelitian yang menggambarkan apa yang sudah dilaksanakan dan yang akan dikerjakan selama waktu yang diusulkan. Format diagram alir dapat berupa file JPG/PNG. Metode penelitian harus dibuat secara utuh dengan penahapan yang jelas, mulai dari awal bagaimana proses dan luarannya, dan indikator capaian yang ditargetkan yang tercermin dalam Rencana Anggaran Biaya (RAB).

Prosedur kerja pada penelitian ini dijelaskan sebagai berikut.

- Modifikasi *benchmark* model IndoNLU dan *fine-tuning* untuk masalah NLP



Gambar 2. Alur kerja model modifikasi untuk *fine-tuning* IndoNLU

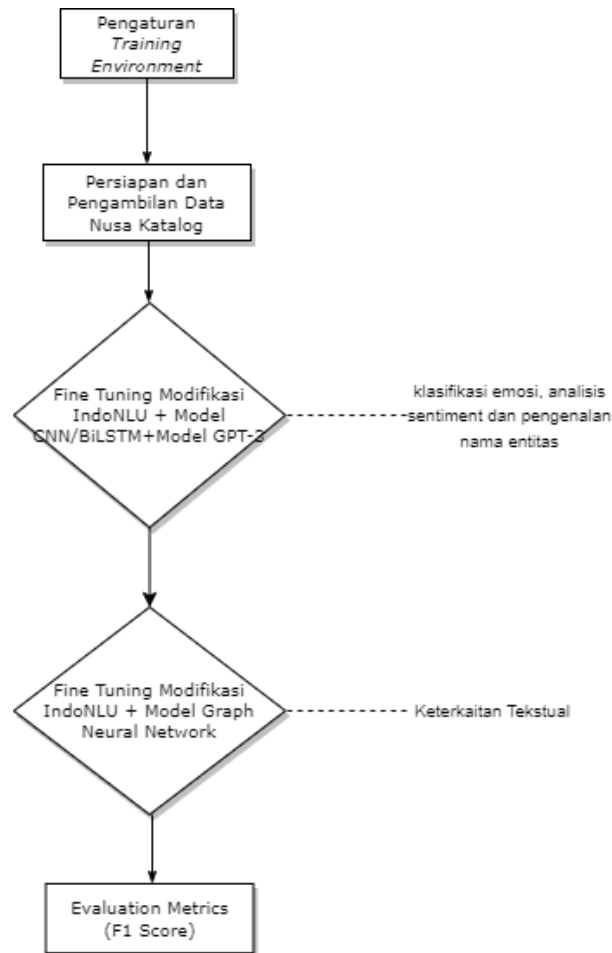
Alur kerja model modifikasi untuk *fine-tuning benchmark* IndoNLU dimulai dari pengaturan *training environment*. Selanjutnya, dilakukan pengambilan data untuk jenis masalah pada dataset Nusa Katalog [2]. Pada tahap ini, dilakukan *pre-processing* dataset sehingga dapat digunakan pada model *benchmark* IndoNLU. Kemudian, dilakukan *fine-tuning* untuk masalah NLP model klasifikasi emosi, analisis sentimen, keterkaitan tekstual, dan pengenalan nama entitas pada setiap dataset yang sudah di *pre-processing*. Evaluasi akan dilakukan menggunakan *F1 score*. Alur kerja ini dapat dilihat pada Gambar 2.

Luaran pada tahap ini adalah model modifikasi *benchmark* IndoNLU. Indikator capaian target adalah peningkatan nilai klasifikasi emosi, analisis sentimen, keterkaitan tekstual, dan pengenalan nama entitas menjadi lebih baik. Saat ini, kinerja dari masing-masing model tersebut adalah 79,47%, 92,03%, 85,41%, dan 79,25%.

Setiap langkah yang sudah dijelaskan sebelumnya, dilakukan pengujian modifikasi *benchmark* IndoNLU pada *Google Collaborative Environment* (Colab) dengan menjalankan notebook pada *Google Cloud Virtual Machines* (VMs) dengan *hardware accelerators* menggunakan *Graphics Processing Unit* (GPU) dengan jumlah akses 100 *compute unit*.

- Model gabungan modifikasi IndoNLU dengan CNN/BiLSTM dan GPT-3

Pada tahap kedua ini, model modifikasi IndoNLU digabungkan dengan model CNN/BiLSTM dan model GPT-3 [34] untuk masalah klasifikasi emosi [35], analisis sentimen, dan pengenalan nama entitas [36]. Pelatihan dan pengujian menggunakan dataset katalog Nusa NLP Bahasa Indonesia [2]. Khusus untuk masalah keterkaitan tekstual pendekatan akan dilakukan menggunakan model *graph neural network* [37]. Proses alur kerja model *hybrid* dapat dilihat pada Gambar 3.



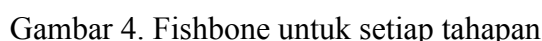
Gambar 3 Alur Kerja Model Hybrid

Pada model penggabungan model (*hybrid*), proses yang dilakukan adalah proses *transfer learning* pada masalah NLP setiap model dengan menggunakan *pretrained benchmark* IndoNLU dan modifikasi *fine-tuning benchmark* IndoNLU. Semua output token dari model *pretrained benchmark* akan digunakan sebagai masukan (*input*) untuk model CNN/BiLSTM dan model GPT-3.

Alur kerja penggabungan model (*hybrid*) diawali oleh pengaturan *training environment*. Selanjutnya dilakukan pengambilan data untuk jenis masalah pada dataset Nusa Katalog [2]. Kemudian, dilakukan *fine-tuning* sebagai bagian proses *transfer learning* dengan menggabungkan model benchmark IndoNLU untuk masalah NLP model klasifikasi emosi dan analisis sentimen dengan model CNN/BiLSTM serta model GPT-3. Pada masalah NLP pengenalan nama entitas, pendekatan akan menggunakan model modifikasi *benchmark* IndoNLU dan model CNN/BiLSTM.

Pada masalah keterkaitan tekstual, diawali dengan melakukan pembuatan *encoder* untuk mengekstrak penyematan tekstual, menggunakan *benchmark* IndoNLU, dan keterkaitan tekstual *encoder* untuk mengekstrak *graph embedding* dari model *graph neural network*. Setiap *encoder* dilakukan *pretrained* dengan data kontekstual, dan *graph* keterkaitan tekstual yang diekstraksi dari dataset. Kemudian, data dimasukkan ke dalam model *pretrained* dan proses penyematan gabungan (*concatenated embedding*) dihitung oleh masing-masing *encoder*. Setiap masalah NLP pada model *hybrid* ini akan dihitung kinerjanya menggunakan *F1 score*.

Setiap tahapan akan dilakukan pengujian pada *Google Colab* dengan menjalankan notebook pada *Google Cloud Virtual Machines* (VMs) dengan *hardware accelerators* menggunakan *Graphics Processing Unit* (GPU) dengan jumlah akses lebih kurang 100 *compute unit*.



Gambar 4 menggambarkan pendekatan yang akan dilakukan untuk mendapatkan kinerja model *benchmark* IndoNLU yang lebih baik. Proses tersebut dilakukan dalam waktu 2 tahun. Pada tahun pertama, kegiatan diawali dengan memodifikasi *benchmark* IndoNLU dan menggabungkan model CNN/BiLSTM dan model GPT-3 untuk masalah NLP klasifikasi emosi dan analisis sentimen. Selanjutnya, pada tahun kedua, penggabungan model *graph neural network* dilakukan untuk masalah NLP keterkaitan tekstual, sedangkan untuk masalah NLP pengenalan nama entitas penggabungan model CNN/BiLSTM dan GPT-3 dilakukan.

Jadwal penelitian disusun berdasarkan pelaksanaan penelitian, harap disesuaikan berdasarkan lama tahun pelaksanaan penelitian

[illegible]

	IndoNLU												
3	Memodifikasi model arsitektur <i>benchmark</i> IndoNLU												
4	Penulisan artikel prosiding dan pengiriman ke konferensi internasional												
5	Persiapan data untuk pengujian <i>fine-tuning</i> masalah NLP hasil modifikasi model												
6	Pengujian <i>fine-tuning</i> menggunakan modifikasi model <i>benchmark</i> IndoNLU dan model CNN/BiLSTM dan GPT-3 untuk masalah klasifikasi emosi												
7	Pengujian <i>fine-tuning</i> menggunakan modifikasi model gabungan <i>benchmark</i> IndoNLU dan model CNN/BiLSTM dan GPT-3 untuk masalah analisis sentimen												
8	Penulisan artikel jurnal dan mengirimkannya ke jurnal bereputasi internasional												

Tahun ke-2

No	Nama Kegiatan	Bulan											
		1	2	3	4	5	6	7	8	9	10	11	12
1	Persiapan data untuk pengujian <i>fine-tuning</i> masalah NLP hasil modifikasi model												
2	Pengujian <i>fine-tuning</i> menggunakan modifikasi model <i>benchmark</i> IndoNLU dan model CNN/BiLSTM dan GPT-3 untuk masalah pengenalan nama entitas												
3	Pengujian <i>fine-tuning</i> menggunakan modifikasi model <i>benchmark</i> IndoNLU dan model <i>Graph Neural Network</i> untuk masalah keterkaitan tekstual												
4	Penulisan artikel jurnal yang kedua dan <i>submit</i> ke jurnal internasional bereputasi												

DAFTAR PUSTAKA

Sitasi disusun dan ditulis berdasarkan sistem nomor sesuai dengan urutan pengutipan, mengikuti format Vancouver. Hanya pustaka yang disitasi pada usulan penelitian yang dicantumkan dalam Daftar Pustaka.

- [1] B. Wilie *et al.*, "IndoNLU: Benchmark and Resources for Evaluating Indonesian Natural Language Understanding," Sep. 2020, [Online]. Available: <http://arxiv.org/abs/2009.05387>
- [2] S. Cahyawijaya *et al.*, "NusaCrowd: Open Source Initiative for Indonesian NLP Resources." 2022.
- [3] F. Koto, A. Rahimi, J. H. Lau, and T. Baldwin, "IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP," Nov. 2020, [Online]. Available: <http://arxiv.org/abs/2011.00677>

- [4] J. Gehring, M. Auli, D. Grangier, D. Yarats, and Y. N. Dauphin, "Convolutional Sequence to Sequence Learning," 2017.
- [5] N. Kalchbrenner, E. Grefenstette, and P. Blunsom, "A Convolutional Neural Network for Modelling Sentences," 2014. [Online]. Available: www.nal.co
- [6] Y. Kim, "Convolutional Neural Networks for Sentence Classification," 2014. [Online]. Available: <http://nlp.stanford.edu/sentiment/>
- [7] P. Liu, X. Qiu, and X. Huang, "Recurrent Neural Network for Text Classification with Multi-Task Learning," 2016.
- [8] I. Sutskever Google, O. Vinyals Google, and Q. v le Google, "Sequence to Sequence Learning with Neural Networks," 2014.
- [9] D. Marcheggiani, J. Bastings, and I. Titov, "Exploiting Semantics in Neural Machine Translation with Graph Convolutional Networks," Apr. 2018, [Online]. Available: <http://arxiv.org/abs/1804.08313>
- [10] R. Socher *et al.*, "Recursive Deep Models for Semantic Compositionality Over a Sentiment Treebank," 2013. [Online]. Available: <http://nlp.stanford.edu/>
- [11] K. S. Tai, R. Socher, and C. D. Manning, "Improved Semantic Representations From Tree-Structured Long Short-Term Memory Networks," Feb. 2015, [Online]. Available: <http://arxiv.org/abs/1503.00075>
- [12] D. Bahdanau, K. Cho, and Y. Bengio, "Neural Machine Translation by Jointly Learning to Align and Translate," Sep. 2014, [Online]. Available: <http://arxiv.org/abs/1409.0473>
- [13] A. Vaswani *et al.*, "Attention Is All You Need," Jun. 2017, [Online]. Available: <http://arxiv.org/abs/1706.03762>
- [14] X. Han *et al.*, "Pre-Trained Models: Past, Present and Future," Jun. 2021, [Online]. Available: <http://arxiv.org/abs/2106.07139>
- [15] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," Oct. 2018, [Online]. Available: <http://arxiv.org/abs/1810.04805>
- [16] M. E. Peters *et al.*, "Deep contextualized word representations," Feb. 2018, [Online]. Available: <http://arxiv.org/abs/1802.05365>
- [17] A. Wang, A. Singh, J. Michael, F. Hill, O. Levy, and S. R. Bowman, "GLUE: A Multi-Task Benchmark and Analysis Platform for Natural Language Understanding," Apr. 2018, [Online]. Available: <http://arxiv.org/abs/1804.07461>
- [18] A. Wang *et al.*, "SuperGLUE: A Stickier Benchmark for General-Purpose Language Understanding Systems," May 2019, [Online]. Available: <http://arxiv.org/abs/1905.00537>
- [19] L. Xu *et al.*, "CLUE: A Chinese Language Understanding Evaluation Benchmark," Apr. 2020, [Online]. Available: <http://arxiv.org/abs/2004.05986>
- [20] B. Wilie *et al.*, "IndoNLU: Benchmark and Resources for Evaluating Indonesian Natural Language Understanding," Sep. 2020, [Online]. Available: <http://arxiv.org/abs/2009.05387>
- [21] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," Oct. 2018, [Online]. Available: <http://arxiv.org/abs/1810.04805>
- [22] Z. Lan *et al.*, "ALBERT: A LITE BERT FOR SELF-SUPERVISED LEARNING OF LANGUAGE REPRESENTATIONS." [Online]. Available: <https://github.com/google-research/ALBERT>.
- [23] A. Abuzayed and H. Al-Khalifa, "Sarcasm and sentiment detection in Arabic tweets using BERT-based models and data augmentation," in *Proceedings of the sixth Arabic natural language processing workshop*, 2021, pp. 312–317.
- [24] Z. Ding, R. Xia, and J. Yu, "ECPE-2D: Emotion-Cause Pair Extraction based on Joint Two-Dimensional Representation, Interaction and Prediction," in *Proceedings of the*

58th Annual Meeting of the Association for Computational Linguistics, Online: Association for Computational Linguistics, Jul. 2020, pp. 3161–3170. doi: 10.18653/v1/2020.acl-main.288.

- [25] R. Xia, M. Zhang, and Z. Ding, “RTHN: A RNN-Transformer Hierarchical Network for Emotion Cause Extraction,” in *Proceedings of the Twenty-Eighth International Joint Conference on Artificial Intelligence, IJCAI-19*, International Joint Conferences on Artificial Intelligence Organization, Mar. 2019, pp. 5285–5291. doi: 10.24963/ijcai.2019/734.
- [26] M. I. Prabha and G. Umarani Srikanth, “Survey of Sentiment Analysis Using Deep Learning Techniques,” in *2019 1st International Conference on Innovations in Information and Communication Technology (ICIICT)*, 2019, pp. 1–9. doi: 10.1109/ICIICT1.2019.8741438.
- [27] G. Alexandridis, I. Varlamis, K. Korovesis, G. Caridakis, and P. Tsantilas, “A Survey on Sentiment Analysis and Opinion Mining in Greek Social Media,” *Information*, vol. 12, no. 8, 2021, doi: 10.3390/info12080331.
- [28] N. Reimers and I. Gurevych, “Sentence-bert: Sentence embeddings using siamese bert-networks,” *arXiv preprint arXiv:1908.10084*, 2019.
- [29] X. Sun *et al.*, “Sentence Similarity Based on Contexts,” *Trans Assoc Comput Linguist*, vol. 10, pp. 573–588, May 2022, doi: 10.1162/tacl_a_00477.
- [30] V. Yadav and S. Bethard, “A survey on recent advances in named entity recognition from deep learning models,” *arXiv preprint arXiv:1910.11470*, 2019.
- [31] J. Li, A. Sun, J. Han, and C. Li, “A survey on deep learning for named entity recognition,” *IEEE Trans Knowl Data Eng*, vol. 34, no. 1, pp. 50–70, 2020.
- [32] M.-S. Huang, P.-T. Lai, P.-Y. Lin, Y.-T. You, R. T.-H. Tsai, and W.-L. Hsu, “Biomedical named entity recognition and linking datasets: survey and our recent development,” *Brief Bioinform*, vol. 21, no. 6, pp. 2219–2238, 2020.
- [33] L. Akhtyamova, “Named entity recognition in Spanish biomedical literature: Short review and BERT model,” in *2020 26th Conference of Open Innovations Association (FRUCT)*, IEEE, 2020, pp. 1–7.
- [34] T. Brown *et al.*, “Language models are few-shot learners,” *Adv Neural Inf Process Syst*, vol. 33, pp. 1877–1901, 2020.
- [35] M. A. H. Wadud, M. Mridha, J. Shin, K. Nur, and A. K. Saha, “Deep-BERT: Transfer Learning for Classifying Multilingual Offensive Texts on Social Media,” *Comput. Syst. Sci. Eng*, vol. 44, pp. 1775–1791, 2022.
- [36] Y. Chen, J. Ding, D. Li, and Z. Chen, “Joint BERT Model Based Cybersecurity Named Entity Recognition,” in *2021 The 4th International Conference on Software Engineering and Information Management*, in ICSIM 2021. New York, NY, USA: Association for Computing Machinery, 2021, pp. 236–242. doi: 10.1145/3451471.3451508.
- [37] C. Jeong, S. Jang, E. Park, and S. Choi, “A context-aware citation recommendation model with BERT and graph convolutional networks,” *Scientometrics*, vol. 124, pp. 1907–1922, 2020.
- [38] S. Lintang, “IndoBERT: Transformer-based Model for Indonesian Language,” Universitas Gajah Mada, Yogyakarta, 2020. Accessed: Apr. 11, 2023. [Online]. Available: <http://etd.repository.ugm.ac.id/penelitian/detail/190630>.
- [39] F. Koto, A. Rahimi, J. H. Lau, and T. Baldwin, “IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP,” Nov. 2020, [Online]. Available: <http://arxiv.org/abs/2011.00677>.

PERSETUJUAN PENGUSUL

Tanggal Pengiriman	Tanggal Persetujuan	Nama Pimpinan Pemberi Persetujuan	Sebutan Jabatan Unit	Nama Unit Lembaga Pengusul
-	-	-	-	-

Komentar : -