

BAB II

TINJAUAN PUSTAKA

2.1 Penelitian yang Relevan

Penelitian yang berkaitan ini dilakukan untuk menganalisis dan membandingkan hasil penelitian terdahulu dari jurnal serta makalah yang relevan sebagai referensi guna memperkuat dasar penelitian ini. Pembahasan mencakup topik penelitian, metode yang diterapkan, serta temuan yang diperoleh.

Penelitian pertama, melakukan perbandingan dua metode peramalan deret waktu, yaitu ARIMA dan (LSTM), dalam memprediksi harga saham PT Gojek Tokopedia Tbk (GOTO.JK). Data yang digunakan berupa harga *closing price* saham GOTO dari Desember 2022 hingga Maret 2024 yang diambil dari *Yahoo Finance*. Data tersebut dibagi menjadi 16 kelompok, masing-masing terdiri dari data latih dan data uji, kemudian dianalisis menggunakan *Python*. Model ARIMA dikembangkan dengan parameter (p,d,q) yang ditentukan berdasarkan hasil uji *stasioneritas*, sedangkan model LSTM dibangun dengan enam lapisan, meliputi dua lapisan LSTM, dua lapisan *dropout*, dan dua lapisan *dense*, menggunakan *Adam optimizer* serta fungsi *loss* MSE. Hasil evaluasi menunjukkan bahwa ARIMA memiliki akurasi yang lebih tinggi dibandingkan LSTM, dengan nilai RMSE sebesar 3,259 dan MAPE 3,57% (akurasi 96,53%), sedangkan LSTM mencatat RMSE 3,843 dan MAPE 4,04% (akurasi 95,96%). Dengan demikian, dapat disimpulkan bahwa ARIMA lebih efektif dalam memprediksi harga saham GOTO dalam jangka pendek karena menghasilkan kesalahan prediksi yang lebih kecil, sementara LSTM kurang efisien untuk data dengan pola fluktuasi yang sangat tinggi (Adwin Adam et al., 2024).

Penelitian yang kedua, difokuskan pada layanan Garuda Indonesia *Cargo*, terutama di kantor cabang *Cargo Service Center (CSC) Tangerang City*, membahas pengembangan model prediksi pendapatan melalui penerapan algoritma LSTM. Pemilihan algoritma LSTM ini didasarkan pada kemampuannya yang unggul dalam memproses data deret waktu secara efektif, sehingga cocok untuk menangani pola temporal dalam transaksi. Data utama yang dimanfaatkan berasal dari catatan transaksi pengiriman barang harian, yang telah melalui proses *preprocessing* lengkap meliputi empat tahap utama, yaitu perhitungan *subtotal* untuk agregasi awal, deteksi *outlier* guna membersihkan anomali, proses *differencing* untuk menstabilkan variasi, serta *scaling* untuk normalisasi skala data. Berdasarkan hasil analisis, komposisi pembagian data dengan 90% untuk

pelatihan dan 10% untuk pengujian terbukti memberikan performa optimal, di mana nilai RMSE mencapai 641.375,70 pada data latih dan 594.197,70 pada data uji, yang secara keseluruhan menandakan tingkat akurasi prediksi yang cukup baik dan andal (Aprian et al., 2020).

Penelitian ketiga, membahas penerapan metode peramalan penjualan untuk mendukung pengambilan keputusan bisnis pada usaha mikro, dengan studi kasus penjualan bulanan *Gociko Snack*. Model yang digunakan adalah ARIMA, yang dipilih karena data yang dianalisis berupa data *time series* dengan jumlah observasi relatif terbatas serta memiliki ketergantungan antar periode waktu. Proses pemodelan dilakukan melalui analisis *stasioneritas* dan identifikasi parameter ARIMA yang sesuai, kemudian dievaluasi menggunakan metrik MAPE. Hasil evaluasi menunjukkan bahwa model ARIMA mampu menghasilkan peramalan dengan tingkat kesalahan yang rendah, ditunjukkan oleh nilai MAPE sebesar 9,18%, sehingga hasil peramalan dinilai cukup akurat dan dapat diandalkan. Berdasarkan hasil tersebut, penelitian ini menyimpulkan bahwa model ARIMA efektif dan relevan untuk digunakan dalam peramalan penjualan atau pendapatan pada skala usaha kecil hingga menengah, terutama ketika data historis yang tersedia relatif pendek dan diperlukan metode peramalan yang sederhana serta stabil (Hartanti & Permatasari, 2024).

Penelitian keempat, membandingkan efektivitas beberapa model dalam memprediksi produksi energi di Tiongkok pada masa pandemi COVID-19. Model yang diuji mencakup pendekatan statistik seperti ARIMA dan ETS, serta model *machine learning* seperti *LightGBM*, NNAR, dan LSTM. Penilaian kinerja menggunakan empat metrik utama, yaitu RMSE, MAE, MPE, dan MAPE. Hasilnya menunjukkan bahwa *LightGBM* secara keseluruhan memiliki akurasi paling baik, dengan nilai MAPE antara 1,79% hingga 4,00%, sehingga dianggap paling akurat untuk menangani data yang tidak stabil akibat pandemi. Namun, saat membandingkan ARIMA dan LSTM, keduanya memiliki kelebihan pada jenis data tertentu. ARIMA lebih baik untuk deret waktu yang stabil, seperti produksi tenaga air dan gas alam, dengan MAPE terendah mencapai 1,66%. Sedangkan LSTM lebih efektif pada data yang berfluktuasi, berkat kemampuannya menangkap pola waktu yang rumit dan menggabungkan faktor eksternal seperti dampak pandemi COVID-19. Meskipun performa LSTM kurang maksimal dibandingkan dengan *LightGBM*, model ini tetap memberikan prediksi yang solid dengan MAPE di bawah 10%. Secara umum, penelitian ini menyimpulkan bahwa *LightGBM* adalah metode terbaik secara keseluruhan, tetapi ARIMA cocok untuk data stabil, sementara LSTM lebih unggul pada data yang dinamis dan tidak beraturan (Lu et al., 2021).

Penelitian keempat, menganalisis dan membandingkan tiga metode peramalan deret waktu, yaitu ARIMA, LSTM, dan *Prophet*, untuk memprediksi penjualan harian roti putih di Bandung *White Bread Factory* yang berlokasi di Denpasar, Bali, menggunakan 822 data harian yang dikumpulkan dari Maret 2021 hingga Mei 2023, dengan pembagian 80% data pelatihan dan 20% data pengujian guna memastikan kestabilan hasil. Evaluasi dilakukan melalui tiga metrik utama, yakni MAPE untuk mengukur akurasi dalam persentase MSE untuk rata-rata kuadrat kesalahan, serta RMSE untuk menilai besarnya penyimpangan prediksi terhadap nilai sebenarnya. Hasil menunjukkan bahwa model ARIMA (1,0,2) memberikan performa terbaik dengan MSE 5,27%, dan RMSE 47,4139, diikuti LSTM dengan MAPE 7,3275%, MSE 6,33%, serta RMSE 56,9761, sementara *Prophet* mencatat MAPE 7,402%, MSE 8,83%, dan RMSE 79,5137. Secara keseluruhan, ARIMA lebih unggul pada data penjualan yang stabil, tidak terlalu berfluktuasi, dan memiliki tren jangka pendek yang jelas, meskipun LSTM dan *Prophet* yang lebih kompleks dan mampu menangani pola *non-linear* serta musiman cenderung kurang optimal di sini karena LSTM sedikit mengalami *overfitting* pada data sederhana, sedangkan *Prophet* lebih cocok untuk data dengan musiman jangka panjang seperti penjualan tahunan atau *e-commerce* skala besar. Dengan demikian, penelitian ini menekankan bahwa pemilihan metode peramalan harus disesuaikan dengan karakteristik data, di mana ARIMA paling ideal untuk data stabil, sementara LSTM dan *Prophet* lebih sesuai untuk data yang kompleks dan berfluktuasi tinggi (Suryawan et al., 2024).

Penelitian keenam menekankan pada implementasi dan perbandingan dua pendekatan peramalan, yakni ARIMA dan *Exponential Smoothing*, guna meramalkan anggaran belanja lingkungan di Provinsi Jawa Barat untuk periode 2022-2025. Analisis dilaksanakan dengan memanfaatkan data realisasi belanja lingkungan dari tahun 1994 hingga 2021, didukung oleh perangkat lunak *Eviews* 9. Temuan penelitian mengungkapkan bahwa metode ARIMA menunjukkan akurasi yang lebih unggul dibandingkan *Exponential Smoothing*, berkat nilai kesalahan yang jauh lebih rendah dengan MAPE sebesar 19,042%. Nilai tersebut mengindikasikan keunggulan ARIMA dalam memprediksi data deret waktu (*time series*) dan menjadikannya sebagai metode yang paling sesuai untuk memperkirakan anggaran fungsi lingkungan di Provinsi Jawa Barat (Firdayati et al., 2024).

Penelitian ketujuh membahas survei dan analisis praktis tentang berbagai cara memproses data awal pada deret waktu berupa angka, dengan tujuan membuat data lebih baik dan meningkatkan keakuratan model kecerdasan buatan, terutama LSTM. Mereka menggunakan

pendekatan berbasis angka dengan data tambahan dari repositori UCI *Machine Learning*, yaitu *dataset* Kualitas Udara yang punya 9.358 pengamatan *multivariat*, di mana 90% untuk latihan dan 10% untuk uji coba. Penulis menjelaskan bahwa data deret waktu yang diambil langsung dari sensor biasanya penuh dengan gangguan, nilai ekstrem, dan data yang hilang. jadi perlu langkah awal seperti *smoothing*, penyesuaian skala, pengisian data hilang, dan pendeteksian nilai ekstrem sebelum digunakan model peramalan. Untuk melihat dampaknya, semua teknik ini diuji dengan model LSTM utama, menggunakan ukuran evaluasi seperti RMSE, MAE, dan MAPE. Hasilnya menunjukkan bahwa data yang sudah diproses lengkap membuat performa model naik drastis, dengan MAPE turun dari sekitar 51,41% pada data asli menjadi 25,26% setelah diproses, plus penurunan RMSE dan MAE hampir 50%. Penelitian ini menyimpulkan bahwa memproses data awal, termasuk penghalusan dan perubahan distribusi data, sangat penting dalam analisis deret waktu karena bisa kurangi gangguan, stabilkan pola data, dan langsung tingkatkan akurasi model peramalan berbasis LSTM, sehingga cocok diterapkan pada data yang fluktuatif dan jumlahnya terbatas (Zbezhkhovska & Chumachenko, 2025).

Penelitian kedelapan, menerapkan pendekatan komputasional berbasis *deep learning* dengan metode kuantitatif untuk meningkatkan interpretasi *neural-biomarker* melalui proses *data preprocessing*. Penelitian ini menggunakan data sekunder berupa citra otak (*neuroimaging*) yang bersifat kompleks, berukuran relatif terbatas, serta mengandung noise tinggi, sehingga berpotensi menurunkan kinerja model *deep learning*. Untuk mengatasi permasalahan tersebut, peneliti menerapkan *Gaussian smoothing* dan *modified histogram normalization* sebagai tahap awal pengolahan data sebelum dimasukkan ke dalam model jaringan saraf. *Gaussian smoothing* digunakan dengan tujuan mereduksi *noise* dan memperhalus pola data, sedangkan normalisasi histogram diterapkan untuk menyelaraskan distribusi intensitas data agar lebih konsisten antar sampel. Model *deep learning* kemudian digunakan untuk melakukan klasifikasi dan interpretasi *biomarker*, dengan evaluasi kinerja menggunakan metrik seperti *acuration*, *precision*, *recall*, dan *error classification*. Hasil penelitian menunjukkan bahwa penerapan *Gaussian smoothing* secara signifikan meningkatkan kualitas sinyal data dan menghasilkan peningkatan akurasi model hingga sekitar 94,7%, serta menurunkan kesalahan klasifikasi dibandingkan data tanpa *smoothing*. Berdasarkan hasil tersebut, penelitian ini menyimpulkan bahwa *Gaussian smoothing* merupakan teknik *preprocessing* yang efektif untuk *dataset* berukuran kecil dan bising, karena mampu meningkatkan stabilitas data dan membantu model *deep learning* dalam mengenali pola secara

lebih akurat, sehingga pendekatan ini relevan untuk diterapkan pada berbagai kasus pemodelan data kompleks dengan keterbatasan jumlah observasi (Usman et al., 2021).

Penelitian ketujuh menerapkan pendekatan komputasional berbasis *machine learning* dan *deep learning* dengan metode kuantitatif untuk memprediksi kekuatan tekan beton *geopolymer* yang disusun dari limbah *fly ash* hasil pembakaran batu bara serta *ground granulated blast furnace slag* (GGBFS). Sumber data yang digunakan bersifat sekunder, diperoleh dari hasil eksperimen terdahulu yang melibatkan 156 sampel dengan variasi komposisi material, meliputi *fly ash*, *slag*, agregat halus dan kasar, larutan alkali (NaOH dan Na_2SiO_3), tingkat molaritas, rasio aktivator terhadap *binder*, serta kondisi *curing*. Data tersebut diolah dengan pembagian 80% untuk tahap pelatihan (*training*) dan 20% untuk tahap pengujian (*testing*), disertai proses normalisasi guna mengurangi potensi bias pada model. Peneliti membandingkan kinerja lima algoritma prediksi, yaitu *Deep Long Short-Term Memory (LSTM)*, *Artificial Neural Network (ANN)*, *Bootstrap Aggregating (Bagging)*, *Least-Squares Boosting (LSBoost)*, dan *K-Nearest Neighbors (kNN)*, dengan evaluasi berdasarkan metrik R^2 , RMSE, dan MAPE. Hasil penelitian menunjukkan bahwa model *Deep LSTM* memberikan performa paling unggul dengan nilai R^2 sebesar 0,9819 dan MAPE sebesar 4,56%, diikuti oleh ANN dengan R^2 sebesar 0,9501 serta *LSBoost* dengan R^2 sebesar 0,9473. Analisis sensitivitas juga mengindikasikan bahwa faktor yang paling berpengaruh terhadap kekuatan tekan adalah agregat halus (99,57%), kandungan SiO_2 dan CaO pada *slag* (99,4%), serta kandungan SiO_2 pada *fly ash* (96,8%). Dengan demikian, pendekatan *Deep LSTM* terbukti paling andal dalam memodelkan hubungan *nonlinier* antara komposisi bahan dan kekuatan beton *geopolymer*, sekaligus memperlihatkan bahwa penerapan *machine learning* mampu menggantikan sebagian besar pengujian laboratorium konvensional dengan hasil prediksi yang presisi dan efisien (Kina et al., 2025).

Penelitian sepuluh mengeksplorasi potensi pemanfaatan limbah *fly ash* dan *bottom ash* hasil pembakaran batu bara abu dengan proporsi yang berbeda, kemudian menguji karakteristik fisik, mekanik, dan serta *pet-coke* sebagai komponen utama dalam pembuatan panel tahan api yang ramah lingkungan. Melalui serangkaian uji laboratorium, peneliti merancang berbagai variasi formulasi ketahanan panas untuk memperoleh data empiris yang menjadi dasar proses pemodelan. Setelah data diperoleh, tahap selanjutnya melibatkan penerapan sejumlah algoritma *machine learning* guna memprediksi dan mengoptimalkan kinerja material berdasarkan hasil eksperimental tersebut. Pada model regresi yang digunakan untuk memprediksi durasi ketahanan terhadap api

(t_{180}), algoritma *Random Forest Regressor* menunjukkan performa terbaik dengan nilai koefisien determinasi (R^2) sebesar 0,87 dan tingkat kesalahan maksimum sekitar 11% antara hasil prediksi dan pengukuran aktual. Hasil prediksi waktu ketahanan api pada masing-masing komposisi panel menunjukkan tingkat konsistensi yang tinggi, misalnya pada sampel B-0 (0% *bottom ash*) dengan waktu aktual 28,2 menit yang diprediksi sebesar 28,4 menit (kesalahan 0,7%), sedangkan pada sampel B-40 (40% *bottom ash*) waktu aktual 25,3 menit diprediksi sebesar 23,2 menit (kesalahan 8,5%). Sementara itu, pada model regresi untuk kekuatan tekan (*compressive strength*), algoritma *XGBoost* menghasilkan nilai R^2 tertinggi yaitu 0,79, meskipun menunjukkan sedikit kecenderungan *overfitting* akibat keterbatasan jumlah data. Prediksi kekuatan tekan umumnya berada di bawah nilai aktual dengan selisih rata-rata sekitar -40%, contohnya pada sampel B-0 dengan nilai aktual 2,02 MPa dan hasil prediksi 1,15 MPa. Pada analisis klasifikasi, model *Random Forest Classifier* berhasil mengelompokkan kategori *Fire Resistance Rating (FRR)* dengan akurasi mencapai 97%, di mana hanya terjadi satu kesalahan klasifikasi antara kelas FRR-15 dan FRR-30. Selain itu, model *Decision Tree Classifier* memperoleh akurasi tertinggi sebesar 99% dalam mengklasifikasikan rentang kekuatan tekan material. Berdasarkan hasil tersebut, dapat disimpulkan bahwa model berbasis *ensemble learning* seperti *Random Forest* dan *XGBoost* sangat efektif dalam menangkap hubungan *nonlinier* antara komposisi abu, densitas, ketebalan, dan sifat *termomekanik* material. Penerapan *machine learning* dalam penelitian ini terbukti mampu mempercepat proses perancangan material tahan api dengan tingkat akurasi yang tinggi, sekaligus mengurangi ketergantungan pada pengujian laboratorium yang memerlukan waktu dan biaya besar (Guirado et al., 2025).

Tabel 2. 1 Matriks Penelitian

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
Perbandingan Model ARIMA dan LSTM dalam Peramalan Harga Saham: Studi Kasus GOTO.JK				
Membandingkan model ARIMA dan LSTM dalam konteks peramalan data deret	Penelitian tersebut hanya berfokus pada data harga saham	Hasil penelitian menunjukkan bahwa model ARIMA memiliki	2024	Adwin Adam, H., Raditia

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
waktu. Relevansinya terletak pada fokus terhadap pengujian performa kedua model dalam memprediksi data yang bersifat fluktuatif dan <i>non-linear</i> , serupa dengan karakteristik data FABA yang digunakan pada penelitian ini. Studi ini memberikan dasar empiris untuk menilai perbandingan kemampuan model statistik konvensional dengan model berbasis jaringan saraf dalam menghasilkan prediksi jangka pendek yang akurat.	bersifat finansial dan berfluktuasi harian, tanpa meninjau penerapan model pada konteks lingkungan atau industri seperti data FABA. Selain itu, penelitian ini belum membahas secara mendalam mengenai pengaruh parameter hiper (<i>hyperparameter tuning</i>) maupun potensi penerapan model gabungan atau hibrida. Dengan demikian, penelitian ini berupaya mengisi gap tersebut dengan mengimplementasi kan ARIMA dan LSTM pada data FABA yang	tingkat akurasi yang lebih tinggi dibandingkan LSTM, dengan nilai RMSE sebesar 3,259 dan MAPE 3,57% (akurasi 96,53%), sedangkan LSTM mencatat RMSE 3,843 dan MAPE 4,04% (akurasi 95,96%). Temuan ini menegaskan bahwa ARIMA lebih efektif dalam memprediksi data deret waktu jangka pendek, khususnya pada pola yang cenderung stabil, sementara LSTM kurang efisien ketika menghadapi data dengan fluktuasi yang sangat tinggi.		nsyah, F., Rayyan Imani, M., Faris Fawwa z, M., & Ridho Lubis, A.

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
	memiliki pola non-linear serta fluktuasi yang berbeda dari data keuangan.			
Prediksi Pendapatan Kargo Menggunakan Arsitektur Long Short Term Memory				
Menggunakan model LSTM dalam konteks peramalan data deret waktu. Relevansi penelitian ini terletak pada penerapan LSTM untuk memprediksi pendapatan berdasarkan data transaksi harian yang bersifat fluktuatif dan temporal, serupa dengan karakteristik data FABA yang juga menunjukkan dinamika waktu yang bervariasi. Studi ini turut memberikan landasan empiris terkait efektivitas LSTM dalam mengelola data deret	Penelitian tersebut hanya menggunakan model LSTM tanpa melakukan perbandingan langsung dengan model statistik seperti ARIMA, sehingga belum dapat menggambarkan keunggulan relatif model <i>deep learning</i> terhadap pendekatan tradisional. Selain itu, penelitian ini berfokus pada data keuangan perusahaan penerbangan,	Pembagian data sebesar 90% untuk pelatihan dan 10% untuk pengujian memberikan hasil paling optimal, dengan nilai RMSE sebesar 641.375,70 pada data latih dan 594.197,70 pada data uji. Nilai tersebut menunjukkan bahwa model LSTM mampu menghasilkan prediksi yang cukup akurat dan andal terhadap pendapatan kargo Garuda Indonesia. Temuan ini	2020	Aprian, B. A., Azhar, Y., Rahma yanti, V., & Nastiti, S.

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
waktu yang kompleks dan memerlukan proses <i>preprocessing</i> yang matang sebelum pelatihan model.	sehingga masih terbuka peluang penerapan pada sektor lain seperti pengelolaan limbah industri energi, termasuk data FABAs. Oleh karena itu, penelitian ini berupaya mengisi celah tersebut dengan membandingkan performa model LSTM dan ARIMA pada data FABAs yang memiliki pola fluktuasi dan variasi temporal yang kompleks.	memperkuat efektivitas model LSTM dalam menangani data deret waktu finansial yang memiliki pola musiman serta fluktuasi tinggi, berkat kemampuannya dalam mengenali hubungan temporal jangka panjang.		
Enhancing Sales Performance through ARIMA-Based Predictive Modeling: Insights and Applications Model				
Berkaitan langsung dengan penelitian yang kamu lakukan karena sama-sama membahas peramalan	Penelitian ini hanya menggunakan model ARIMA tunggal	Penelitian menunjukkan bahwa model ARIMA mampu menghasilkan	2024	Dwi Hartanti dan Hanifah

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
pendapatan/penjualan berbasis data <i>time series</i> menggunakan model ARIMA. Penelitian ini memperkuat bahwa ARIMA efektif digunakan pada data bisnis dengan jumlah observasi terbatas, sehingga relevan sebagai pembanding dan landasan empiris dalam peramalan pendapatan FABAs.	melakukan perbandingan dengan model lain seperti LSTM atau metode <i>machine learning</i> , serta belum menerapkan teknik <i>preprocessing</i> lanjutan seperti <i>smoothing</i> atau transformasi data untuk meningkatkan performa model. Selain itu, objek penelitian masih terbatas pada skala UMKM, sehingga belum mengkaji penerapan ARIMA pada sektor industri atau perusahaan besar.	peramalan penjualan dengan tingkat kesalahan yang rendah. Hasil evaluasi menggunakan MAPE menunjukkan nilai 9,18%, yang menandakan bahwa hasil peramalan cukup akurat dan dapat diandalkan untuk mendukung perencanaan produksi dan pengelolaan persediaan usaha.		Permatasari
Time series analysis and forecasting of China's energy production during Covid-19: statistical models vs machine learning models				
Membandingkan model ARIMA dan LSTM dalam konteks	Belum adanya penerapan model ARIMA dan	Model <i>LightGBM</i> memberikan akurasi tertinggi	2021	Lu, Z., Liu, N., Xie, Y.,

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
peramalan deret waktu. Relevansinya terletak pada fokus terhadap pengujian efektivitas model statistik dan model berbasis <i>machine learning</i> dalam menangani data yang fluktuatif, seperti halnya data FABA yang digunakan dalam penelitian ini. Studi tersebut juga menyoroti perbedaan kemampuan model dalam menghadapi data dengan karakteristik stabil dan tidak stabil, sehingga memberikan dasar empiris untuk memahami kondisi di mana masing-masing model menunjukkan performa optimal.	LSTM pada konteks industri spesifik seperti energi limbah FABA, serta belum adanya pembahasan mendalam mengenai pengaruh parameter model terhadap tingkat akurasi prediksi. Selain itu, penelitian tersebut berfokus pada konteks pandemi COVID-19 yang bersifat temporer, sehingga masih terbuka peluang untuk meneliti performa kedua model dalam kondisi data jangka panjang dan berkelanjutan. Penelitian ini berupaya mengisi	secara keseluruhan dengan nilai MAPE antara 1,79% hingga 4,00%, sementara model ARIMA dan LSTM menunjukkan keunggulan pada tipe data tertentu. ARIMA lebih baik digunakan untuk data yang stabil seperti produksi tenaga air dan gas alam, dengan MAPE terendah mencapai 1,66%, sedangkan LSTM lebih unggul untuk data yang fluktuatif dan kompleks karena kemampuannya mengenali pola temporal serta mempertimbangkan faktor eksternal seperti dampak pandemi COVID-		& Xu, J.

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
	kekosongan tersebut dengan menerapkan ARIMA dan LSTM pada data FABA yang memiliki karakteristik non-linear dan berfluktuasi secara alami.	19. Meskipun LSTM belum mampu melampaui performa <i>LightGBM</i> , model ini tetap menghasilkan prediksi yang solid dengan tingkat kesalahan di bawah 10%.		

Performance Comparison of ARIMA, LSTM, and Prophet Methods in Sales Forecasting				
Menggunakan model ARIMA dan LSTM dalam konteks peramalan deret waktu. Relevansinya terletak pada tujuan untuk membandingkan performa model statistik dan model berbasis jaringan saraf dalam menghasilkan prediksi yang akurat, dengan tambahan metode Prophet sebagai pembanding. Studi ini memberikan	Fokus yang masih terbatas pada data penjualan produk dengan pola yang relatif stabil, sehingga belum menguji performa model pada data dengan tingkat fluktuasi non-linear yang tinggi seperti pada data FABA. Selain itu, penelitian tersebut belum menyoroti aspek optimasi	Hasil penelitian menunjukkan bahwa model ARIMA (1,0,2) memberikan performa terbaik dengan nilai MAPE sebesar 4,548%, MSE sebesar 5,27%, dan RMSE sebesar 47,4139, diikuti oleh LSTM dengan MAPE sebesar 7,3275%, MSE sebesar 6,33%, dan RMSE	2024	Suryawan, I. G. T., Putra, I. K. N., Meliana, P. M., & Sudipa, I. G. I.

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
dasar empiris mengenai efektivitas ARIMA dan LSTM dalam mengolah data yang memiliki pola tren dan fluktuasi tertentu, yang relevan untuk diterapkan pada penelitian ini dalam konteks data FABAs yang juga bersifat deret waktu dan dinamis.	parameter model secara mendalam maupun potensi pengembangan model hibrida untuk meningkatkan akurasi prediksi. Oleh karena itu, penelitian ini berupaya mengisi kekosongan tersebut dengan menerapkan ARIMA dan LSTM pada data FABAs yang memiliki variasi pola dan tingkat ketidakstabilan yang lebih kompleks.	sebesar 56,9761, serta Prophet MAPE sebesar 7,402%, MSE sebesar 8,83%, dan RMSE sebesar 79,5137. Temuan ini menegaskan bahwa ARIMA unggul pada data dengan tren jangka pendek dan tingkat fluktuasi rendah, sedangkan LSTM dan Prophet FABAs cenderung lebih sesuai untuk data kompleks dengan pola non-linear atau musiman jangka panjang.		
Analisis Peramalan Anggaran Belanja Daerah dengan Model ARIMA Dalam Upaya Pelaksanaan Green Economy di Provinsi Jawa Barat				
Menggunakan model ARIMA dalam konteks peramalan data deret waktu. Relevansinya terletak	penelitian tersebut hanya menggunakan model ARIMA dan <i>Exponential</i>	metode ARIMA memiliki akurasi yang lebih unggul dibandingkan <i>Exponential</i>	2024	Firdaya ti, N., Politek nik, A., Bandun

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
pada fokus penelitian terhadap efektivitas ARIMA dalam menghasilkan prediksi yang akurat terhadap data anggaran lingkungan, yang memiliki karakteristik tren dan pola musiman seperti halnya data FABA yang digunakan dalam penelitian ini. Studi ini memberikan dasar empiris mengenai keunggulan model statistik konvensional dalam peramalan jangka menengah dan jangka panjang yang memiliki kestabilan pola.	<i>Smoothing</i> tanpa melibatkan metode <i>deep learning</i> seperti LSTM, sehingga belum mengevaluasi potensi model berbasis jaringan saraf dalam menangani data non-linear dan berfluktuasi. Selain itu, penelitian tersebut berfokus pada konteks anggaran lingkungan, belum pada data berbasis produksi atau energi seperti FABA yang cenderung lebih dinamis. Oleh karena itu, penelitian ini berupaya mengisi celah tersebut dengan membandingkan	<i>Smoothing</i> , dengan nilai MAPE sebesar 19,042%, yang menunjukkan kemampuan ARIMA dalam menghasilkan prediksi dengan kesalahan yang relatif rendah. Temuan ini memperkuat posisi ARIMA sebagai metode statistik yang andal dalam memprediksi data deret waktu lingkungan, khususnya untuk perencanaan anggaran berkelanjutan dalam konteks penerapan <i>green economy</i> di Provinsi Jawa Barat.		g, N., Susilaw aty, R., & Politek nik, H.

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
	performa model ARIMA dan LSTM pada data FABA yang memiliki kompleksitas dan pola ketidakstabilan yang lebih tinggi.			

Smoothing Techniques for Improving COVID-19 Time Series Forecasting Across Countries

sama-sama menekankan peran <i>preprocessing</i> data <i>time series</i> melalui teknik <i>smoothing</i> sebelum dilakukan proses peramalan. Dalam penelitian tersebut, digunakan untuk mengurangi <i>noise</i> dan fluktuasi ekstrem pada data COVID-19 agar model peramalan, khususnya LSTM dan model berbasis <i>machine learning</i> , dapat menangkap pola	penelitian tersebut belum mengeksplorasi penggunaan <i>Gaussian smoothing</i> secara spesifik pada data <i>time series</i> . Selain itu, model statistik klasik seperti ARIMA dijadikan fokus utama dalam pengujian, karena penelitian lebih menitikberatkan pada model <i>machine learning</i>	Hasil penelitian menunjukkan bahwa pemilihan model peramalan merupakan faktor utama yang memengaruhi akurasi, namun penerapan <i>smoothing</i> terbukti mampu meningkatkan stabilitas hasil peramalan dan menurunkan nilai error relatif, khususnya pada model neural	2019	Uliana Zbezhk hovska dan Dmytro Chuma chenko
--	---	---	------	--

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
data dengan lebih stabil. Hal ini sejalan dengan penelitian FABAs yang juga menghadapi karakteristik data fluktuatif dan jumlah observasi terbatas, sehingga membutuhkan teknik <i>smoothing</i> untuk meningkatkan kualitas data sebelum pemodelan ARIMA maupun LSTM.	dan <i>deep learning</i> . Jurnal ini juga menggunakan data epidemiologi, sehingga belum menguji efektivitas <i>smoothing</i> pada konteks data bisnis atau pendapatan, yang memiliki karakteristik dan tujuan peramalan yang berbeda. Dengan demikian, masih terdapat celah penelitian terkait penerapan <i>Gaussian smoothing</i> sebagai <i>preprocessing</i> pada model ARIMA dan LSTM untuk peramalan pendapatan atau data ekonomi.	seperti LSTM dan Temporal <i>Fusion Transformer</i> . Teknik <i>smoothing</i> tertentu, terutama STL dan <i>rolling mean</i> , memberikan performa yang lebih konsisten pada horizon peramalan jangka pendek. Uji statistik ANOVA mengonfirmasi bahwa perbedaan model berpengaruh signifikan terhadap nilai MAPE, sedangkan <i>smoothing</i> berperan sebagai pendukung untuk meningkatkan kualitas <i>input</i> data. Secara keseluruhan, penelitian ini menyimpulkan bahwa <i>smoothing</i> sebagai tahap <i>preprocessing</i>		

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
		penting untuk mempersiapkan data <i>time series</i> yang berisik agar model peramalan dapat bekerja lebih andal, meskipun tidak selalu menjadi faktor dominan dibandingkan pemilihan model itu sendiri.		

Gaussian smoothing and modified histogram normalization methods to improve neural-biomarker interpretations for dyslexia classification mechanism

Artikel ini memiliki keterkaitan yang kuat dengan penelitian yang kamu lakukan karena sama-sama menekankan pentingnya tahap <i>preprocessing</i> data, khususnya <i>Gaussian smoothing</i> , dalam meningkatkan performa model berbasis <i>deep learning</i> . Meskipun konteks penelitian berada pada	Penelitian ini belum menguji penerapan <i>Gaussian smoothing</i> pada model peramalan <i>time series</i> seperti ARIMA LSTM dalam konteks <i>forecasting</i> , belum mengevaluasi pengaruh <i>smoothing</i>	Hasil penelitian menunjukkan bahwa penerapan <i>Gaussian smoothing</i> dan <i>modified histogram normalization</i> pada tahap <i>data preprocessing</i> mampu meningkatkan kualitas data <i>neuroimaging</i> dan kinerja model <i>deep learning</i> secara	2021	Opeye mi Lateef Usman, Ravie Chandr en Muniya ndi, Khairu ddin Omar, Mazlyf arina
--	---	--	------	---

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
bidang kesehatan, pendekatan <i>smoothing</i> yang digunakan relevan untuk data <i>time series</i> bisnis atau pendapatan yang memiliki fluktuasi tinggi dan jumlah data terbatas, seperti pada peramalan pendapatan FABAs.	terhadap metrik kesalahan peramalan seperti MAE, RMSE, atau MAPE. Selain itu, objek penelitian terbatas pada klasifikasi medis, sehingga belum mengeksplorasi implementasi <i>Gaussian smoothing</i> pada data bisnis atau industri.	signifikan. Setelah dilakukan <i>preprocessing</i> , model klasifikasi yang digunakan mampu mencapai tingkat akurasi hingga sekitar 94,7%, disertai dengan penurunan nilai kesalahan klasifikasi dan peningkatan stabilitas fitur <i>biomarker</i> yang diekstraksi. Hal ini menunjukkan bahwa <i>preprocessing</i> berbasis <i>smoothing</i> berperan penting dalam meningkatkan kemampuan model dalam mengenali pola data yang kompleks dan berisik.		Mohamad

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
Comparison of Deep LSTM and Machine Learning Models for Predicting Compressive Strength of Fly Ash/Slag-Based Geopolymer Concrete				
Penelitian terdahulu mengenai prediksi kinerja FABA melalui pendekatan <i>komputasional</i> berbasis <i>machine learning</i> memberikan dasar ilmiah yang penting bagi penelitian ini, terutama dalam konteks <i>forecasting</i> pendapatan dari pemanfaatan limbah FABA. Pemanfaatan algoritma <i>Deep LSTM</i> , <i>ANN</i> , serta model <i>ensemble</i> seperti <i>Bagging</i> dan <i>LSBoost</i> memungkinkan identifikasi pola <i>nonlinier</i> dan tren historis dalam data penjualan, harga, serta volume produksi FABA. Hal ini memiliki keterkaitan	Meskipun penelitian sebelumnya menunjukkan bahwa model <i>Deep LSTM</i> mampu memberikan prediksi yang sangat akurat, terdapat beberapa celah yang perlu diatasi. Data yang digunakan masih terbatas pada eksperimen laboratorium dan variabel fisik FABA, sedangkan penerapan <i>machine learning</i> untuk memprediksi pendapatan atau potensi finansial FABA secara langsung masih jarang diteliti. Selain itu, variasi	Hasil penelitian terdahulu menunjukkan bahwa model <i>Deep LSTM</i> memberikan performa prediksi terbaik dengan nilai R^2 sebesar 0,9819 dan MAPE sebesar 4,56%, diikuti oleh <i>ANN</i> dengan R^2 sebesar 0,9501, serta <i>LSBoost</i> dengan R^2 sebesar 0,9473. Analisis sensitivitas juga menegaskan bahwa faktor paling berpengaruh terhadap kinerja prediksi adalah agregat halus (99,57%), kandungan SiO_2 dan CaO pada slag (99,4%), serta kandungan SiO_2	2025	Kina, C., Tanyildizi, H., Al Bakri Abdullah, M. Razak, R. A., & Imjai, T.

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
langsung dengan penelitian saat ini, yang bertujuan mengembangkan model prediksi pendapatan yang presisi mendukung pengambilan keputusan strategis, perencanaan produksi, dan optimalisasi distribusi FABAs. Pendekatan ini juga sejalan dengan prinsip efisiensi dalam pengelolaan sumber daya dan pengurangan ketergantungan pada metode peramalan tradisional yang bersifat subjektif.	komposisi material, kondisi pasar, dan fluktuasi harga yang memengaruhi pendapatan FABAs belum dianalisis secara menyeluruh. Gap ini menekankan perlunya penelitian yang menggabungkan data historis penjualan FABAs, variabel produksi, serta kondisi pasar untuk menghasilkan model <i>forecasting</i> yang lebih <i>robust</i> dan mampu menangkap dinamika pasar dengan tingkat akurasi tinggi.	pada <i>fly ash</i> (96,8%). Temuan ini menegaskan kemampuan <i>machine learning</i> , khususnya <i>Deep LSTM</i> , dalam menangkap hubungan <i>nonlinier</i> yang kompleks antara variabel <i>input</i> dan <i>output</i> , sekaligus menunjukkan bahwa model ini dapat diaplikasikan secara efektif untuk <i>forecasting</i> pendapatan FABAs, memberikan prediksi yang presisi dan efisien dibandingkan metode konvensional.		

Properties and Optimization Process Using Machine Learning for Recycling of Fly and Bottom Ashes in Fire-Resistant Materials

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
Penelitian tersebut memiliki keterkaitan yang substansial dengan penelitian ini karena sama-sama memanfaatkan teknologi <i>machine learning</i> dalam proses pemodelan dan prediksi berbasis data hasil pembakaran batu bara, khususnya <i>fly ash</i> dan <i>bottom ash</i> (FABA). Pada penelitian terdahulu, penerapan algoritma <i>machine learning</i> seperti <i>Random Forest</i> dan <i>XGBoost</i> digunakan untuk memprediksi serta mengoptimalkan karakteristik fisik dan mekanik material tahan api yang dikembangkan dari limbah pembakaran tersebut. Prinsip serupa diadopsi dalam	Meskipun penelitian sebelumnya telah berhasil menerapkan algoritma <i>machine learning</i> untuk memodelkan hubungan nonlinier antara komposisi abu dan sifat termomekanik material, studi tersebut belum menjangkau dimensi ekonomi dan nilai tambah dari pemanfaatan FABA dalam konteks industri energi. Selain itu, pendekatan regresi yang digunakan masih bersifat statis dan belum mempertimbangkan karakteristik temporal pada data, sehingga belum	Hasil penelitian terdahulu menunjukkan bahwa model berbasis <i>ensemble learning</i>, seperti <i>Random Forest</i> dan <i>XGBoost</i>, mampu menangkap hubungan kompleks antarvariabel dengan akurasi tinggi, mencapai 97% pada klasifikasi <i>Fire Resistance Rating (FRR)</i> dan 99% pada klasifikasi kekuatan tekan material. Temuan ini menunjukkan efektivitas <i>machine learning</i> dalam mengelola data nonlinier dan meningkatkan efisiensi proses berbasis	2025	Guirado, E., Ruiz Martine z, J. D., Campo y, M., & Leiva, C.

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
penelitian ini, namun dengan fokus yang berbeda, yaitu pada aspek ekonomi berupa analisis dan prediksi pendapatan <i>beyond kWh</i> dari pemanfaatan FABA. Dengan demikian, penelitian ini mengembangkan pendekatan prediktif serupa namun dengan orientasi pada data finansial, menggunakan metode <i>time series forecasting</i> seperti <i>Long Short-Term Memory (LSTM)</i> dan <i>AutoRegressive Integrated Moving Average (ARIMA)</i> untuk menganalisis dinamika fluktuasi pendapatan dari pemanfaatan FABA di masa mendatang.	dapat menangkap pola tren jangka panjang yang lazim ditemukan dalam data pendapatan. Oleh karena itu, penelitian ini berupaya mengisi kekosongan tersebut dengan menerapkan model <i>forecasting</i> berbasis deret waktu seperti <i>ARIMA</i> dan <i>LSTM</i>, guna memperoleh pemahaman yang lebih komprehensif mengenai proyeksi pendapatan FABA serta potensi peningkatannya dari setiap PLTU dan mitra terkait.	data eksperimen. Temuan tersebut memperkuat dasar metodologis penelitian ini yang berorientasi pada penerapan <i>machine learning</i> untuk tujuan prediksi ekonomi, di mana algoritma <i>LSTM</i> dan <i>ARIMA</i> dimanfaatkan untuk mengidentifikasi pola perubahan pendapatan dari pemanfaatan FABA secara temporal. Dengan demikian, penelitian ini tidak hanya berkontribusi terhadap pengembangan pendekatan <i>data-driven decision making</i> di bidang energi, tetapi juga memberikan perspektif baru		

Keterkaitan Penelitian	Gap Penelitian	Hasil Penelitian	Tahun	Penulis
		dalam optimalisasi nilai ekonomi dari limbah pembakaran batu bara secara berkelanjutan.		

Berdasarkan penelitian terdahulu yang telah diuraikan pada **Tabel 2.1**, dapat disimpulkan bahwa model *Autoregressive Integrated Moving Average* (ARIMA) dan *Long Short-Term Memory* (LSTM) merupakan dua pendekatan utama yang banyak digunakan dalam peramalan data *time series*. Secara umum, ARIMA menunjukkan kinerja yang unggul pada data dengan pola yang stabil, linier, dan tidak memiliki termometer ekstrem, sedangkan LSTM lebih efektif dalam menangani data yang bersifat non-linier, dinamis, serta mengandung variasi musiman. Beberapa penelitian seperti yang dilakukan oleh Taslim & Murwantara (2024) serta Adwin Adam et al. (2024) menyatakan bahwa ARIMA lebih efisien pada *dataset* besar dan stabil, sementara penelitian oleh Gupta & Jaiswal, 2024 dan Lu et al. 2021 menunjukkan bahwa LSTM mampu memberikan akurasi yang lebih tinggi dalam data dengan kompleksitas temporal yang tinggi.

Selain itu, penelitian terdahulu juga menunjukkan bahwa performa kedua model sangat dipengaruhi oleh ukuran data, parameter model, serta karakteristik kinerja yang terkandung dalam *dataset*. Model LSTM, khususnya dalam bentuk *multivariat* seperti yang dikemukakan oleh (Vithisoontorn, 2021), terbukti lebih adaptif terhadap data dengan hubungan variabel yang saling bergantung, sedangkan ARIMA tetap menjadi pilihan yang Andal untuk peramalan jangka pendek yang stabil dan mudah diinterpretasikan. Dari sisi efisiensi waktu, ARIMA umumnya lebih ringan dalam proses komputasi dibandingkan LSTM yang memerlukan waktu pelatihan lebih lama.

Penelitian ini mengacu pada temuan-temuan tersebut dengan tujuan untuk menganalisis dan membandingkan kinerja model ARIMA dan LSTM dalam memprediksi data FABA yang bersifat fluktuatif dan non-linear. Tidak seperti penelitian sebelumnya yang sebagian besar fokus pada sektor keuangan, penjualan, atau energi umum, penelitian ini memperluas konteks penerapan kedua model tersebut pada bidang pengelolaan limbah industri energi , khususnya dalam memprediksi potensi pemanfaatan FABA sebagai bahan konstruksi ramah lingkungan. Dengan

demikian, penelitian ini diharapkan dapat memberikan kontribusi nyata terhadap pengembangan model penyiaran berbasis data lingkungan serta mendukung upaya efisiensi dan penghentian dalam pengelolaan FABA di Indonesia.

2.2 Landasan Teori

Landasan teoritis menggambarkan berbagai teori yang digunakan dalam penelitian ini. Teori disusun dan diperoleh dari berbagai buku, jurnal, dan artikel terpercaya yang membahas, mendukung, dan menjadi referensi penelitian ini.

2.2.1 *Monitoring Of Power Plan (MAPP) POWER*

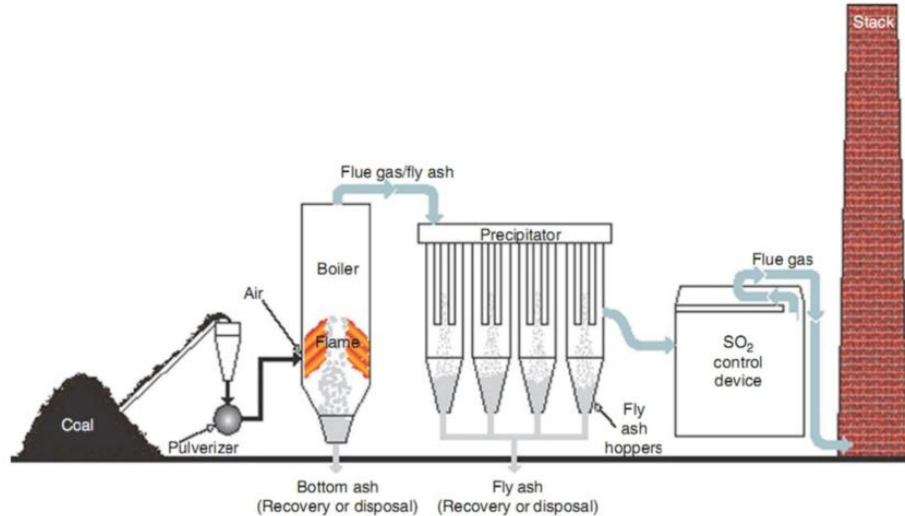


Gambar 2. 1 Aplikasi MAPP Power

Monitoring Of Power Plants (MAPP) Power adalah aplikasi yang digunakan untuk memantau dan mencatat data operasional serta kinerja pembangkit listrik secara langsung. Aplikasi ini berfungsi sebagai pusat pemantauan yang membantu pengelolaan aset, kondisi utama pembangkit, dan indikator kinerja penting. Melalui fitur pemantauan dan pelaporan yang tersedia, *MAPP Power* membantu perusahaan mengelola aset dengan lebih efisien, mengurangi gangguan operasional yang tidak direncanakan, memperpanjang masa pakai aset, serta menekan biaya operasional (pln pusat, 2024).

Salah satu fitur utama *MAPP Power* adalah pemantauan data *Fly Ash* dan *Bottom Ash* (FABA) yang memungkinkan pengguna melihat pendapatan dan volume FABA yang terjual dalam periode tertentu secara langsung. Fitur ini membantu perusahaan dalam mengelola pemanfaatan FABA secara lebih efisien serta mempermudah pengambilan keputusan yang berkaitan dengan pengelolaan limbah pembangkit listrik secara berkelanjutan (pln pusat, 2024).

2.2.2 *Fly Ash and Bottom Ash (FABA)*



Gambar 2. 2 Proses Fly Ash dan Bottom Ash

FABA merupakan residu padat sisa hasil pembakaran batubara pada PLTU. Residu ini terdiri atas partikel halus abu terbang (*Fly Ash*), serta abu dasar (*Bottom Ash*) (Yunita et al., 2017) seperti pada Gambar 2.2.

Di Indonesia, *Fly Ash* dan *Bottom Ash* (FABA) merupakan limbah hasil pembakaran batu bara pada PLTU yang dikategorikan sebagai limbah B3 sesuai PP No. 101 Tahun 2014, sehingga pengelolaannya wajib mendapat izin resmi. Karena volumenya yang besar dan berpotensi menimbulkan masalah lingkungan, berbagai upaya dilakukan untuk memanfaatkannya, salah satunya melalui penerapan FABA dalam berbagai bidang industri (Besari et al., 2022).

2.2.2.1 *Fly Ash (FA)*



Gambar 2. 3 Fly Ash

Fly Ash merupakan produk sampingan dari proses pembakaran batubara pada Pembangkit Listrik Tenaga Uap (PLTU), di mana partikel-partikel halus seperti pada Gambar 2.3 yang terbawa oleh gas buang saat batu bara dibakar di *boiler* dan selanjutnya ditangkap menggunakan sistem *electrostatic precipitator*. Berdasarkan analisis kimia, *Fly Ash* mengandung SiO_2 (56,44%), Al_2O_3 (31,31%), dan Fe_2O_3 (0,51%) dengan CaO rendah (0,78%). Komposisi ini mengklasifikasikan *Fly Ash* sebagai tipe F menurut ASTM C618-12a, yang memiliki sifat *pozzolanik* tinggi dan reaktif terhadap larutan basa. Oleh karena itu, *Fly Ash* berpotensi besar dimanfaatkan sebagai bahan dasar material seperti *geopolimer* dan *zeolit*, serta dapat meningkatkan kekuatan dan ketahanan material berbasis semen. (Alyatikah et al., n.d.).

2.2.2.2 *Bottom Ash (BA)*



Gambar 2. 4 Bottom Ash

Abu dasar (*bottom ash*) merupakan residu padat hasil pembakaran batu bara yang mengendap di bagian dasar ruang bakar (*boiler*). Secara fisik, *bottom ash* memiliki ukuran partikel yang lebih besar serta bentuk yang tidak beraturan dibandingkan *fly ash*, dengan karakteristik menyerupai pasir kasar seperti pada Gambar 2.4. Berdasarkan analisis kimia, *bottom ash* mengandung SiO_2 (66,66%), Al_2O_3 (17,09%), Fe_2O_3 (0,31%), dan CaO (5,40%). Kandungan silika yang tinggi menunjukkan potensi pemanfaatannya dalam pembuatan material konstruksi, sementara kadar kalsium yang lebih besar dibandingkan *fly ash* mengindikasikan sifat yang mendekati tipe C. Meskipun memiliki reaktivitas kimia lebih rendah, *bottom ash* tetap berperan dalam meningkatkan kekuatan mekanik dan stabilitas material berbasis semen maupun *geopolymer* (Alyatikah et al., n.d.).

2.2.3 Data Mining

Data mining merupakan proses terstruktur untuk mengekstraksi pengetahuan dari data melalui tahapan pemahaman bisnis, pemahaman data, persiapan data, pemodelan, evaluasi, dan implementasi, sehingga hasil analisis dapat digunakan sebagai dasar pengambilan keputusan berbasis data (Yessy Asri et al., 2023). Data tersebut dapat berasal dari berbagai sumber seperti sensor *Internet of Things* (IoT), *smart meter*, maupun catatan historis sistem energi. Melalui penerapan metode analisis seperti *classification*, *clustering*, *association rule mining*, dan deteksi anomali, data mentah yang awalnya tidak terstruktur diubah menjadi informasi bermakna (Munawaroh, 2023).

2.2.3.1 Classification (Klasifikasi)

Classification merupakan metode yang berfungsi untuk mengelompokkan data ke dalam kelas atau kategori yang telah ditentukan sebelumnya berdasarkan atribut tertentu. Model yang sering digunakan dalam *classification* adalah *decision tree* dan *random forest*, karena keduanya mampu memberikan hasil interpretasi yang jelas terhadap pola pengambilan keputusan. *classification* membantu sistem dalam mengenali pola historis dan menghasilkan prediksi akurat terhadap data baru (Munawaroh, 2023).

2.2.3.2 Clustering (Klustering)

Clustering merupakan teknik eksplorasi data yang digunakan untuk mengelompokkan data berdasarkan kemiripan pola tanpa memerlukan label awal (*unsupervised learning*). Metode ini membantu mengenali karakteristik utama data sebelum dilakukan pemilihan fitur dan pemodelan lanjutan. Dalam konteks *Decision Support System*, *clustering* dimanfaatkan untuk menemukan pola alami dalam data, seperti pengelompokan perilaku pengguna atau segmentasi pelanggan, dengan algoritma yang umum digunakan antara lain *K-means* dan *Hierarchical Clustering* (Munawaroh, 2023).

2.2.3.3 Regression (Regresi)

Regresi linear adalah metode statistik yang digunakan untuk menggambarkan hubungan linier antara variabel respon dan satu atau lebih variabel penjelas. Metode ini mencari garis terbaik yang mendekati data menggunakan pendekatan *least squares* dan banyak digunakan dalam *machine learning* untuk memperkirakan nilai numerik berdasarkan data historis (Jatnika et al.,

n.d.). Koefisien menunjukkan perubahan nilai respon akibat perubahan variabel prediktor, dengan asumsi variabel lain tetap. Jika hubungan antar variabel tidak linier, maka dapat digunakan pendekatan lain seperti transformasi variabel, regresi polinomial, atau regresi logistik untuk variabel respon biner (Zhang et al., 2025).

2.2.3.4 *Random Forest*

Random Forest merupakan salah satu algoritma *ensemble learning* yang bekerja dengan membangun sejumlah pohon keputusan (*decision trees*) secara acak, kemudian menggabungkan hasil prediksi dari masing-masing pohon untuk menghasilkan keputusan akhir yang lebih akurat dan stabil. Konsep dasar dari metode ini adalah bahwa kombinasi beberapa model sederhana dapat menghasilkan performa prediksi yang lebih baik dibandingkan satu model tunggal. Proses pembentukan *Random Forest* melibatkan dua elemen utama, yaitu *bootstrap sampling* (atau *bagging*) untuk membentuk *subset data* pelatihan secara acak, dan pemilihan acak *subset fitur* pada setiap simpul keputusan guna mengurangi korelasi antar pohon. Pendekatan ini efektif dalam mengurangi risiko *overfitting*, meningkatkan kemampuan generalisasi model, serta mempertahankan akurasi tinggi pada data berukuran besar atau berdimensi tinggi. Karena keandalannya dalam menangani data kompleks dan *nonlinier*, *Random Forest* banyak digunakan dalam berbagai bidang seperti *klasifikasi*, *regresi*, dan analisis fitur dalam penelitian ilmiah maupun industri (Liu et al., 2021).

2.2.3.5 *Association Rule Mining*

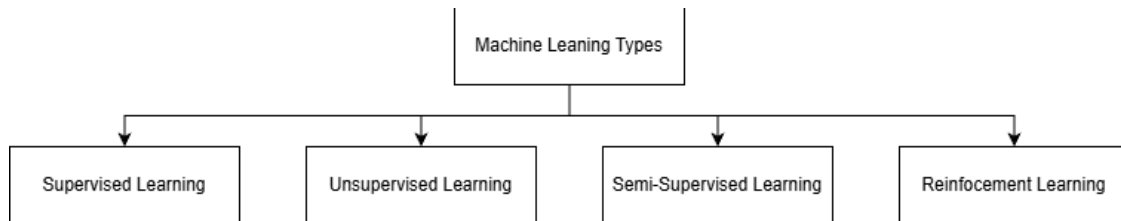
Association rule mining bertujuan menemukan hubungan atau keterkaitan antara variabel dalam *dataset* besar. Teknik ini dikenal luas melalui penerapan *market basket analysis*, di mana sistem mengidentifikasi produk-produk yang sering dibeli bersama. Dalam DSS, *association rule mining* digunakan untuk menghasilkan rekomendasi otomatis berdasarkan pola historis transaksi. Misalnya, sistem dapat memprediksi kemungkinan suatu kejadian berdasarkan frekuensi kemunculan item lain, sehingga membantu organisasi dalam pengambilan keputusan berbasis prediksi perilaku (Munawaroh, 2023).

2.2.3.6 *Anomaly Detection (Deteksi Anomali)*

Detection Anomaly berfokus pada identifikasi data yang menyimpang dari pola umum. Teknik ini penting dalam sistem yang membutuhkan keamanan tinggi, seperti deteksi penipuan finansial atau serangan siber. *Anomaly detection* membantu mengenali kondisi tidak biasa yang

mungkin menandakan risiko atau peluang baru. Dengan menggunakan pendekatan statistik dan algoritma *machine learning*, sistem dapat memantau data secara *real-time* untuk menemukan penyimpangan yang signifikan dari tren normal (Munawaroh, 2023).

2.2.4 *Machine Learning*



Gambar 2. 5 Machine Learning Types

Machine learning (ML) merupakan cabang kecerdasan buatan yang menggunakan algoritma untuk mempelajari pola dari data historis dan memanfaatkannya dalam menghasilkan keputusan atau hasil pada data baru. Pendekatan ini mencakup beberapa paradigma utama, seperti *supervised learning*, *unsupervised learning*, *semi-supervised learning*, dan *reinforcement learning*, dengan *deep learning* sebagai bagian yang mampu menangani data dalam jumlah besar seperti pada Gambar 2.5 sebagai *subset* yang memfasilitasi analisis skala besar (Razzaq & Shah, 2025). Secara umum, ML digunakan untuk berbagai tugas seperti *classification*, *regression*, dan *clustering*, melalui proses pelatihan menggunakan data historis yang kemudian diterapkan pada data baru. Dengan kemampuannya dalam menganalisis data dan menghasilkan peramalan, *machine learning* banyak dimanfaatkan sebagai dasar pengambilan keputusan berbasis data di berbagai bidang (Palupiningsih & Prayitno, 2020). Dengan demikian, ML menjadi fondasi utama bagi pengembangan sistem berbasis data yang mampu melakukan pengambilan keputusan secara cerdas dan tepat waktu.

2.2.4.1 *Supervised Learning*

Supervised learning merupakan pendekatan *machine learning* di mana model dilatih menggunakan data berlabel untuk mempelajari hubungan antara masukan dan keluaran. Proses pelatihan dilakukan dengan menyesuaikan bobot model agar kesalahan (*loss*) semakin kecil, sehingga hasil yang diperoleh lebih akurat. Contoh metode *supervised learning* meliputi *Deep Neural Network* (DNN), *Convolutional Neural Network* (CNN), dan *Recurrent Neural Network* (RNN).

2.2.4.2 *Unsupervised Learning*

Unsupervised learning merupakan pendekatan *machine learning* yang tidak menggunakan data berlabel dan bertujuan untuk menemukan pola atau struktur tersembunyi dalam data. Metode ini digunakan untuk mengelompokkan atau merepresentasikan data tanpa informasi keluaran yang jelas. Contoh model *unsupervised learning* antara lain *Autoencoder* (AE), *Restricted Boltzmann Machine* (RBM), *Deep Belief Network* (DBN), dan *Generative Adversarial Network* (GAN).

2.2.4.3 *Semi-Supervised Learning*

Reinforcement learning adalah metode pembelajaran mesin di mana agen belajar melalui interaksi langsung dengan lingkungan untuk memaksimalkan *reward* kumulatif yang diperoleh. Pendekatan ini tidak memerlukan model matematis eksplisit dari sistem, melainkan berfokus pada proses *trial and error* berdasarkan konsep *state* (keadaan), *action* (tindakan), dan *reward* (imbalan), sehingga agen dapat menemukan strategi optimal dalam pengambilan keputusan.

2.2.4.4 *Reinforcement Learning*

Semi-supervised learning merupakan pendekatan hibrida yang menggabungkan data berlabel dan tidak berlabel untuk meningkatkan performa model, terutama ketika data berlabel terbatas. Metode ini memadukan keunggulan dari supervised dan unsupervised learning atau bahkan reinforcement learning, sehingga model dapat mempelajari representasi yang lebih akurat dan generalisasi yang lebih baik terhadap data baru.

2.2.5 *Deep Learning*

Deep Learning (DL) merupakan pengembangan dari *Artificial Neural Network* (ANN) yang menggunakan banyak lapisan untuk mempelajari pola data. Metode ini memungkinkan sistem mengenali pola yang lebih kompleks langsung dari data dalam jumlah besar. Sejak diperkenalkan pada tahun 2006, *deep learning* telah digunakan di berbagai bidang, seperti pengenalan objek, diagnosis medis, analisis media sosial, dan sistem rekomendasi (Razzaq & Shah, 2025).

2.2.6 *Dataset dan Processing Data*

Dataset merupakan kumpulan data terstruktur maupun tidak terstruktur yang berfungsi sebagai dasar dalam proses pembelajaran mesin (*machine learning*). *Dataset* biasanya terdiri atas

dua komponen utama, yaitu fitur (*features*) sebagai variabel *input* dan *label (target)* sebagai variabel output yang ingin diprediksi. Kualitas *dataset* sangat menentukan performa model karena data yang tidak bersih atau tidak seimbang dapat menurunkan akurasi prediksi. Dalam konteks penelitian energi atau pemanfaatan FABA, *dataset* dapat mencakup data historis produksi, volume pemanfaatan, hingga pendapatan yang dihasilkan oleh setiap PLTU (Fan et al., 2021).

Sementara itu, *data processing* merupakan tahapan krusial yang dilakukan sebelum model *machine learning* dibangun. Tahap *preprocessing* bagian penting dalam pengolahan data karena berperan dalam meningkatkan kualitas input dan kinerja model *machine learning* dengan mengurangi kompleksitas serta mempertahankan informasi yang relevan (Yosrita et al., 2021). Tujuan utamanya adalah untuk meningkatkan kualitas data agar model dapat belajar secara optimal. Tahap-tahap umum dalam *data processing* mencakup:

- 1) *Data Cleaning* (Pembersihan Data), untuk menghapus data duplikat, memperbaiki nilai hilang (*missing values*) dan mengatasi *outlier* yang dapat mengganggu hasil analisis.
- 2) *Data Integration* (Integrasi Data), untuk menghubungkan berbagai sumber data menjadi satu kesatuan yang konsisten dan komperhensif.
- 3) *Data Transformation* (Transformasi Data), yaitu tahap mengubah format atau skala data, melalui *normalization*, *standardization* atau *encoding* untuk menyesuaikan dengan algoritma yang digunakan.

Tahapan ini memastikan bahwa data yang akan dimasukkan ke model berada dalam kondisi optimal, bebas dari bias, dan siap digunakan dalam tahap *modeling*. Proses *data preprocessing* yang tepat terbukti mampu meningkatkan kinerja model *machine learning* secara signifikan hingga 20–30% dibandingkan dengan data mentah yang belum diolah (Fan et al., 2021).

2.2.7 Cross-Industry Standard Process for Data Mining (CRISP-DM)



Gambar 2. 6 The CRISP-DM life cycle

Model CRISP-DM (*Cross-Industry Standard Process for Data Mining*) merupakan salah satu kerangka kerja yang paling banyak digunakan dalam penelitian berbasis *data mining* karena memiliki struktur yang sistematis dan fleksibel. Model ini terdiri atas enam tahapan utama yang saling berhubungan dan berulang secara iteratif, yaitu: *business understanding*, *data understanding*, *data preparation*, *modeling*, *evaluation*, dan *deployment* seperti pada Gambar 2.6. Tujuan utamanya adalah untuk memastikan proses penggalian data dilakukan secara sistematis mulai dari pemahaman masalah hingga implementasi hasil model secara nyata (Schröer et al., 2021).

- 1) *Business Understanding*, yaitu proses untuk memahami tujuan utama penelitian, mendefinisikan permasalahan yang ingin diselesaikan, dan merumuskan sasaran analisis yang ingin dicapai. Tahapan ini berperan penting untuk memastikan seluruh proses *data mining* berjalan sesuai dengan konteks dan kebutuhan penelitian.
- 2) *Data Understanding*, yang mencakup kegiatan pengumpulan data, eksplorasi awal, serta analisis deskriptif terhadap karakteristik data. Pada tahap ini, peneliti mengidentifikasi struktur data, hubungan antarvariabel, tren, pola, serta potensi masalah seperti data hilang, anomali, atau ketidakkonsistenan yang dapat memengaruhi kualitas hasil analisis.
- 3) *Data Preparation*, yaitu proses menyiapkan data agar siap digunakan dalam pemodelan. Langkah-langkah utamanya meliputi pembersihan data (*data cleaning*), transformasi variabel

(*data transformation*), penggabungan data (*data agregation*), normalisasi, serta pemilihan fitur yang relevan (*feature selection*). Tujuan utama dari tahap ini adalah menghasilkan *dataset* berkualitas tinggi yang representatif dan sesuai dengan kebutuhan model yang akan dibangun.

- 4) *Modeling*, di mana peneliti membangun model analisis menggunakan algoritma statistik atau *machine learning*. Pada tahap ini, dilakukan penyesuaian parameter model serta proses pelatihan dan pengujian data untuk memperoleh hasil yang optimal.
- 5) *Evaluation*, yaitu proses untuk menilai kinerja dan keandalan model yang telah dibangun. Evaluasi dilakukan dengan menggunakan berbagai metrik seperti akurasi, *Mean Absolute Error (MAE)*, *Root Mean Square Error (RMSE)*, atau *Mean Absolute Percentage Error (MAPE)*. Selain menilai ketepatan model, tahap ini juga memastikan bahwa hasil yang diperoleh sesuai dengan tujuan awal penelitian.
- 6) *Deployment*, yang berfokus pada penerapan model dalam konteks nyata. Pada tahap ini, hasil analisis diimplementasikan untuk mendukung proses pengambilan keputusan, perencanaan strategi, atau otomatisasi sistem berbasis data. *Deployment* juga dapat mencakup pembuatan laporan, *dashboard* analitik, atau integrasi model ke dalam sistem operasional agar hasil penelitian dapat dimanfaatkan secara berkelanjutan (Schröer et al., 2021).

2.2.8 *Gaussian Smoothing*

Gaussian smoothing merupakan teknik *data preprocessing* yang digunakan untuk mengurangi *noise* dan fluktuasi ekstrem pada data dengan memanfaatkan distribusi *Gaussian*, sehingga data menjadi lebih halus dan stabil. Teknik ini bekerja dengan menghitung nilai suatu titik berdasarkan nilai di sekitarnya, di mana data yang lebih dekat diberi bobot lebih besar dibandingkan data yang lebih jauh, sehingga gangguan acak dapat ditekan tanpa menghilangkan pola utama (Zbezhkhovska & Chumachenko, 2025). *Gaussian smoothing* efektif dalam meningkatkan kualitas data yang mengandung *noise* tinggi, sehingga membantu model *deep learning* dalam mengenali pola dengan lebih baik. Oleh karena itu, *Gaussian smoothing* cocok digunakan pada data *time series* yang tidak stabil, seperti data operasional, data sensor, atau data pendapatan yang mengalami naik dan turun tajam (Usman et al., 2021)

2.2.9 *Box-Jenkins Models*

Metode *Box-Jenkins*, yang dikembangkan oleh George Box dan Gwilym Jenkins pada tahun 1976, merupakan pendekatan sistematis untuk membangun model deret waktu, terutama

Autoregressive Integrated Moving Average (ARIMA), melalui tiga tahap utama, yaitu identifikasi model, estimasi parameter, dan pengujian kecocokan model (Banaś & Utnik-Banaś, 2021).

2.2.9.1 Autoregressif (AR)

Model *Autoregressive* (AR) merupakan dasar dari model Box–Jenkins. Ia memprediksi nilai saat ini dari suatu deret waktu berdasarkan kombinasi linear dari nilai-nilai masa lalunya. Secara umum ditulis sebagai:

$$y_t = c + \sum_{i=1}^p \phi_i y_{t-i} + \varepsilon_t \quad (2.1)$$

Dengan :

- y_t : nilai variabel pada waktu ke- t
- c : konstanta (intersep)
- ϕ_i : koefisien autoregresif pada lag ke- i
- y_{t-i} : nilai variabel pada waktu sebelumnya (lag ke- i)
- p : orde model AR (jumlah lag)
- ε_t : error acak pada waktu ke- t

2.2.9.2 Moving Average (MA)

Model *Moving Average* (MA) memprediksi nilai saat ini berdasarkan kombinasi linear dari kesalahan prediksi masa lalu.

$$y_t = \mu + \sum_{i=1}^q \theta_i \varepsilon_{t-i} + \varepsilon_t \quad (2.2)$$

Dengan :

- μ : rata-rata dari proses deret waktu
- θ_i : koefisien moving average pada lag ke- i
- ε_{t-i} : error acak pada periode sebelumnya
- q : orde model MA

2.2.9.3 Autoregressive Integrated Moving Average (ARIMA)

Model ARIMA dikembangkan oleh Box dan Jenkins. Ia menggabungkan elemen AR dan MA, serta menambahkan proses *differencing* (I) untuk membuat data menjadi stasioner. Dalam jurnal dijelaskan bahwa model ini memerlukan tiga parameter: p (orde AR), d (orde *differencing*), dan q (orde MA). Rumus dasar ARIMA (sebelum *differencing*) dengan operator p dan q adalah :

$$\Phi(L)^p \Delta^d y_t = \phi(L)^q \varepsilon_t \quad (2.3)$$

Rumus dasar ARIMA (setelah *differencing*) dengan operator p , b dan q adalah :

$$\begin{aligned} (1 - \phi_1 L - \phi_2 L^2 - \dots - \phi_p L^p)(1 - L)^d y_t \\ = (1 + \theta_1 L + \theta_2 L^2 + \dots + \theta_q L^q) \varepsilon_t \end{aligned} \quad (2.4)$$

Dengan :

L : operator lag, di mana $Ly_t = y_{t-1}$

$\Phi(L)^p$: operator AR orde- p

$\theta(L)^q$: operator MA orde- q

$(1 - L)^d$: operator diferensiasi orde- d

d : operator diferensiasi (*degree of differencing*)

2.2.10 Durbin Levinson

Algoritma *Durbin-Levinson* merupakan sebuah metode perhitungan yang digunakan untuk mendapatkan nilai korelasi parsial (PACF) dengan memanfaatkan informasi dari nilai *autokorelasi* (ACF) yang sudah ada. Sederhananya, algoritma ini bekerja seperti sebuah sistem berantai (*rekursif*), di mana hasil perhitungan pada tahap sebelumnya digunakan kembali untuk mendapatkan hasil di tahap berikutnya.

Perhitungan dimulai dari tahap yang paling sederhana, yaitu pada jarak waktu pertama atau *lag* 1. Pada tahap ini, nilai korelasi parsial (PACF) dianggap sama dengan nilai *autokorelasinya* (ACF) karena belum ada pengaruh dari data di antara mereka, sehingga bisa diselesaikan dengan persamaan :

$$\phi_{11} = \rho_1 \quad (2.5)$$

Dengan :

ϕ_{11} : Nilai PACF untuk *lag* pertama.

ρ_1 : Nilai ACF untuk *lag* pertama.

Setelah langkah awal diketahui, kita dapat mencari nilai korelasi parsial untuk jarak waktu selanjutnya. Rumus berikut digunakan untuk menghitung nilai PACF secara bertahap dengan cara "membersihkan" pengaruh korelasi dari masa lalu yang sudah dihitung sebelumnya:

$$\phi_{kk} = \frac{\rho_k - \sum_{j=1}^{k-1} \phi_{k-1,j} \rho_{k-j}}{1 - \sum_{j=1}^{k-1} \phi_{k-1,j} \rho_j} \quad (2.6)$$

Dengan :

ϕ_{kk} : Koefisien PACF pada *lag* ke- k yang sedang dihitung.

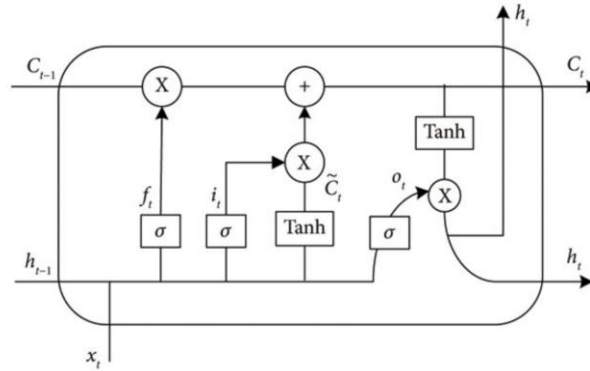
ρ_k : Nilai ACF pada *lag* ke- k

$\sum_{j=1}^{k-1}$: Simbol sigma yang menunjukkan penjumlahan hasil penjumlahan dari $j = 1$ hingga $k - 1$

$\phi_{k-1,j}$: Nilai koefisien yang telah didapatkan dari perhitungan pada tahap (iterasi) sebelumnya.

2.2.11 Model Long Short-Term Memory (LSTM)

Long Short-Term Memory (LSTM) adalah salah satu arsitektur jaringan saraf tiruan yang termasuk dalam keluarga *Recurrent Neural Network* (RNN). Long Short-Term Memory (LSTM) merupakan model *deep learning* yang efektif dalam mempelajari ketergantungan jangka panjang pada data deret waktu dan mampu menangkap pola *nonlinier* (Ainul Idham & Efy Yosrita, 2025). Arsitektur ini dirancang khusus untuk memproses dan mempelajari data berurutan, seperti teks, sinyal suara, atau waktu henti. LSTM bekerja dengan menambahkan mekanisme gerbang (*gating mekanisme*) yang mengatur informasi aliran, sehingga jaringan bisa memilih informasi mana yang harus disimpan, *diupdate*, atau dilupakan seiring berjalannya waktu. Berkat mekanisme ini, LSTM mampu mempertahankan konteks temporal jangka panjang tanpa mengalami masalah ketidakstabilan gradien selama pelatihan. Oleh karena itu, LSTM sering digunakan dalam berbagai aplikasi, seperti prediksi deret waktu, pengenalan ucapan, penerjemahan bahasa alami, dan analisis data berdasarkan waktu lainnya (Ghojogh & Ghodsi, 2023).



Gambar 2. 7 Diagram Skematik LSTM

Pada Gambar 2. 7 menunjukkan arsitektur dasar dari sel *Long Short-Term Memory* (LSTM), yaitu jenis jaringan saraf tiruan yang mampu mengolah data berurutan dengan mengingat informasi jangka panjang. Struktur LSTM dilengkapi dengan tiga gerbang utama yang mengatur aliran informasi, yaitu *input gate* yang menentukan informasi baru mana yang akan disimpan, *forget gate* yang memutuskan informasi lama mana yang perlu dihapus, dan *output gate* yang mengatur informasi mana yang akan diteruskan ke langkah waktu berikutnya (Ghojogh & Ghodsi, 2023).

Secara matematis, proses kerja ketiga gerbang pada arsitektur LSTM dapat dijelaskan melalui serangkaian persamaan yang merepresentasikan mekanisme pembaruan memori dan keluaran jaringan pada setiap langkah waktu t .

1 Input Gate

Gate ini mengatur seberapa besar informasi baru dari *input* x_t masuk ke memori.

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f) \quad (2.7)$$

Dengan :

f_t : *Forget gate*, menentukan informasi mana yang akan dibuang

σ : Fungsi aktivasi sigmoid

W_f : Bobot pada *input gate*

$[h_{t-1}, x_t]$: Konkatenasi antara *output* sebelumnya dan input saat ini

b_f : Bias pada *input gate*

2 Forget Gate

Gate ini memutuskan informasi lama mana yang akan dihapus dari sel memori.

$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i) \quad (2.8)$$

Dengan :

i_t : *input gate*, menentukan informasi

W_i : Bobot pada *forget gate*

b_i : Bias pada *forget gate*

3 New Memory Cell (Candidate Cell)

Ini adalah informasi baru yang akan ditambahkan ke memori jika disetujui oleh *input gate*.

$$\tilde{C}_t = \tanh(W_c \cdot [h_{t-1}, x_t] + b_c) \quad (2.9)$$

Dengan :

$\tilde{C}_t$: Kandidat memori baru, hasil dari aktivasi tanh.

$\tanh$: Fungsi aktivasi sigmoid

W_c : Bobot

b_c : Bias

4 Final Memory Update (Cell State)

Rumus ini menunjukkan kombinasi antara informasi lama dan baru yang dipilih.

$$C_t = f_t \odot C_{t-1} + i_t \odot \tilde{C}_t \quad (2.10)$$

Dengan :

C_t : Sel memori (*cell state*) yang menyimpan informasi jangka panjang

$\odot$: Operasi perkalian elemen per elemen (*Hadamard product*)

C_{t-1} : Memori jangka panjang dari waktu sebelumnya

5 Output Gate

Gate ini menentukan bagian mana dari memori yang akan diteruskan ke output.

$$o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o) \quad (2.11)$$

Dengan :

o_t : *Output gate*, mengontrol informasi yang diteruskan ke *output*.

W_o : Bobot pada *Output gate*

b_o : Bias pada *forget gate*

6 Hidden State (Output)

Output akhir dari sel LSTM, yang juga digunakan sebagai input ke langkah waktu berikutnya.

$$h_t = o_t \odot \tanh(C_t) \quad (2.12)$$

Dengan :

h_t : *Hidden state*, keluaran akhir LSTM pada waktu t .

2.2.12 Evaluasi Model dan Validasi

Penilaian model dan validasi bertujuan untuk mengukur tingkat akurasi model prediksi terhadap data aktual. Pengimplementasian evaluasi MAE, RMSE, dan MAPE bertujuan untuk menilai sejauh mana penyimpangan hasil prediksi dari nilai sebenarnya (Manandhar et al., 2024).

2.2.12.1 Mean Absolute Percentage Error (MAPE)

MAPE digunakan untuk menghitung rata-rata kesalahan absolut dalam bentuk persentase terhadap nilai aktual. Rumus yang digunakan adalah:

$$MAPE = \frac{100\%}{n} \sum_{t=1}^n \left| \frac{y_t - \hat{y}_t}{y_t} \right| \quad (2.13)$$

Dengan :

n : Jumlah total data

y_t : nilai aktual (observasi sebenarnya)

$\hat{y}_t$: nilai prediksi dari model

Evaluasi hasil peramalan dilakukan menggunakan MAPE untuk mengukur besarnya kesalahan relatif antara nilai prediksi dan nilai aktual, sehingga dapat digunakan sebagai dasar

pemilihan metode peramalan yang paling akurat (Sudirman et al., 2020). Dalam jurnal ini, MAPE digunakan untuk memberikan gambaran umum tentang seberapa besar deviasi relatif model terhadap nilai sebenarnya.

2.2.12.2 Root Mean Square Error (RMSE)

RMSE mengukur besarnya kesalahan dengan memberi penalti lebih besar pada kesalahan yang ekstrem, karena menghitung rata-rata kuadrat selisih antara nilai aktual dan nilai prediksi, lalu diakarkan. Rumus yang digunakan adalah:

$$RMSE = \sqrt{\frac{1}{n} \sum_{t=1}^n (y_t - \hat{y}_t)^2} \quad (2.14)$$

RMSE sangat berguna untuk menilai stabilitas model, karena semakin kecil nilai RMSE maka semakin baik performa model. Dalam jurnal ini, RMSE digunakan untuk menilai ketepatan prediksi beban listrik dengan mempertimbangkan besarnya *error* dalam satuan yang sama dengan data asli.

2.2.12.3 Mean Absolute Error (MAE)

MAE menghitung rata-rata dari selisih absolut antara nilai aktual dan hasil prediksi, dengan rumus:

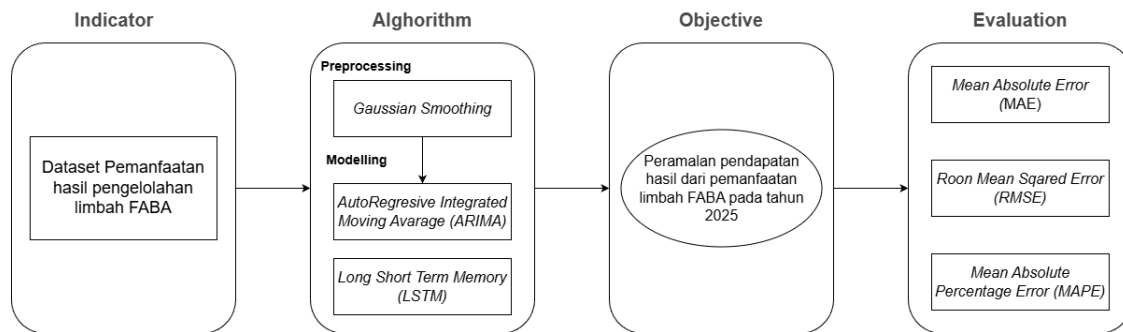
$$MAE = \frac{1}{n} \sum_{t=1}^n |y_t - \hat{y}_t| \quad (2.15)$$

MAE dianggap sebagai metrik yang paling mudah diinterpretasikan karena menunjukkan rata-rata kesalahan dalam satuan yang sama dengan data. Keunggulannya adalah tidak terlalu dipengaruhi oleh *outlier* seperti RMSE, sehingga lebih stabil untuk data dengan variasi besar. Dalam jurnal ini, MAE digunakan sebagai ukuran dasar untuk membandingkan performa model, karena memberikan pandangan langsung tentang rata-rata *error* yang dihasilkan oleh setiap metode peramalan.

2.3 Kerangka Pemikiran

Kerangka pemikiran dalam penelitian ini dirancang untuk menggambarkan alur logis dalam proses pengolahan data hingga pengukuran performa model deteksi. Setiap komponen dalam

kerangka ini saling berkesinambungan dan memiliki peran masing-masing dalam mencapai tujuan penelitian. Adapun penjabaran tiap komponen dapat dijelaskan sebagai berikut:



Gambar 2. 8 Kerangka Pemikiran

Pada **Gambar 2.8** menunjukkan kerangka pemikiran penelitian yang menggambarkan alur proses dalam melakukan prediksi pemanfaatan FABA. Proses dimulai dari pengumpulan *dataset* pemanfaatan limbah FABA yang berisi data historis hasil pemanfaatan *Fly Ash* dan *Bottom Ash*. Data kemudian diproses pada tahap *preprocessing*, yaitu dengan membandingkan data asli dan data yang telah diperhalus menggunakan *Gaussian smoothing* untuk melihat pengaruh penghalusan data terhadap hasil analisis. Selanjutnya, kedua jenis data tersebut dimodelkan menggunakan metode *Autoregressive Integrated Moving Average* (ARIMA) dan *Long Short-Term Memory* (LSTM) untuk melakukan peramalan pendapatan dari pemanfaatan FABA pada periode selanjutnya hingga tahun 2025. Pada tahap akhir, kinerja masing-masing model dievaluasi menggunakan metrik *Mean Absolute Error* (MAE), *Root Mean Squared Error* (RMSE), dan *Mean Absolute Percentage Error* (MAPE) guna menilai tingkat akurasi hasil peramalan yang diperoleh.

2.3.1 Indicators

Indikator utama penelitian ini adalah pemanfaatan *Fly Ash* dan *Bottom Ash* (FABA) yang dihasilkan dari proses pembakaran batu bara di Pembangkit Listrik Tenaga Uap (PLTU). Data yang digunakan berupa pendapatan hasil pemanfaatan FABA pada rentang waktu Januari 2022 – Desember 2024, yang akan digunakan untuk memprediksi pendapatan pada Januari 2025 – September 2025. Indikator ini merepresentasikan *beyond kWh revenue* atau pendapatan *non-listrik* yang berasal dari pemanfaatan limbah FABA sebagai bahan bernilai ekonomi.

2.3.2 *Algorithm*

Algoritma yang digunakan terdiri dari dua model pembandingan utama, yaitu ARIMA dan LSTM. Pemilihan kedua model ini didasarkan pada perbedaan kemampuan masing-masing dalam mengolah karakteristik data deret waktu. Karena sifat data produksi atau pendapatan FABA belum dapat dipastikan bersifat linier maupun non-linier, maka dilakukan pendekatan komparatif untuk mengetahui model yang paling sesuai.

2.3.2.1 *Preprocessing*

Pada tahap *data preprocessing*, dilakukan beberapa proses untuk mempersiapkan data sebelum pemodelan, salah satunya adalah *Gaussian smoothing*. Teknik ini digunakan untuk menghaluskan data dengan mengurangi fluktuasi ekstrem sehingga pola utama data tetap terjaga. Dengan data yang lebih stabil, proses analisis dan pemodelan dapat dilakukan dengan lebih optimal.

2.3.2.2 *Modelling*

Pada tahap *modelling*, dilakukan pemodelan menggunakan dua metode, yaitu *Autoregressive Integrated Moving Average* (ARIMA) dan *Long Short-Term Memory* (LSTM).

1. ARIMA (*AutoRegressive Integrated Moving Average*). Merupakan metode statistik klasik untuk memprediksi data deret waktu (*time series*) berdasarkan hubungan linier antara nilai saat ini dengan nilai masa lalu. ARIMA digunakan sebagai model dasar (*baseline*) karena mampu menangkap pola trend dan musiman pada data historis.
2. LSTM (*Long Short-Term Memory*). Termasuk dalam kategori *Recurrent Neural Network* (RNN) yang dirancang untuk mempelajari pola jangka panjang (*long-term dependencies*) pada data deret waktu. LSTM digunakan sebagai model utama (*advanced model*) karena memiliki kemampuan memproses data *non-linear* dan menangkap fluktuasi pendapatan FABA yang kompleks.

Kedua algoritma ini digunakan untuk membandingkan tingkat akurasi hasil prediksi, guna mengetahui model yang paling sesuai dalam memprediksi pendapatan hasil pemanfaatan FABA.

2.3.3 *Objective*

Tujuan dari penelitian ini adalah untuk memprediksi hasil pemanfaatan FABA pada periode selanjutnya (Januari – September 2025) berdasarkan data historis dari Januari 2022 – Desember

2024. Penelitian ini juga bertujuan untuk mengetahui tren perubahan pendapatan (apakah meningkat atau menurun), serta menilai keakuratan hasil prediksi antara metode ARIMA dan LSTM. Hasil prediksi ini diharapkan dapat menjadi dasar informasi bagi pengelola PLTU atau pihak terkait dalam mengambil keputusan strategis terhadap pengelolaan dan pemanfaatan FABA.

2.3.4 *Measurement*

Evaluasi performa model dilakukan menggunakan tiga metrik utama diantaranya :

- 1) *Mean Absolute Percentage Error* (MAPE) untuk mengukur tingkat kesalahan relatif antara nilai aktual dan nilai prediksi dalam bentuk persentase.
- 2) *Root Mean Square Error* (RMSE) untuk mengukur besarnya deviasi atau kesalahan prediksi dengan memberikan penalti lebih besar pada kesalahan yang besar.
- 3) *Mean Absolute Error* (MAE) untuk menunjukkan rata-rata kesalahan absolut tanpa memperhatikan arah kesalahan.

Ketiga metrik ini digunakan untuk membandingkan akurasi hasil prediksi antara model ARIMA dan LSTM, di mana nilai kesalahan yang lebih kecil menandakan performa model yang lebih baik.