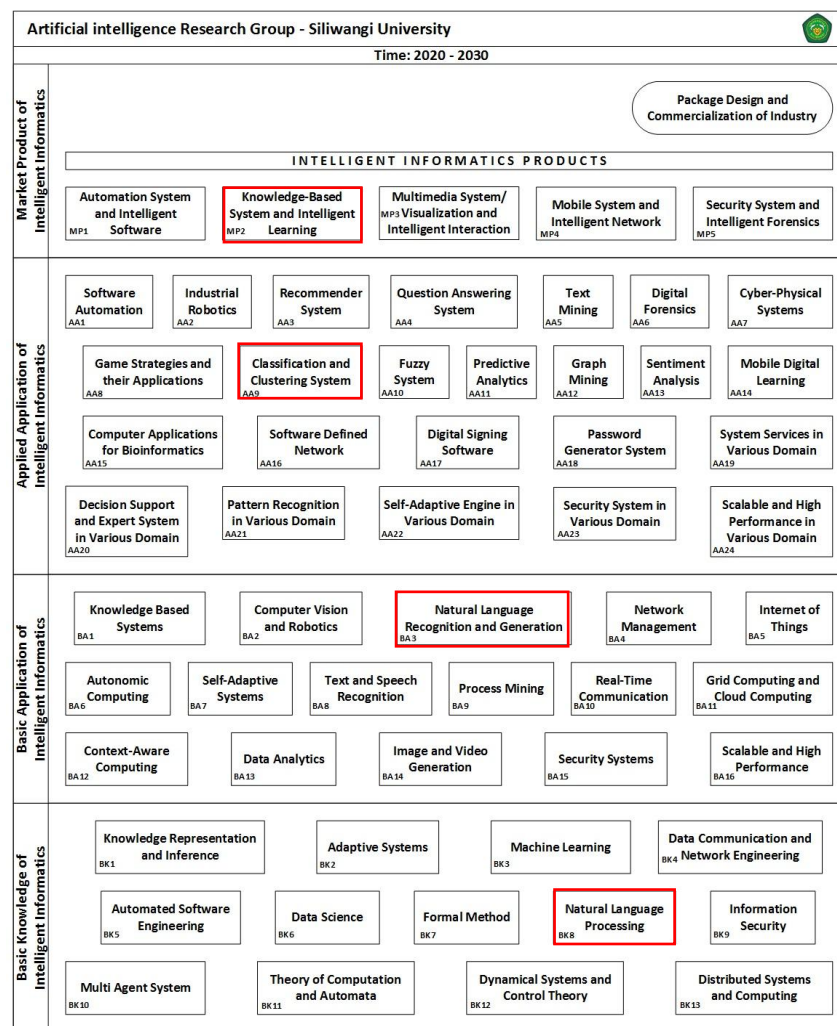


BAB III

METODOLOGI

3.1 Peta Jalan (*Roadmap*) Penelitian

Berdasarkan *roadmap* penelitian Universitas Siliwangi, penelitian ini secara umum sejalan dengan peta jalan yang ada pada sub bidang *Artificial Intelligence*. *Roadmap* tersebut dapat dilihat pada Gambar 3.1 yang menunjukkan bagaimana penelitian ini berada dalam konteks pengembangan teknologi berbasis kecerdasan buatan di masa depan.



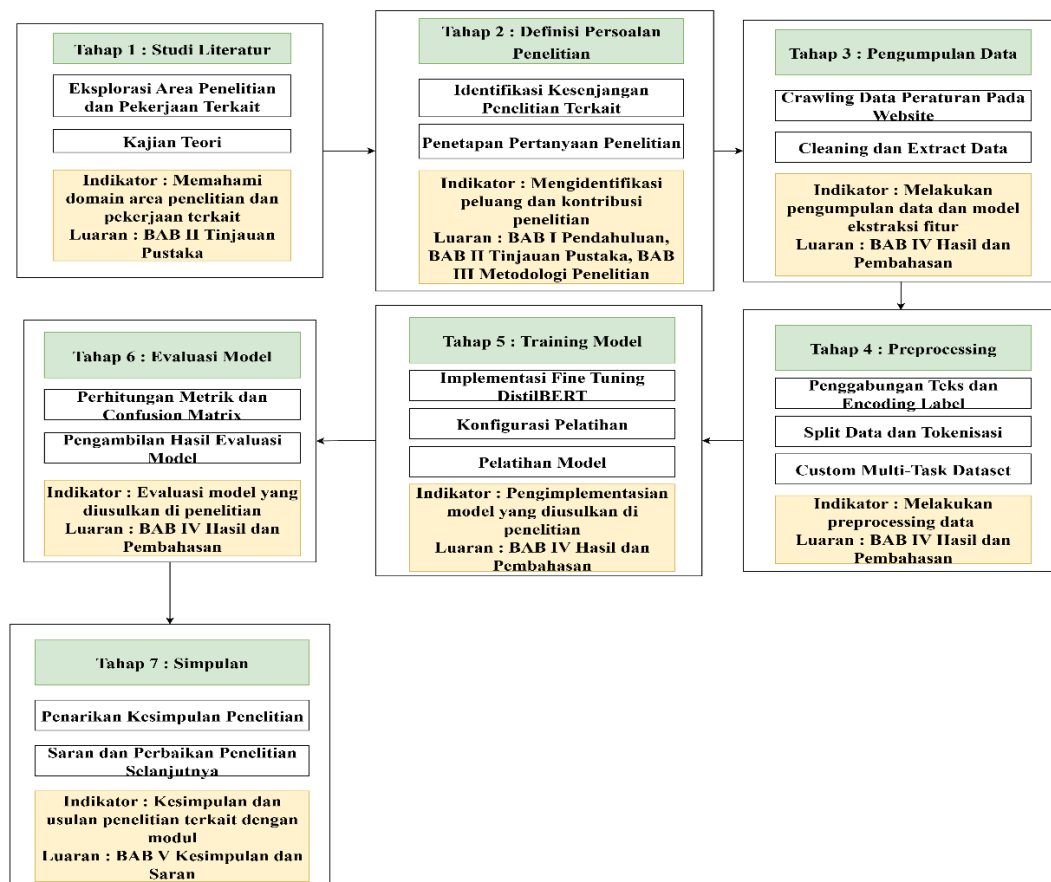
Gambar 3.1 *Roadmap* Penelitian AIS Universitas Siliwangi

Topik penelitian mengenai optimasi model *fine-tuned* DistilBERT untuk klasifikasi peraturan perundang-undangan sangat relevan dengan bidang *Natural Language Processing*, yang termasuk dalam kategori “*Basic Knowledge of Intelligent Informatics*”. Dengan memanfaatkan teknik-teknik NLP, model seperti DistilBERT mampu memahami dan mengklasifikasikan dokumen hukum. Selanjutnya, topik ini juga berhubungan dengan kategori “*Basic Application of Intelligent Informatics*” dalam hal *Natural Language Recognition and Generation*, dimana teknik pengenalan teks hukum digunakan untuk menentukan status dan jenis peraturan. Selain itu, penelitian ini masuk dalam kategori “*Applied Application of Intelligent Informatics*” di bidang *Classification and Clustering System*, dengan penggunaan model DistilBERT yang telah dilatih pada data hukum untuk mengklasifikasikan status dan jenis peraturan berdasarkan teks yang ada. Selain itu, penelitian ini masuk dalam kategori “*Market Product of Intelligent Informatics*” di bidang *Knowledge Based Systems and Intelligent Learning*, dengan melibatkan penerapan sistem berbasis pengetahuan dan teknik pembelajaran untuk mengklasifikasikan status dan jenis peraturan berdasarkan teks.

Dengan dasar yang kuat pada bidang *Artificial Intelligence* ini, diharapkan dapat tercipta penerapan pengetahuan yang lebih efektif dan berdaya saing tinggi dalam menghadapi tantangan global. Penelitian ini, yang fokus pada optimasi model *fine-tuned* DistilBERT untuk klasifikasi peraturan perundang-undangan, dapat mampu memberikan kontribusi signifikan dalam pengembangan teknologi *Natural Language Processing*, serta meningkatkan kualitas penelitian-penelitian lain yang berada dalam satu keilmuan.

3.2 Metode Penelitian

Penelitian merupakan penelitian kuantitatif dengan pendekatan eksperimen. Pendekatan eksperimen digunakan dikarenakan penelitian ini bertujuan untuk mengetahui pengaruh sebab akibat antara variabel-variabel yang ada seperti dataset, arsitektur model, dan parameter model (Lê & Schmid, 2022). Kebaruan yang ditargetkan dari penelitian yang diusulkan ini adalah penggunaan data peraturan perundang-undangan untuk melakukan *fine-tuning* menggunakan model DistilBERT dengan tujuan untuk klasifikasi peraturan perundang-undangan, khususnya dalam konteks hukum dengan bahasa Indonesia. Tahapan penelitian disajikan secara keseluruhan pada Gambar 3.2.



Gambar 3.2 Tahapan Penelitian

3.2.1 Studi Literatur

Studi literatur merupakan tahapan untuk mengumpulkan konsep, teori serta data-data dari berbagai sumber yang berhubungan dengan penelitian. Studi literatur dilakukan untuk pemahaman konsep, teori yang berhubungan dengan penelitian seperti teori mengenai arsitektur BERT, DistilBERT, dan terkait peraturan perundang-undangan. Pencarian informasi menggunakan sumber kedua (*web*, jurnal, *e-book*, artikel, dan lainnya). Selain itu, pada tahap studi literatur juga dilakukan “*review paper*” atau menganalisis penelitian terdahulu yang terkait dengan topik penelitian yang dilakukan.

3.2.2 Definisi Persoalan Penelitian

Definisi persoalan penelitian merupakan tahapan lanjutan yang dilakukan berdasarkan analisis penelitian terdahulu. Berdasarkan penelitian terdahulu yang telah dianalisis, dilakukan identifikasi kesenjangan dari penelitian-penelitian tersebut atau mengidentifikasi kekurangan dari penelitian, sehingga dapat dilakukan perbaikan. Setelah mengidentifikasi kesenjangan atau kekurangan penelitian terdahulu, kemudian menetapkan pertanyaan penelitian untuk mendapatkan tujuan dari penelitian yang akan dilakukan.

3.2.3 Pengumpulan Data

Pengumpulan data merupakan tahapan pertama sebelum dilakukan eksperimen. Dataset yang dikumpulkan akan digunakan untuk membuat model dan menguji model. Sumber yang dipilih untuk pengambilan data adalah *website* resmi pemerintah pada URL <https://peraturan.go.id>, yang menyediakan data hukum dalam jumlah besar. Pada tahapan ini data akan di *crawl* dari sumber yang telah

ditentukan dengan bentuk PDF. Data kemudian akan di ekstrak dan di *cleaning*, setelah itu data kemudian akan diformat sesuai dengan kebutuhan *training*. Peraturan-peraturan tersebut akan dipilih kembali signifikansinya terhadap hukum untuk selanjutnya dapat di *crawling* dan di ekstrak sesuai dengan *flow* yang digambarkan pada Gambar 3.3.



Gambar 3.3 Tahap Pengumpulan Dataset

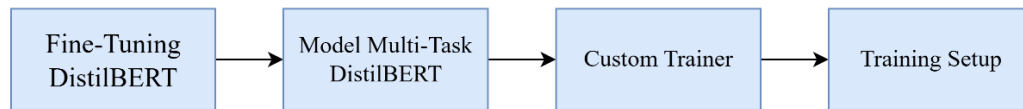
3.2.4 Preprocessing

Preprocessing merupakan tahapan persiapan sebelum melakukan eksperimen. Pada tahapan ini melakukan penggabungan beberapa kolom penting seperti *title*, *tentang*, dan *pemrakarsa* menjadi satu kolom teks terpadu guna memperkaya konteks *input*. Selanjutnya, label status peraturan dikonversi ke dalam format biner dan label jenis peraturan dikodekan menggunakan teknik label *encoding* agar dapat diproses oleh model. Data kemudian dibagi menjadi data pelatihan dan data uji untuk memastikan evaluasi model yang objektif dengan proporsi 80:20. Setelah itu, teks *input* ditokenisasi menggunakan *tokenizer* dari DistilBERT dengan proses *padding* dan *truncation* agar sesuai dengan panjang *input* maksimal. Seluruh data dikemas ke dalam format *custom* dataset yang mendukung *multi-task learning*, di mana setiap data mencakup input ID, *attention mask*, dan dua label target sekaligus, yaitu status dan jenis/bentuk peraturan. *Preprocessing* digambarkan pada Gambar 3.4.

Gambar 3.4 Tahap *Preprocessing*

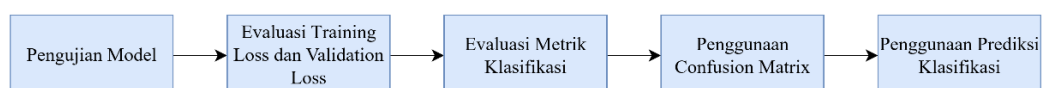
3.2.5 Training Model

Training model merupakan tahapan selanjutnya setelah dilakukan *preprocessing* untuk mengimplementasikan model DistilBERT. Implementasi digunakan pendekatan *multi-task learning* untuk dua tugas klasifikasi sekaligus, yaitu klasifikasi status peraturan dan klasifikasi jenis/bentuk peraturan. Model dilatih menggunakan dataset yang telah melalui proses tokenisasi dan *encoding* label, serta disusun dalam format *custom multi-task* dataset. Selama pelatihan, model dioptimasi menggunakan algoritma AdamW, yang dirancang khusus untuk mengatasi masalah *weight decay* pada model *transformer*, serta dilengkapi dengan teknik *dropout* untuk mencegah *overfitting*. Parameter pelatihan yang digunakan meliputi *learning rate* sebesar $5e-5$, *batch size* sebesar 40, jumlah *epoch* sebanyak 3, dan *weight decay* sebesar 0.01. *Training* dilakukan menggunakan modul *Trainer* dari *Hugging Face* dengan strategi evaluasi dan penyimpanan model pada setiap *epoch*. Model dilatih untuk meminimalkan dua fungsi *loss* sekaligus dari dua klasifikasi, sehingga dapat mengoptimalkan pembelajaran kedua label target secara paralel. Selama proses pelatihan, *training loss* dan *validation loss* dimonitor untuk mengevaluasi kestabilan model, serta untuk memastikan bahwa model tidak mengalami *overfitting* atau *underfitting*. *Training* model digambarkan pada Gambar 3.5.

Gambar 3.5 Tahap *Training* Model

3.2.6 Evaluasi Model

Evaluasi model merupakan pengukuran dan pengambilan hasil. Pada tahapan ini model yang telah di *fine-tuning* akan diuji. Evaluasi dilakukan terhadap dua *output* klasifikasi secara terpisah menggunakan metrik *accuracy*, *precision*, *recall*, *F1-score*, dan *confusion matrix*. Hasil prediksi terhadap data uji diekstrak dengan fungsi prediksi khusus, kemudian dibandingkan dengan label sebenarnya. *Confusion matrix* disertakan untuk memvisualisasikan jumlah prediksi benar dan salah dari setiap kelas, baik untuk label status (“Berlaku” dan “Tidak Berlaku”) maupun label jenis peraturan yang terdiri dari 12 kelas. Selain itu, performa model juga divisualisasikan dengan grafik *training loss* dan *validation loss* untuk mengamati stabilitas pelatihan. Disediakan pula fungsi prediksi teks untuk menguji hasil klasifikasi pada *input* baru secara *real-time*, termasuk prediksi terhadap status dan jenis peraturan beserta tingkat kepercayaannya (*confidence score*). Pengujian dilakukan juga pada data baru yang tidak termasuk dalam data pelatihan maupun pengujian, untuk mengamati kemampuan model dalam menangani teks yang benar-benar baru. Evaluasi model digambarkan pada Gambar 3.6.



Gambar 3.6 Tahap Evaluasi Model

3.2.7 Simpulan

Penarikan kesimpulan merupakan tahapan akhir yang dilakukan untuk memberikan gambaran umum terhadap analisis data dan hasil evaluasi model yang mencakup keseluruhan penelitian. Selain itu, kesimpulan berfungsi sebagai ringkasan utama dari temuan penelitian yang dapat menjadi dasar untuk rekomendasi, pengambilan keputusan, atau pengembangan lebih lanjut.