

BAB II

LANDASAN TEORI

2.1 Penelitian Terdahulu

Dalam mendukung penelitian ini, dilakukan tinjauan terhadap penelitian terdahulu yang berkaitan dengan *Central Bank Digital Currency* (CBDC), Rupiah Digital, keberlanjutan uang tunai, privasi pembayaran digital, opini pengguna di media sosial, serta analisis sentimen berbasis *Natural Language Processing*. Penelitian yang dikaji dipilih berdasarkan keterkaitan topik, metode, dan kontribusinya terhadap analisis sentimen serta klasifikasi teks. Sebagian besar rujukan berasal dari publikasi tahun 2023 hingga 2026 agar tetap relevan dengan perkembangan terbaru. Penelitian terdahulu membahas CBDC dan Rupiah Digital dari aspek hukum, sosial, privasi, kesiapan adopsi, kepercayaan pengguna, hingga narasi kebijakan. Selain itu, terdapat penelitian yang menggunakan algoritma *machine learning* dan *deep learning* seperti Naive Bayes, *Support Vector Machine*, dan model berbasis *transformer*. Evaluasi model umumnya menggunakan *accuracy*, *precision*, *recall*, *F1-score*, dan *confusion matrix*. Oleh karena itu, penelitian terdahulu digunakan sebagai dasar dalam pemilihan topik, metode, model klasifikasi, dan evaluasi penelitian. Adapun 10 penelitian terdahulu yang menjadi rujukan utama disajikan pada Tabel 2.1.

Tabel 2. 1 Penelitian Terdahulu

Penelitian Terdahulu 1	
Judul	Central Bank Digital Currency (CBDC): A Sentiment Analysis and Legal Perspective
Penulis	Aisyah As-Salafiyah, Aam Slamet Rusydiana, Ihsanul Ikhwan
Sumber jurnal	<i>Journal of Central Banking Law and Institutions</i> , Vol. 2 No. 2, Bank Indonesia.
Tahun	2023
Permasalahan	CBDC masih menjadi isu baru dan memerlukan kajian dari sisi persepsi serta legalitas.
Framework	Pengumpulan 50 paper Scopus → analisis sentimen menggunakan <i>SentiStrength</i> → kajian perspektif hukum.
Temuan	Sentimen terhadap CBDC didominasi netral sebesar 44%, positif 30%, dan negatif 26%.

Pembahasan	Hasil menunjukkan bahwa CBDC belum sepenuhnya dipersepsikan positif atau negatif, melainkan masih berada pada tahap evaluasi kritis.
Gap Penelitian	Belum menggunakan opini langsung dari masyarakat di media sosial dan belum berfokus pada Rupiah Digital berbasis komentar YouTube.
Relevansi terhadap penelitian ini	Relevan karena membahas CBDC, analisis sentimen, dan aspek hukum yang berkaitan dengan pengembangan Rupiah Digital.
Penelitian Terdahulu 2	
Judul	Measuring the Urgency of a Digital Rupiah: A Socio-Legal Review
Penulis	Sendy Pratama Firdaus
Sumber jurnal	<i>Journal of Central Banking Law and Institutions</i> , Vol. 4 No. 3, pp. 505–532.
Tahun	2025
Permasalahan	Urgensi Rupiah Digital perlu dikaji tidak hanya dari sisi teknologi, tetapi juga dari aspek hukum dan sosial masyarakat.
Framework	Kajian sosio-legal → analisis normatif, kontekstual, kritis, dan komparatif.
Temuan	Rupiah Digital memiliki potensi mengubah kebiasaan normatif masyarakat dalam menggunakan uang dan sistem pembayaran.
Pembahasan	Implementasi Rupiah Digital perlu mempertimbangkan kesiapan sosial, regulasi, dan dampaknya terhadap masyarakat.
Gap Penelitian	Belum melakukan klasifikasi sentimen berbasis komentar publik dan belum membandingkan model machine learning serta deep learning.
Relevansi terhadap penelitian ini	Sangat relevan karena membahas Rupiah Digital secara langsung dalam konteks Indonesia.
Penelitian Terdahulu 3	
Judul	<i>Privacy Implications of Central Bank Digital Currencies (CBDCs): A Systematic Review of Literature</i>
Penulis	Guneet Kaur
Sumber jurnal	<i>EDPACS</i> , Vol. 69 No. 9, pp. 87–123.
Tahun	2024
Permasalahan	Implementasi CBDC menimbulkan tantangan privasi, terutama terkait data transaksi dan anonimitas pengguna.
Framework	<i>Systematic Literature Review</i> → identifikasi isu privasi → analisis implikasi desain CBDC.
Temuan	Penelitian menekankan bahwa privasi merupakan salah satu isu utama dalam pengembangan CBDC.
Pembahasan	CBDC berbasis akun berpotensi mengurangi anonimitas pengguna, sehingga desain perlindungan data menjadi penting.

Gap Penelitian	Belum menganalisis bagaimana isu privasi dan keamanan data muncul dalam opini pengguna media sosial.
Relevansi terhadap penelitian ini	Relevan karena isu privasi dan anonimitas menjadi salah satu isu dominan dalam komentar publik terhadap Rupiah Digital dan uang tunai.
Penelitian Terdahulu 4	
Judul	Sentiment Analysis on YouTube Comments Using Machine Learning and Deep Learning with PCA- and LDA-Based Feature Selection
Penulis	Gulay Cicek, Nazlı Buldağ, Elif Aydın
Sumber Jurnal	<i>International Journal of Information Technology, Springer.</i>
Tahun	2025
Permasalahan	Komentar YouTube bersifat dinamis dan kompleks sehingga memerlukan perbandingan berbagai model klasifikasi sentimen.
Framework	Akuisisi komentar YouTube <i>real-time</i> → <i>preprocessing</i> → <i>PCA/LDA feature selection</i> → <i>11 model ML dan DL</i> → <i>evaluasi</i> .
Temuan	Penelitian membandingkan 11 model, termasuk <i>Naive Bayes</i> , <i>SVM</i> , <i>Logistic Regression</i> , <i>Random Forest</i> , <i>LSTM</i> , <i>CNN</i> , <i>BiLSTM</i> , <i>GRU</i> , <i>BiGRU</i> , dan <i>CNN-LSTM</i> .
Pembahasan	Hasil menunjukkan bahwa pemilihan fitur dan reduksi dimensi berpengaruh terhadap performa klasifikasi sentimen.
Gap Penelitian	Belum berfokus pada isu CBDC atau Rupiah Digital di Indonesia serta belum mengkaji bentuk kekhawatiran publik terhadap uang digital bank sentral.
Relevansi terhadap penelitian ini	Sangat relevan karena menggunakan komentar YouTube dan membandingkan model klasik serta deep learning.
Penelitian Terdahulu 5	
Judul	Comparative Analysis of Transformer Models for Sentiment Classification of UK CBDC Discourse on X
Penulis	Guneet Kaur, Saemundur Haraldsson, Andrea Bracciali [26].
Sumber jurnal	<i>Discover Analytics</i> , Vol. 3, Article 7, Springer.
Tahun	2025
Permasalahan	Diskursus publik mengenai CBDC di media sosial membutuhkan model sentimen yang mampu memahami konteks bahasa finansial.
Framework	Pengumpulan data <i>X</i> terkait <i>CBDC</i> → <i>fine-tuning DistilBERT, RoBERTa, XLM-RoBERTa</i> → evaluasi model.
Temuan	RoBERTa dengan 3 epoch memberikan generalisasi terbaik; model dievaluasi menggunakan <i>accuracy</i> , <i>precision</i> , <i>recall</i> , <i>F1-score</i> , serta <i>training/validation loss</i> .

Pembahasan	Pelatihan model yang terlalu lama dapat meningkatkan risiko <i>overfitting</i> , sehingga pemilihan epoch perlu diperhatikan.
Gap Penelitian	Belum menggunakan konteks Indonesia, belum membahas Rupiah Digital, dan belum membandingkan dengan model klasik seperti Naive Bayes dan SVM.
Relevansi terhadap penelitian ini	Relevan sebagai pembanding penggunaan model transformer untuk analisis sentimen CBDC di media sosial.
Penelitian Terdahulu 6	
Judul	CBDC in a Privacy Sensitive World: Extending UTAUT with Digital Financial Literacy and Anonymity Insight
Penulis	Vikrant Singh, Mayank Yadav, Ashish Gupta [27].
Sumber Jurnal	<i>Journal of Financial Services Marketing, Springer/Palgrave.</i>
Tahun	2026
Permasalahan	Adopsi CBDC dipengaruhi oleh privasi, anonimitas, literasi keuangan digital, dan kepercayaan pengguna.
Framework	<i>Model UTAUT</i> → survei 397 pengguna pembayaran <i>mobile</i> → <i>PLS-SEM, ANN, dan IPMA</i> .
Temuan	<i>Effort expectancy</i> , literasi keuangan digital, dan faktor anonimitas berpengaruh terhadap intensi adopsi <i>CBDC</i> .
Pembahasan	Desain CBDC perlu memperhatikan privasi dan literasi pengguna agar meningkatkan kepercayaan dan adopsi.
Gap Penelitian	Berbasis survei, bukan komentar media sosial, serta tidak melakukan klasifikasi sentimen menggunakan model NLP.
Relevansi terhadap penelitian ini	Relevan untuk memperkuat isu dominan seperti privasi, anonimitas, literasi digital, dan kepercayaan terhadap regulator.
Penelitian Terdahulu 7	
Judul	Sentiment Analysis of TikTok User Comments on QRIS Adoption in Indonesia Using IndoBERT
Penulis	Edi Supriyadi, Putra Nurhuda Makatita [28].
Sumber Jurnal	<i>Procedia Computer Science, Vol. 269, pp. 121–130.</i>
Tahun	2025
Permasalahan	Adopsi QRIS di Indonesia menimbulkan beragam opini publik yang dapat dianalisis melalui komentar media sosial.
Framework	<i>Scraping</i> komentar <i>TikTok</i> → <i>preprocessing</i> ringan → klasifikasi sentimen menggunakan <i>IndoBERT</i> .
Temuan	Penelitian menggunakan 1.128 komentar <i>TikTok</i> dan mengklasifikasikan sentimen menjadi positif, netral, dan negatif.
Pembahasan	Isu yang muncul mencakup kegagalan transaksi, keamanan, dan kebutuhan edukasi pengguna.

Gap Penelitian	Membahas QRIS, bukan Rupiah Digital, dan belum mengkaji kategori kekhawatiran publik terhadap CBDC atau keberlanjutan uang tunai.
Relevansi terhadap penelitian ini	Relevan karena membahas sistem pembayaran digital Indonesia dan menggunakan IndoBERT pada komentar media sosial.
Penelitian Terdahulu 8	
Judul	<i>Analysis of CBDC Narrative by Central Banks Using Large Language Models</i>
Penulis	Andres Alonso-Robisco, José Manuel Carbó [29].
Sumber jurnal	<i>Finance Research Letters</i> , Vol. 58, Article 104643.
Tahun	2023
Permasalahan	Narasi bank sentral terkait CBDC perlu dianalisis untuk memahami arah komunikasi kebijakan dan respons publik.
Framework	Pengumpulan berita <i>LexisNexis</i> → pemrosesan teks → analisis narasi menggunakan <i>BERT</i> dan ChatGPT.
Temuan	<i>LLM</i> dinilai efektif untuk meningkatkan pengukuran sentimen pada teks kebijakan, meskipun tetap memiliki risiko seperti sensitivitas terhadap panjang teks dan <i>prompt engineering</i> .
Pembahasan	Sentimen positif banyak muncul pada narasi peluncuran atau uji coba CBDC, sedangkan isu politik, regulasi, dan privasi dapat memicu sentimen negatif.
Gap Penelitian	Berfokus pada narasi bank sentral dan teks kebijakan, bukan komentar publik dari YouTube.
Relevansi terhadap penelitian ini	Relevan karena menunjukkan bahwa isu CBDC dapat dianalisis menggunakan pendekatan NLP dan model berbasis transformer.
Penelitian Terdahulu 9	
Judul	Sentimen Analysis Social Media for Disaster Using Naïve Bayes and IndoBERT
Penulis	Sri Mulyani Anugerah, Rifki Wijaya, Moch Arif Bijaksana [30].
Sumber Jurnal	<i>INTEK Jurnal Penelitian</i> , Vol. 11 No. 1.
Tahun	2024
Permasalahan	Data media sosial berbahasa Indonesia memiliki variasi bahasa yang kompleks sehingga memerlukan model klasifikasi yang tepat.
Framework	Pengumpulan data media sosial → <i>data preparation</i> → <i>labeling</i> → TF-IDF → klasifikasi <i>Naive Bayes</i> dan <i>IndoBERT</i> .
Temuan	IndoBERT memperoleh akurasi 91%, sedangkan Naive Bayes memperoleh akurasi 74%.
Pembahasan	IndoBERT lebih unggul karena mampu memahami konteks bahasa media sosial dibandingkan pendekatan probabilistik sederhana.

Gap Penelitian	Tidak membahas isu Rupiah Digital atau CBDC, serta belum menggunakan SVM dan analisis kategori kekhawatiran publik.
Relevansi terhadap penelitian ini	Relevan sebagai acuan komparasi model klasik dan transformer pada teks media sosial berbahasa Indonesia.
Penelitian Terdahulu 10	
Judul	Prediction Sentiment Analysis Grab Reviews Using SVM Linear Based <i>Streamlit</i>
Penulis	M. T. Hidayat, M. Arifin, S. Muzid [31].
Sumber Jurnal	<i>IJCCS (Indonesian Journal of Computing and Cybernetics Systems)</i> , Vol. 19 No. 2, pp. 189–200.
Tahun	2025
Permasalahan	Hasil analisis sentimen perlu disajikan secara interaktif agar mudah dipahami oleh pengguna dan pemangku kepentingan.
Framework	<i>Data Collection</i> → <i>Text Preprocessing</i> → <i>Feature Extraction</i> → <i>Supervised Learning</i> (SVM, <i>Random Forest</i> , <i>Naïve Bayes</i>) → Evaluasi Metrik.
Temuan	Model memperoleh akurasi 83%, precision 85%, recall 83%, dan F1-score 81%.
Pembahasan	<i>Streamlit</i> digunakan untuk menampilkan distribusi sentimen, confusion matrix, classification report, word cloud, bar chart, dan pie chart.
Gap Penelitian	Objek penelitian berbeda, yaitu ulasan aplikasi Grab, dan belum membahas isu uang digital bank sentral maupun kategori kekhawatiran publik.
Relevansi terhadap penelitian ini	Relevan karena penelitian ini juga menggunakan TF-IDF, SVM, evaluasi model, dan <i>Dashboard Streamlit</i> untuk visualisasi hasil.

Seluruh penelitian terdahulu pada Tabel 2.1 menunjukkan bahwa kajian mengenai *Central Bank Digital Currency* (CBDC), Rupiah Digital, dan sistem pembayaran digital tidak hanya berkaitan dengan aspek teknologi, tetapi juga mencakup privasi, kepercayaan terhadap regulator, literasi digital, kesiapan masyarakat, serta penerimaan pengguna. Beberapa penelitian menekankan bahwa privasi, anonimitas, dan kepercayaan menjadi faktor penting dalam adopsi CBDC. Sementara itu, penelitian lain menunjukkan bahwa opini pengguna pada media sosial dapat dianalisis menggunakan pendekatan *Natural Language Processing* dan model klasifikasi sentimen.

Dari sisi metode, penelitian terdahulu menunjukkan bahwa *Naive Bayes* dan *Support Vector Machine* masih banyak digunakan sebagai model pembanding

karena sederhana dan efisien dalam klasifikasi teks. Namun, model berbasis *transformer* seperti *IndoBERT* dinilai lebih mampu memahami konteks bahasa, terutama pada data media sosial yang mengandung bahasa informal dan tidak terstruktur. Selain itu, penggunaan *Dashboard* seperti *Streamlit* juga dapat membantu menyajikan hasil analisis secara lebih interaktif dan mudah dipahami.

Meskipun demikian, masih terdapat celah penelitian yang dapat dikembangkan. Sebagian besar penelitian mengenai CBDC dan Rupiah Digital lebih banyak membahas aspek hukum, privasi, kebijakan, atau kesiapan adopsi, sedangkan penelitian analisis sentimen pada media sosial belum banyak yang secara khusus membahas Rupiah Digital dan keberlanjutan uang tunai berdasarkan komentar YouTube berbahasa Indonesia. Oleh karena itu, penelitian ini dilakukan untuk menganalisis sentimen pengguna YouTube, mengidentifikasi bentuk kekhawatiran dan tanggapan terhadap keberlanjutan uang tunai, membandingkan performa *Naive Bayes*, *Support Vector Machine*, dan *IndoBERT*, serta menyajikan hasilnya melalui *Dashboard Streamlit*.

2.2 Teori yang berkaitan

2.2.1 Central Bank Digital Currency

Central Bank Digital Currency atau CBDC merupakan bentuk uang digital yang diterbitkan oleh bank sentral dan menjadi kewajiban langsung bank sentral. CBDC muncul seiring dengan perkembangan digitalisasi sistem keuangan, meningkatnya penggunaan pembayaran non-tunai, serta kebutuhan bank sentral untuk tetap menyediakan uang yang aman dan terpercaya di era digital. CBDC berbeda dengan uang elektronik, dompet digital, maupun aset kripto karena diterbitkan oleh otoritas moneter resmi dan dinyatakan dalam satuan mata uang nasional. Secara umum, CBDC dikembangkan untuk mendukung efisiensi sistem pembayaran, memperkuat inovasi keuangan, serta menjaga peran uang bank sentral dalam sistem moneter modern. Namun, penerapan CBDC juga perlu mempertimbangkan berbagai aspek, seperti stabilitas moneter, stabilitas sistem keuangan, keamanan data, privasi pengguna, regulasi, dan kesiapan infrastruktur. Selain itu, CBDC tidak dimaksudkan untuk langsung menggantikan seluruh bentuk uang yang sudah

ada, melainkan dapat berjalan berdampingan dengan uang tunai dan instrumen pembayaran lain [32].

2.2.2 Rupiah Digital

Rupiah Digital merupakan bentuk CBDC yang dikembangkan oleh Bank Indonesia sebagai representasi digital dari mata uang Rupiah. Pengembangan Rupiah Digital dilakukan melalui Proyek Garuda, yaitu inisiatif Bank Indonesia untuk mengeksplorasi desain, arsitektur, teknologi, serta dukungan regulasi dalam penerapan mata uang digital bank sentral di Indonesia. Rupiah Digital berbeda dengan uang elektronik atau dompet digital karena diterbitkan oleh Bank Indonesia sebagai bank sentral, sedangkan uang elektronik dan dompet digital diterbitkan atau dikelola oleh penyedia jasa pembayaran. Rupiah Digital dikembangkan sebagai respons terhadap perkembangan ekonomi dan keuangan digital, meningkatnya kebutuhan transaksi yang cepat dan efisien, serta pentingnya menjaga kedaulatan Rupiah di era digital. Melalui Rupiah Digital, Bank Indonesia berupaya menyediakan bentuk uang bank sentral yang sesuai dengan perkembangan teknologi dan kebutuhan sistem pembayaran masa depan. Namun, pengembangannya tetap perlu memperhatikan aspek keamanan, privasi, kesiapan infrastruktur, regulasi, serta penerimaan masyarakat [33].

2.2.2.1 Faktor Kekhawatiran dan Tanggapan Publik terhadap Rupiah Digital

Tanggapan publik terhadap *Central Bank Digital Currency* dapat terbentuk melalui persepsi masyarakat terhadap manfaat, keterbatasan, dan kekhawatiran atas sistem pembayaran digital yang sudah digunakan sebelumnya. Sikap terhadap CBDC juga dipengaruhi oleh konteks penggunaan pembayaran serta karakteristik pengguna. Pada penelitian mengenai Digital Euro, beberapa bentuk kekhawatiran yang paling sering muncul berkaitan dengan privasi dan keamanan. Responden juga menilai bahwa keamanan merupakan prasyarat penting sebelum menggunakan mata uang digital bank sentral. Selain itu, kepercayaan terhadap bank sentral

dapat memengaruhi harapan masyarakat terhadap perlindungan data dan privasi transaksi [34].

Kekhawatiran terhadap CBDC juga berkaitan dengan persepsi masyarakat mengenai kerentanan data, kontrol transaksi, dan kemampuan individu dalam menghadapi risiko privasi. Persepsi kerentanan terhadap penyalahgunaan data berhubungan dengan meningkatnya kekhawatiran privasi, sedangkan kemampuan individu dalam mengelola ancaman privasi berhubungan dengan penurunan tingkat kekhawatiran tersebut. Selain itu, kekhawatiran privasi terbukti dapat menurunkan niat masyarakat untuk menggunakan CBDC. Faktor kepercayaan, khususnya kredibilitas dan citra bank sentral, juga berperan dalam membentuk kekhawatiran privasi serta memengaruhi niat penggunaan CBDC [35].

Selain privasi dan kepercayaan, pemahaman masyarakat mengenai CBDC juga berhubungan dengan penerimaan terhadap mata uang digital bank sentral. Pengetahuan mengenai CBDC dan kepercayaan terhadap bank sentral ditemukan memiliki hubungan positif dengan niat masyarakat untuk menggunakan CBDC. Selain itu, kebutuhan masyarakat terhadap privasi dan keamanan perlu diperhatikan dalam perancangan CBDC serta komunikasi publik mengenai karakteristiknya [36]. Dalam penelitian ini, faktor kekhawatiran dan tanggapan publik diidentifikasi melalui kategori keamanan data dan kontrol, kepercayaan terhadap regulator, literasi dan kesiapan masyarakat, akses teknologi dan infrastruktur, kebutuhan terhadap uang tunai, kekhawatiran uang tunai dihapus, efisiensi dan kemudahan transaksi, serta tanggapan mendukung digitalisasi.

2.2.3 Analisis Sentimen

Analisis sentimen adalah bidang dalam *Natural Language Processing* yang bertujuan mengidentifikasi opini, sikap, emosi, atau penilaian dalam teks. Tujuannya adalah mengubah teks tidak terstruktur menjadi informasi terukur dengan mengelompokkan sentimen menjadi kelas positif, negatif, atau netral. Analisis ini membantu memahami persepsi pengguna terhadap produk, layanan,

kebijakan, isu sosial, atau fenomena tertentu. Pendekatannya meliputi *lexicon-based*, *machine learning*, dan *deep learning*. Pada data media sosial, tantangannya lebih besar karena teks sering menggunakan bahasa informal, singkatan, slang, emoji, sarkasme, dan konteks yang tidak eksplisit [37]. Analisis sentimen juga dapat digunakan untuk mengukur penerimaan publik terhadap suatu isu melalui data teks dengan membandingkan beberapa model klasifikasi, seperti *Naive Bayes*, *Support Vector Machine*, dan *Long Short-Term Memory* [38]. Dalam penelitian ini, analisis sentimen digunakan untuk mengklasifikasikan komentar YouTube terkait Rupiah Digital dan keberlanjutan uang tunai ke dalam kelas Positif, Negatif, dan Netral. Hasil klasifikasi tersebut kemudian digunakan sebagai dasar untuk melihat kecenderungan opini pengguna.

2.2.3.1 Pelabelan Sentimen dan Anotasi Data

Pelabelan sentimen dan anotasi data merupakan proses pemberian kategori pada data teks berdasarkan kelas yang telah ditentukan. Dalam analisis sentimen, kelas yang umum digunakan meliputi positif, negatif, dan netral. Label tersebut menjadi dasar bagi model *machine learning* untuk mempelajari hubungan antara isi teks dan kategori sentimennya. Oleh karena itu, kualitas anotasi perlu diperhatikan karena kesalahan, bias, atau ketidakkonsistenan label dapat memengaruhi proses pelatihan serta hasil evaluasi model [39].

Pelabelan data dapat dilakukan secara manual, otomatis, maupun semi-otomatis. Pelabelan manual dilakukan oleh manusia berdasarkan interpretasi terhadap isi dan konteks teks. Pelabelan otomatis menggunakan aturan, kamus sentimen, atau metode komputasional. Sementara itu, pelabelan semi-otomatis menggabungkan pembentukan label awal secara otomatis dengan pemeriksaan manusia pada data tertentu. Dalam analisis sentimen, pelibatan manusia tetap diperlukan karena komentar dapat mengandung negasi, sentimen campuran, bahasa informal, sarkasme, atau konteks implisit yang sulit ditangkap sepenuhnya oleh aturan otomatis [40]. Dalam penelitian ini, pelabelan sentimen dilakukan melalui pendekatan

semi-otomatis. Label awal dibentuk menggunakan *lexicon-based labelling*, kemudian disempurnakan melalui *context rules* dan kalibrasi. Sebagian data juga divalidasi secara manual untuk memastikan konsistensi hasil pelabelan.

2.2.3.2 Pedoman, Kompetensi, dan Objektivitas Annotator

Pedoman anotasi diperlukan untuk membantu annotator menerapkan kategori sentimen secara konsisten. Pedoman perlu menjelaskan definisi setiap kelas dan arahan untuk menghadapi komentar yang memiliki penafsiran berbeda. Dalam anotasi sentimen, instruksi dapat memengaruhi pengalaman annotator serta tingkat kesepakatan label. Namun, pedoman yang terlalu rinci tidak selalu meningkatkan konsistensi. Oleh karena itu, pedoman perlu disusun secara jelas, sesuai dengan kebutuhan tugas, dan mudah diterapkan oleh annotator [41].

Objektivitas anotasi dapat ditingkatkan melalui penggunaan skema anotasi yang sama, pelabelan secara independen, dan pengukuran kesepakatan antar-annotator. Pada komentar daring, konteks juga perlu diperhatikan karena makna komentar dapat bergantung pada komentar sebelumnya atau isi percakapan. Ketersediaan konteks dapat meningkatkan kualitas anotasi manual, terutama pada komentar balasan yang lebih bergantung pada konteks diskusi [42].

Kompetensi annotator perlu disesuaikan dengan karakteristik data dan kompleksitas tugas anotasi. Pemilihan annotator, jumlah pilihan label, serta rancangan proses pengumpulan anotasi dapat memengaruhi kualitas data berlabel dan performa model *machine learning*. Oleh karena itu, proses anotasi perlu dirancang secara cermat agar menghasilkan data yang dapat digunakan secara lebih valid dalam pengembangan model [43]. Dalam penelitian ini, annotator menggunakan pedoman pelabelan sentimen positif, negatif, dan netral. Validasi dilakukan secara independen oleh dua annotator pada sampel komentar, kemudian tingkat kesepakatan label diukur menggunakan *Cohen's Kappa*.

2.2.4 Machine Learning

Machine learning merupakan cabang dari kecerdasan buatan yang memungkinkan sistem mempelajari pola dari data untuk menghasilkan prediksi, klasifikasi, atau pengambilan keputusan tanpa harus diprogram secara eksplisit untuk setiap aturan. Dalam pendekatan ini, model dibangun melalui proses pelatihan menggunakan data, kemudian pola yang telah dipelajari digunakan untuk memprediksi data baru. *Machine learning* banyak digunakan dalam berbagai bidang, termasuk pengenalan pola, klasifikasi dokumen, analisis teks, dan sistem rekomendasi. Dalam klasifikasi teks, *machine learning* digunakan untuk mempelajari hubungan antara fitur teks dan label kelas tertentu. Karena teks tidak dapat diproses secara langsung oleh sebagian besar algoritma klasik, data teks perlu diubah terlebih dahulu ke dalam bentuk numerik melalui proses representasi fitur, seperti TF-IDF [44]. Pada penelitian ini, *machine learning* digunakan untuk mengklasifikasikan komentar YouTube ke dalam sentimen Positif, Negatif, dan Netral menggunakan model *Naive Bayes* dan *Support Vector Machine*.

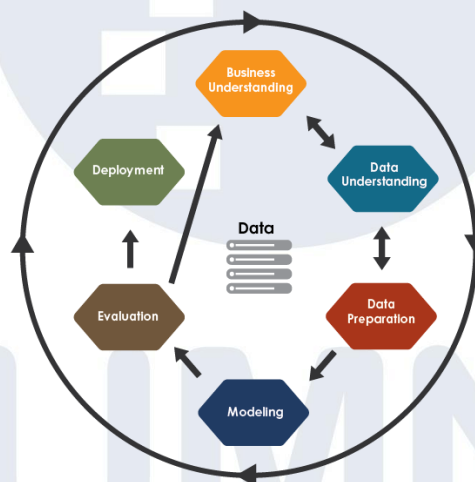
2.2.5 Deep Learning

Deep learning merupakan bagian dari *machine learning* yang menggunakan jaringan saraf tiruan dengan banyak lapisan untuk mempelajari representasi data secara bertingkat. Melalui beberapa lapisan pemrosesan, model *deep learning* dapat mempelajari pola sederhana hingga pola yang lebih kompleks dari data. Pendekatan ini banyak digunakan pada berbagai bidang, seperti pengenalan citra, pengenalan suara, pemrosesan bahasa alami, dan klasifikasi teks. Dalam pengolahan bahasa alami, *deep learning* digunakan karena mampu mempelajari hubungan antarkata dan konteks dalam kalimat secara lebih kompleks dibandingkan pendekatan berbasis fitur manual. Perkembangan model berbasis *transformer*, seperti *BERT*, semakin meningkatkan kemampuan model dalam memahami teks karena dapat memperhatikan hubungan kata dalam suatu kalimat secara kontekstual [45]. Pada penelitian ini, pendekatan *deep learning* digunakan melalui model *IndoBERT* untuk klasifikasi sentimen komentar berbahasa Indonesia.

2.3 Framework/Algoritma yang digunakan

2.3.1 CRISP-DM

Dalam penelitian berbasis *data mining* dan *machine learning*, diperlukan suatu kerangka kerja yang sistematis agar proses analisis dapat berjalan secara terstruktur dan terarah. Salah satu metodologi yang banyak digunakan dalam pengembangan proyek analitik data adalah CRISP-DM (*Cross-Industry Standard Process for Data Mining*). CRISP-DM merupakan standar proses pengolahan data yang dirancang untuk membantu peneliti dalam memahami permasalahan, mengolah data, membangun model, serta mengimplementasikan hasil analisis secara sistematis dan iteratif. Metodologi ini bersifat fleksibel dan dapat diterapkan pada berbagai bidang, termasuk analisis sentimen berbasis teks [46].



Gambar 2. 1 Tahapan CRISP-DM

Sumber: [47]

Berdasarkan Gambar 2.1, CRISP-DM terdiri dari enam tahapan utama yang membentuk suatu siklus berulang (*iterative cycle*). Setiap tahapan saling berkaitan dan memungkinkan perbaikan berkelanjutan hingga diperoleh model yang optimal. Adapun tahapan tersebut adalah sebagai berikut:

1. Business Understanding

Tahap ini merupakan tahap awal untuk memahami tujuan, kebutuhan, dan permasalahan yang ingin diselesaikan. Pada tahap ini, permasalahan didefinisikan secara jelas, tujuan analisis ditentukan, serta kriteria keberhasilan proyek dirumuskan. Tahap ini penting agar proses analisis data memiliki arah yang sesuai dengan kebutuhan yang ingin dicapai.

2. **Data Understanding**

Tahap ini berfokus pada proses pengumpulan dan pemahaman awal terhadap data yang akan digunakan. Kegiatan pada tahap ini mencakup identifikasi sumber data, eksplorasi struktur data, pemeriksaan kualitas data, serta pengenalan pola awal yang terdapat dalam data. Tahap ini membantu peneliti atau analis memahami karakteristik data sebelum dilakukan proses pengolahan lebih lanjut.

3. **Data Preparation**

Tahap ini bertujuan untuk menyiapkan data agar dapat digunakan dalam proses pemodelan. Proses yang dilakukan dapat meliputi pembersihan data, pemilihan atribut, transformasi data, integrasi data, serta pembentukan format data yang sesuai dengan kebutuhan model. Tahap ini umumnya membutuhkan waktu yang cukup besar karena kualitas data sangat memengaruhi hasil analisis.

4. **Modeling**

Tahap *modeling* merupakan proses pemilihan dan penerapan teknik pemodelan yang sesuai dengan tujuan analisis. Pada tahap ini, beberapa algoritma atau metode dapat diuji untuk memperoleh model yang paling sesuai. Proses ini juga dapat melibatkan pengaturan parameter, pelatihan model, serta pengujian awal terhadap performa model.

5. **Evaluation**

Tahap evaluasi dilakukan untuk menilai apakah model yang dibangun telah memenuhi tujuan dan kriteria keberhasilan yang telah ditentukan sebelumnya. Di samping meninjau aspek teknis performa model, tahap evaluasi ini turut berfungsi untuk memastikan bahwa hasil analisis telah menjawab kebutuhan yang direncanakan. Apabila hasil belum memadai, proses dapat kembali ke tahap sebelumnya untuk dilakukan perbaikan.

6. Deployment

Tahap terakhir adalah *deployment*, yaitu penerapan hasil analisis atau model ke dalam bentuk yang dapat digunakan. Bentuk *deployment* dapat berupa laporan, visualisasi, sistem pendukung keputusan, aplikasi, atau integrasi model ke dalam proses operasional. Langkah ini diambil guna memastikan hasil analisis bukan sekadar berhenti di ranah pemodelan abstrak, namun dapat diimplementasikan secara nyata oleh pihak pengguna.

2.3.2 Term Frequency-Inverse Document Frequency

Term Frequency-Inverse Document Frequency atau TF-IDF merupakan metode pembobotan kata yang digunakan untuk mengubah data teks menjadi representasi numerik. Metode ini bekerja dengan menghitung tingkat kepentingan suatu kata dalam sebuah dokumen berdasarkan frekuensi kemunculan kata tersebut pada dokumen tertentu dan seberapa jarang kata tersebut muncul pada keseluruhan kumpulan dokumen. Dengan demikian, kata yang sering muncul pada satu komentar tetapi tidak terlalu umum pada seluruh dataset akan memiliki bobot yang lebih tinggi [48].

Secara umum, rumus TF-IDF dapat dituliskan sebagai berikut:

$$TFIDF(t, d) = TF(t, d) \times IDF(t) \quad (2.1)$$

Keterangan:

$TF(t, d)$ menunjukkan frekuensi kemunculan kata t dalam dokumen d , sedangkan $IDF(t)$ menunjukkan bobot kebalikan dari jumlah dokumen yang mengandung kata tersebut. Semakin tinggi nilai TF-IDF, maka semakin penting kata tersebut dalam merepresentasikan dokumen tertentu.

2.3.3 IndoBERT

IndoBERT merupakan model *pre-trained language model* berbasis arsitektur *Bidirectional Encoder Representations from Transformers* atau *BERT* yang dikembangkan khusus untuk pemrosesan bahasa Indonesia. Model ini dilatih menggunakan korpus teks berbahasa Indonesia dalam jumlah besar

sehingga mampu memahami karakteristik bahasa Indonesia, seperti struktur kalimat, variasi kata, imbuhan, konteks semantik, serta penggunaan kata dalam berbagai bentuk kalimat. Berbeda dengan model pembelajaran mesin tradisional seperti *Naive Bayes* dan *Support Vector Machine* yang bergantung pada representasi fitur statis, *IndoBERT* mampu menghasilkan representasi kata yang bersifat kontekstual atau *contextual embeddings*.

Secara umum, *BERT* bekerja menggunakan arsitektur *transformer encoder* yang memanfaatkan mekanisme *self-attention*. Mekanisme ini memungkinkan model untuk memperhatikan hubungan antar kata dalam suatu kalimat secara dua arah atau *bidirectional*. Artinya, makna sebuah kata tidak hanya dipahami berdasarkan kata sebelumnya, tetapi juga berdasarkan kata sesudahnya. Kemampuan ini penting dalam analisis sentimen karena makna komentar sering kali dipengaruhi oleh konteks kalimat secara keseluruhan. Misalnya, kata “aman” dapat bermakna positif, tetapi dalam frasa “tidak aman” maknanya berubah menjadi negatif. Dengan pendekatan *bidirectional*, *IndoBERT* dapat menangkap perubahan makna tersebut dengan lebih baik.

Dalam tugas klasifikasi sentimen, *IndoBERT* digunakan dengan menambahkan lapisan klasifikasi pada keluaran model. Teks komentar terlebih dahulu diproses menggunakan *tokenizer* untuk mengubah kalimat menjadi token yang dapat dibaca oleh model. Setelah itu, token tersebut diproses melalui beberapa lapisan *transformer encoder* untuk menghasilkan representasi kalimat. Representasi tersebut kemudian digunakan untuk memprediksi kelas sentimen, yaitu positif, negatif, atau netral. Proses penyesuaian model terhadap dataset penelitian dilakukan melalui tahap *fine-tuning*, yaitu melatih kembali model *IndoBERT* menggunakan data komentar yang telah diberi label.

Keunggulan utama *IndoBERT* dalam analisis sentimen terletak pada kemampuannya memahami konteks bahasa Indonesia yang sering tidak eksplisit. Komentar YouTube umumnya menggunakan bahasa informal, singkatan, pertanyaan, negasi, kekhawatiran, dan sentimen campuran. TF-IDF hanya berfokus pada kemunculan kata, sedangkan *IndoBERT*

mempertimbangkan urutan dan hubungan antar kata dalam kalimat. Hal ini membuat *IndoBERT* lebih efektif untuk menangani komentar dengan konteks kompleks, seperti isu keamanan data, kontrol transaksi, kepercayaan terhadap regulator, dan kekhawatiran terhadap Rupiah Digital [49]. Penelitian ini menggunakan *IndoBERT* sebagai model klasifikasi sentimen yang dibandingkan dengan *Naive Bayes* dan *Support Vector Machine*. Data untuk *IndoBERT* tidak melalui *stemming* dan *stopword removal*, hanya light cleaning, agar struktur dan konteks asli komentar tetap terjaga. Pendekatan ini diperlukan karena model transformer memanfaatkan urutan kata untuk memahami makna teks. Dengan cara ini, *IndoBERT* diharapkan memberikan hasil klasifikasi yang lebih akurat, terutama pada komentar yang mengandung konteks, negasi, dan makna tidak langsung.

2.3.4 Naive Bayes

Naive Bayes merupakan algoritma klasifikasi berbasis probabilistik yang menggunakan Teorema Bayes untuk menentukan kemungkinan suatu data termasuk ke dalam kelas tertentu. Algoritma ini disebut “naive” karena memiliki asumsi bahwa setiap fitur bersifat independen terhadap fitur lainnya. Dalam konteks analisis sentimen, fitur yang dimaksud dapat berupa kata atau bobot kata hasil representasi TF-IDF, sedangkan kelas yang diprediksi adalah sentimen positif, negatif, atau netral [50]. Secara umum, Teorema Bayes dapat dirumuskan sebagai berikut:

$$P(C|X) = \frac{P(X|C) \times P(C)}{P(X)} \quad (2.2)$$

Keterangan:

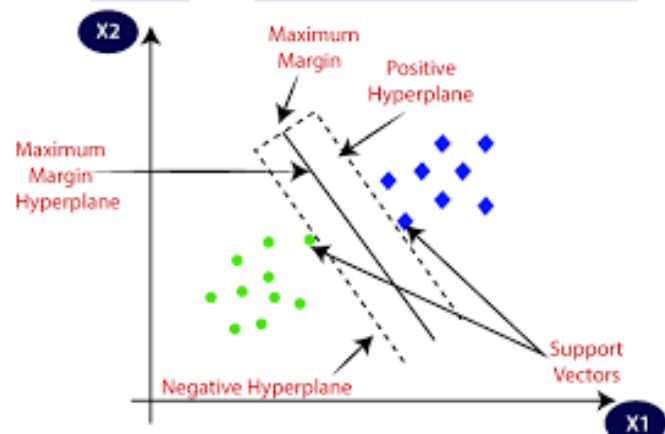
$P(C | X)$ adalah probabilitas kelas C berdasarkan data X , $P(X | C)$ adalah probabilitas data X muncul pada kelas C , $P(C)$ adalah probabilitas awal kelas, dan $P(X)$ adalah probabilitas data.

Pada analisis sentimen, *Naive Bayes* masih banyak digunakan karena memiliki proses klasifikasi yang sederhana, efisien, dan sesuai untuk data teks

berbasis fitur numerik seperti TF-IDF. Penelitian sebelumnya menunjukkan bahwa Naive Bayes dapat digunakan dalam evaluasi metode analisis sentimen pada aplikasi media sosial dan dibandingkan dengan *Support Vector Machine* untuk menilai performa klasifikasi teks [51]. Oleh karena itu, *Naive Bayes* digunakan dalam penelitian ini sebagai salah satu model pembandingan dalam klasifikasi sentimen komentar YouTube terhadap Rupiah Digital.

2.3.5 Support Vector Machine

Support Vector Machine atau SVM merupakan algoritma *supervised learning* yang digunakan untuk klasifikasi maupun regresi. Dalam proses klasifikasi, SVM bekerja dengan mencari garis atau bidang pemisah terbaik yang disebut *hyperplane* untuk memisahkan data ke dalam kelas yang berbeda. Tujuan utama SVM adalah membentuk *hyperplane* dengan jarak pemisah atau *margin* terbesar antara kelas yang satu dengan kelas lainnya [52].



Gambar 2. 2 Ilustrasi SVM

Sumber: [53]

Gambar 2.2 menunjukkan ilustrasi cara kerja SVM dalam membentuk *hyperplane* untuk memisahkan dua kelas data. Titik-titik data yang berada paling dekat dengan *hyperplane* disebut *support vectors*. Garis pemisah utama berada di antara dua kelas, sedangkan jarak antara *hyperplane* dan *support vectors* disebut *margin*. Semakin besar nilai *margin*, semakin baik kemampuan model dalam memisahkan kelas. Dalam penelitian ini, konsep tersebut

digunakan untuk memisahkan komentar YouTube berdasarkan pola fitur teks hasil pembobotan TF-IDF.

2.3.6 Lexicon-Based Labelling dan Penyempurnaan Label Berbasis

Aturan Konteks

Lexicon-based labelling merupakan pendekatan pelabelan teks yang menggunakan kamus sentimen atau *lexicon* untuk mengidentifikasi kecenderungan sentimen suatu kata atau frasa. Dalam pendekatan ini, teks dipecah menjadi token, kemudian token yang sesuai dengan *lexicon* diberi nilai atau kategori sentimen tertentu. Nilai tersebut selanjutnya dapat digabungkan untuk memperkirakan kecenderungan sentimen suatu teks. Pendekatan ini relatif mudah diimplementasikan, lebih efisien dibanding pelabelan manual penuh, dan tetap dapat memberikan hasil yang cukup informatif untuk analisis teks, terutama ketika peneliti membutuhkan proses pelabelan yang terstruktur dan dapat dijelaskan [54].

Meskipun demikian, pendekatan *lexicon-based* memiliki keterbatasan ketika diterapkan pada data media sosial. Teks media sosial sering mengandung bahasa informal, struktur kalimat yang tidak baku, serta nuansa sentimen yang dipengaruhi oleh konteks. Selain itu, pendekatan ini juga dapat mengalami kesulitan dalam membedakan opini yang bergantung pada domain atau konteks tertentu. Oleh karena itu, hasil pelabelan berbasis *lexicon* dapat disempurnakan melalui aturan berbasis konteks agar interpretasi sentimen tidak hanya bergantung pada kemunculan kata dalam kamus, tetapi juga mempertimbangkan pola kalimat dan isyarat konteks yang menyertainya [55].

Dalam penelitian ini, label awal dibentuk menggunakan pendekatan *lexicon-based labelling*. Selanjutnya, label tersebut disempurnakan melalui *context rules* dan kalibrasi untuk meninjau kembali komentar yang masih mengandung ambiguitas atau sinyal sentimen yang belum tertangkap secara optimal oleh pelabelan awal.

2.3.7 Hyperparameter Tuning dengan GridSearchCV

Hyperparameter tuning merupakan proses pencarian kombinasi parameter terbaik pada suatu model agar model dapat menghasilkan performa yang lebih optimal. Berbeda dengan parameter yang dipelajari langsung oleh model selama proses pelatihan, *hyperparameter* ditentukan sebelum proses pelatihan dilakukan. Pemilihan *hyperparameter* yang tepat dapat memengaruhi kemampuan model dalam mengenali pola data dan melakukan prediksi terhadap data baru. Salah satu metode yang dapat digunakan untuk melakukan *hyperparameter tuning* adalah *GridSearchCV*. Metode ini bekerja dengan mencoba beberapa kombinasi nilai *hyperparameter* yang telah ditentukan, kemudian mengevaluasi setiap kombinasi menggunakan *cross validation*. Kombinasi parameter dengan nilai evaluasi terbaik akan dipilih sebagai parameter optimal untuk model [56].

2.3.8 Evaluasi Model Klasifikasi

Evaluasi model klasifikasi bertujuan untuk mengukur sejauh mana model mampu mengklasifikasikan data secara akurat dan konsisten. Dalam penelitian ini, evaluasi dilakukan menggunakan *confusion matrix*, yaitu tabel yang menggambarkan perbandingan antara hasil prediksi model dan label sebenarnya. Berdasarkan *confusion matrix*, dapat dihitung beberapa metrik evaluasi yang umum digunakan dalam klasifikasi, yaitu *accuracy*, *precision*, *recall*, dan *F1-score* [57].

Confusion Matrix

Secara umum, *confusion matrix* terdiri dari empat komponen utama:

- **True Positive (TP):** data positif yang diprediksi benar sebagai positif.
- **True Negative (TN):** data negatif yang diprediksi benar sebagai negatif.
- **False Positive (FP):** data negatif yang diprediksi sebagai positif.
- **False Negative (FN):** data positif yang diprediksi sebagai negatif.

Accuracy

Accuracy mengukur proporsi prediksi yang benar terhadap seluruh data.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (2.3)$$

Precision

Precision mengukur seberapa tepat model dalam memprediksi kelas positif.

$$Precision = \frac{TP}{TP + FP} \quad (2.4)$$

Recall

Recall (atau *Sensitivity*) mengukur kemampuan model dalam mendeteksi seluruh data positif.

$$Recall = \frac{TP}{TP + FN} \quad (2.5)$$

F1-Score

F1-score merupakan rata-rata harmonis antara *precision* dan *recall*.

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (2.6)$$

2.3.9 Cohen's Kappa

Cohen's Kappa merupakan ukuran statistik untuk menilai tingkat kesepakatan antara dua annotator pada data kategorikal. Berbeda dari persentase *agreement*, *Cohen's Kappa* mempertimbangkan kemungkinan kesamaan label yang terjadi secara kebetulan, sehingga lebih kuat digunakan untuk menilai reliabilitas pelabelan. *Cohen's Kappa* membandingkan kesepakatan aktual antar-annotator dengan kesepakatan yang diperkirakan terjadi secara kebetulan. Dalam analisis sentimen, ukuran ini penting karena data teks sering memiliki makna yang tidak selalu tegas, seperti negasi, sindiran, bahasa informal, pertanyaan tidak langsung, atau sentimen campuran. Dengan demikian, *Cohen's Kappa* dapat digunakan untuk melihat apakah pedoman pelabelan telah diterapkan secara konsisten oleh annotator [58].

Rumus *Cohen's Kappa* adalah sebagai berikut:

$$\kappa = \frac{P_o - P_e}{1 - P_e} \quad (2.7)$$

Keterangan:

κ

= nilai Cohen's Kappa

= $\frac{P_o}{P_e}$ proporsi kesepakatan aktual antar-annotator

= proporsi kesepakatan yang terjadi secara kebetulan

Cohen's Kappa	Interpretation
0	No agreement
0.10 - 0.20	Slight agreement
0.21 - 0.40	Fair agreement
0.41 - 0.60	Moderate agreement
0.61 - 0.80	Substantial agreement
0.81 - 0.99	Near perfect agreement
1	Perfect agreement

Gambar 2. 3 Interpretasi Nilai Cohen's Kappa

Sumber: [59]

Berdasarkan Gambar 2.3, *Cohen's Kappa* digunakan untuk mengukur tingkat kesepakatan antar-annotator, dengan nilai 0 menunjukkan tidak ada kesepakatan dan nilai 1 menunjukkan kesepakatan sempurna. Semakin tinggi nilainya, semakin kuat tingkat kesepakatan dalam pemberian label. Interpretasinya terbagi ke dalam beberapa kategori, yaitu slight, fair, moderate, substantial, hingga near perfect agreement, yang digunakan untuk menilai reliabilitas pelabelan data kategorikal.

Dalam penelitian ini, interpretasi *Cohen's Kappa* digunakan untuk menilai tingkat kesepakatan antara annotator pertama dan annotator kedua pada proses validasi manual label sentimen. Hasil pengujian tersebut menjadi dasar untuk melihat apakah proses pelabelan telah dilakukan secara konsisten sebelum data digunakan pada tahap pemodelan.

2.3.10 McNemar Test

McNemar Test merupakan metode pengujian statistik yang digunakan untuk membandingkan dua model klasifikasi yang diuji pada dataset yang sama. Uji ini digunakan ketika peneliti ingin mengetahui apakah perbedaan hasil prediksi antara dua model benar-benar signifikan secara statistik atau hanya

terjadi karena variasi hasil prediksi biasa. Dalam konteks evaluasi model klasifikasi, dua model dapat memiliki nilai *accuracy* yang berbeda, tetapi perbedaan tersebut belum tentu cukup kuat untuk menyatakan bahwa salah satu model benar-benar lebih baik. Oleh karena itu, *McNemar Test* digunakan sebagai pengujian tambahan untuk memperkuat hasil perbandingan model. Berbeda dengan evaluasi umum seperti *accuracy*, *precision*, *recall*, dan *F1-score*, *McNemar Test* tidak berfokus langsung pada nilai metrik performa, melainkan pada pola kesalahan prediksi antara dua model. Uji ini membandingkan jumlah data yang diprediksi benar oleh Model A tetapi salah oleh Model B, serta jumlah data yang diprediksi salah oleh Model A tetapi benar oleh Model B. Dengan demikian, *McNemar Test* dapat menunjukkan apakah dua model memiliki perbedaan kemampuan prediksi yang signifikan pada data uji yang sama [60].

Dalam penelitian ini, *McNemar Test* digunakan untuk membandingkan performa model *Naive Bayes*, *Support Vector Machine*, dan *IndoBERT* secara berpasangan. Pengujian ini penting karena setiap model memiliki pendekatan berbeda dalam memproses teks, sehingga perbandingan tidak hanya berdasarkan metrik evaluasi, tetapi juga didukung oleh bukti statistik terhadap perbedaan hasil prediksi antar model.

Rumus *McNemar Test* ditunjukkan sebagai berikut:

$$\chi^2 = \frac{(b - c)^2}{b + c} \quad (2.8)$$

Keterangan:

- *b* adalah jumlah data yang benar pada Model A tetapi salah pada Model B.
- *c* adalah jumlah data yang salah pada Model A tetapi benar pada Model B.
- χ^2 adalah nilai statistik uji.

Apabila nilai *p-value* lebih kecil dari tingkat signifikansi yang digunakan, misalnya 0,05, maka terdapat perbedaan performa yang signifikan antara kedua model. Sebaliknya, apabila nilai *p-value* lebih besar dari 0,05, maka perbedaan performa kedua model tidak signifikan secara statistik. Dalam

konteks penelitian ini, hasil pengujian tersebut digunakan untuk memperkuat analisis komparasi model klasifikasi sentimen.

2.4 Tools/software yang digunakan

2.4.1 Python

Python merupakan bahasa pemrograman tingkat tinggi yang banyak digunakan dalam pengembangan perangkat lunak, analisis data, *machine learning*, *artificial intelligence*, dan komputasi ilmiah. *Python* memiliki sintaks yang relatif sederhana dan mudah dipahami, sehingga dapat digunakan oleh pemula maupun pengguna yang sudah berpengalaman dalam pemrograman. Selain itu, *Python* didukung oleh banyak pustaka yang dapat membantu proses pengolahan data, visualisasi, pemodelan, dan pengembangan aplikasi [61]. Dalam penelitian ini, *Python* digunakan sebagai bahasa pemrograman utama untuk melakukan pengumpulan data, pembersihan data, *preprocessing*, pelabelan, pemodelan, evaluasi, serta pembuatan *Dashboard*. Penggunaan *Python* didukung oleh berbagai pustaka seperti *pandas*, *scikit-learn*, *transformers*, *matplotlib*, dan *Streamlit*.

2.4.2 Google Colab

Google Colab dan *Jupyter Notebook* digunakan sebagai lingkungan pengembangan dalam penelitian ini karena menyediakan platform interaktif untuk menulis, menjalankan, dan mendokumentasikan kode *Python* secara terstruktur. *Google Colab* memiliki keunggulan berupa akses berbasis *cloud* yang memungkinkan penggunaan sumber daya komputasi seperti GPU tanpa memerlukan perangkat keras dengan spesifikasi tinggi, sehingga sangat mendukung proses pelatihan model *Deep Learning* seperti *IndoBERT*. Sementara itu, *Jupyter Notebook* memungkinkan eksplorasi data, visualisasi, serta pengujian model secara bertahap dalam satu dokumen yang terintegrasi. Penggunaan kedua tools ini memudahkan proses eksperimen, dokumentasi, serta reproduksibilitas penelitian secara keseluruhan [62].

2.4.3 Tools Pengumpulan Data

Proses pengumpulan data dalam penelitian ini dilakukan menggunakan *YouTube Data API v3* yang disediakan oleh *Google Cloud Console*. API ini

digunakan untuk mengambil komentar publik dari video YouTube yang telah dipilih berdasarkan relevansi topik penelitian. Melalui *YouTube Data API v3*, data komentar dapat diperoleh secara lebih terstruktur, termasuk atribut seperti *VideoID*, *CommentDate*, *UserName*, *Comment*, *Like*, dan *ScrapeDate* [63].

2.4.4 Streamlit



Gambar 2. 4 Streamlit

Sumber: [64]

Streamlit merupakan salah satu framework berbasis *Python* yang digunakan untuk membangun aplikasi web interaktif secara cepat dan sederhana, khususnya dalam bidang *data science* dan *machine learning*. *Streamlit* memungkinkan pengguna untuk mengubah script *Python* menjadi aplikasi web tanpa memerlukan pengetahuan mendalam tentang pengembangan web, sehingga memudahkan dalam menampilkan hasil analisis dalam bentuk visual seperti grafik, tabel, dan output model secara interaktif.

2.4.5 Visual Studio Code

Visual Studio Code merupakan *source code editor* yang dikembangkan oleh Microsoft untuk menulis, mengelola, dan menjalankan kode program. Aplikasi ini mendukung berbagai bahasa pemrograman, seperti *Python*, *JavaScript*, *HTML*, dan *CSS*. *Visual Studio Code* memiliki fitur seperti *syntax highlighting*, *debugging*, *code completion*, *integrated terminal*, serta dukungan *extension* yang dapat membantu proses pengembangan perangkat lunak menjadi lebih mudah, efisien, dan terstruktur [65].