

**PENERAPAN LSTM (*LONG SHORT TERM MEMORY*) UNTUK
MENDETEKSI PENIPUAN MENGGUNAKAN KARTU
KREDIT**

TUGAS AKHIR

**Diajukan oleh:
TEUKU GHIFARI TS
NIM:210705079**

**Mahasiswa Fakultas Sains dan Teknologi
Program Studi Teknologi Informasi**



**PROGRAM STUDI TEKNOLOGI INFORMASI
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI AR-RANIRY
BANDA ACEH
2026 M/1447 H**

LEMBAR PERSETUJUAN

PENERAPAN LSTM (*LONG SHORT TERM MEMORY*) UNTUK MENDETEKSI PENIPUAN MENGGUNAKAN KARTU KREDIT

TUGAS AKHIR

Diajukan Kepada Fakultas Sains dan Teknologi
Universitas Islam Negeri (UIN) Ar-Raniry Banda Aceh
Sebagai Salah Satu Beban Memperoleh Gelar Sarjana
pada Program Studi Teknologi Informasi

Oleh:

TEUKU GHIFARI TS

NIM:210705079

**Mahasiswa Fakultas Sains dan Teknologi
Program Studi Teknologi Informasi**

Disetujui untuk Dimunakaqasyahkan Oleh::

Pembimbing I



Ghufran Ibnu Yasa, M.T
NIP. 198409262014031005

Pembimbing II



Hendri Ahmadian, S.Si., M.IM
NIP. 198301042014031002

Mengetahui:

Ketua Program Studi Teknologi Informasi



Malahayati, M.T
NIP. 198301272015032003

LEMBAR PENGESAHAN

PENERAPAN LSTM (*LONG SHORT TERM MEMORY*) UNTUK MENDETEKSI PENIPUAN MENGGUNAKAN KARTU KREDIT

TUGAS AKHIR

Telah Diuji Oleh Panitia Ujian Munaqasyah Tugas Akhir
Fakultas Sains dan Teknologi UIN Ar-Raniry Banda Aceh dan Dinyatakan Lulus
Serta Diterima Sebagai Salah Satu Beban Studi Program Sarjana (S1)
Dalam Program Studi Teknologi Informasi

Pada hari/Tanggal: Kamis, 12 Februari 2026

24 Sya'ban 1447 H

Di Darussalam, Banda Aceh

Panitia Ujian Munaqasyah Tugas Akhir:

Ketua,

Ghufran Ibnu Yasa, M.T
NIP. 198409262014031005

Sekretaris

Hendri Ahmadian, S.Si., M.IM
NIP. 198301042014031002

Penguji I,

Mursyidin, M.T
NIP. 196812301996031001

Penguji II,

Firmansyah, M.T
NIP. 196709261995031003

Mengetahui:

Dekan Fakultas Sains dan teknologi
UIN Ar-Raniry Banda Aceh



Prof. Dr. Ir. Muhammad Dirhamsyah, M.T., IPU
NIP. 196210021980111001

LEMBAR PERNYATAAN KEASLIAN

Yang bertanda tangan di bawah ini:

Nama : Teuku Ghifari TS

NIM : 210705079

Program Studi : Teknologi Informasi

Fakultas : Sains dan Teknologi

Judul : Penerapan Lstm (*Long Short Term Memory*) Untuk Mendeteksi Penipuan Menggunakan Kartu Kredit.

Dengan ini menyatakan bahwa dalam penulisan tugas akhir ini, saya:

1. Tidak menggunakan ide orang lain tanpa mampu mengembangkan dan mempertanggung jawabkan.
2. Tidak melakukan plagiasi terhadap masalah orang lain
3. Tidak menggunakan karya orang lain tanpa menyebutkan sumber asli atau tanpa izin pemilik karya.
4. Tidak memanipulasi dan memalsukan data.
5. Mengerjakan sendiri karya ini dan mampu bertanggung jawab atas karya ini.

Bila dikemudian hari ada tuntutan dari pihak lain atas karya saya, dan telah melalui pembuktian yang dapat dipertanggung jawabkan dan ternyata memang ditemukan bukti bahwa saya telah melanggar pernyataan ini, maka saya siap dikenakan sanksi berdasarkan aturan yang berlaku di Fakultas Sains dan Teknologi UIN Ar-Raniry Banda Aceh.

Demikian Pernyataan ini saya buat dengan sesungguhnya dan tanpa paksaan dari pihak manapun.

جامعة الرانيري

A R - R A N I R Y

Banda Aceh, 10 Februari 2026

Yang Menyatakan,



Teuku Ghifari TS

ABSTRAK

Nama : Teuku Ghifari TS
NIM : 210705079
Program Studi : Teknologi Informasi
Judul : Penerapan Lstm (*Long Short Term Memory*) Untuk Mendeteksi Penipuan Menggunakan Kartu Kredit.
Tanggal Sidang : 12 Februari 2026
Jumlah Halaman : 68 halaman
Pembimbing I : Ghufran Ibnu Yasa', M.T
Pembimbing II : Dr. Hendri Ahmadian, M.I.M
Kata Kunci : Deteksi penipuan, kartu kredit, *Long Short-Term Memory*, *deep learning*, *machine learning*.

Perkembangan teknologi digital telah mendorong peningkatan signifikan dalam penggunaan transaksi non-tunai, khususnya kartu kredit. Sistem deteksi penipuan konvensional yang berbasis aturan dan metode statistik tradisional dinilai kurang efektif dalam menghadapi pola penipuan yang semakin kompleks dan dinamis. Oleh karena itu, diperlukan pendekatan yang lebih adaptif dan akurat, salah satunya melalui penerapan metode deep learning. Penelitian ini bertujuan untuk menerapkan algoritma Long Short-Term Memory (LSTM) dalam mendeteksi penipuan transaksi kartu kredit serta menganalisis kinerja model berdasarkan metrik evaluasi yang relevan. Data yang digunakan dalam penelitian ini merupakan dataset sekunder, yaitu European Credit Card Fraud Dataset 2023 yang diperoleh dari platform Kaggle, dengan total lebih dari 500.000 transaksi yang telah dilabeli sebagai transaksi fraud dan non-fraud. Tahapan penelitian meliputi pengumpulan data, prapemrosesan data, pembentukan data sekuensial, pembangunan dan pelatihan model LSTM, serta evaluasi performa model. Model LSTM dibangun menggunakan bahasa pemrograman Python dengan bantuan pustaka TensorFlow dan Keras. Evaluasi kinerja model dilakukan menggunakan metrik accuracy, precision, recall, F1-score, dan Area Under Curve–Receiver Operating Characteristic (AUC-ROC). Hasil penelitian menunjukkan bahwa model LSTM mampu mengenali pola temporal dalam transaksi kartu kredit secara efektif dan memberikan performa yang baik dalam mendeteksi transaksi penipuan. Penerapan LSTM juga terbukti mampu mengurangi tingkat kesalahan deteksi dibandingkan metode konvensional. Berdasarkan hasil yang diperoleh, dapat disimpulkan bahwa algoritma LSTM memiliki potensi yang besar untuk digunakan sebagai dasar pengembangan sistem deteksi penipuan kartu kredit yang lebih akurat dan adaptif. Penelitian ini diharapkan dapat memberikan kontribusi bagi pengembangan sistem keamanan transaksi digital serta menjadi referensi bagi penelitian selanjutnya di bidang deteksi penipuan berbasis pembelajaran mesin.

ABSTRACT

Name : Teuku Ghifari TS
NIM : 210705079
Study Program : Teknologi Informasi
Title : *Application of LSTM (Long Short Term Memory) to Detect Fraud Using Credit Cards.*
Session Date : 12 February 2026
Number of Pages : 68 pages
Supervisor I : Ghufrani Ibnu Yasa', M.T
Supervisor II : Dr. Hendri Ahmadian, M.I.M
Keyword : *Fraud detection, credit cards, Long Short-Term Memory, deep learning, machine learning.*

The development of digital technology has driven a significant increase in the use of non-cash transactions, particularly credit cards. Conventional fraud detection systems based on rules and traditional statistical methods are considered ineffective in dealing with increasingly complex and dynamic fraud patterns. Therefore, a more adaptive and accurate approach is needed, one of which is through the application of deep learning methods. This study aims to apply the Long Short-Term Memory (LSTM) algorithm to detect fraudulent credit card transactions and analyze model performance based on relevant evaluation metrics. The data used in this study is a secondary dataset, namely the European Credit Card Fraud Dataset 2023 obtained from the Kaggle platform, with a total of over 500,000 transactions that have been labeled as fraudulent and non-fraudulent transactions. The research stages include data collection, data preprocessing, sequential data formation, building and training an LSTM model, and evaluating model performance. The LSTM model was built using the Python programming language with the help of the TensorFlow and Keras libraries. Model performance evaluation was carried out using accuracy, precision, recall, F1-score, and Area Under Curve–Receiver Operating Characteristic (AUC-ROC) metrics. The results of the study show that the LSTM model is able to recognize temporal patterns in credit card transactions effectively and provides good performance in detecting fraudulent transactions. The application of LSTM is also proven to be able to reduce the detection error rate compared to conventional methods. Based on the results obtained, it can be concluded that the LSTM algorithm has great potential to be used as a basis for developing a more accurate and adaptive credit card fraud detection system. This research is expected to contribute to the development of digital transaction security systems and become a reference for further research in the field of machine learning-based fraud detection.

KATA PENGANTAR

Bismillahirrahmanirrahim

Segala puji bagi Allah SWT atas limpahan rahmat, karunia, dan petunjukNya, sehingga penulis dapat menyelesaikan penyusunan Tugas Akhir yang berjudul: **“PENERAPAN LSTM (*LONG SHORT TERM MEMORY*) UNTUK MENDETEKSI PENIPUAN MENGGUNAKAN KARTU KREDIT”** Tugas Akhir ini disusun sebagai salah satu syarat untuk memperoleh gelar Sarjana (S1) pada Program Studi Teknologi Informasi, Fakultas Sains dan Teknologi, Universitas Islam Negeri Ar-Raniry Darussalam Banda Aceh.

Shalawat serta salam semoga senantiasa tercurahkan kepada Nabi Muhammad SAW, keluarga, para sahabat, dan seluruh umatnya yang terus memperjuangkan ajaran Islam hingga akhir zaman.

Penyelesaian Tugas Akhir ini tidak lepas dari dukungan, bimbingan, serta doa dari berbagai pihak. Oleh karena itu, dengan segala hormat dan rasa syukur, penulis menyampaikan ucapan terima kasih yang sebesar-besarnya kepada:

1. Ibunda Cut Salmina dan Ayahanda Teuku Samsul, terima kasih atas kasih sayang yang tulus, doa yang tak pernah putus, serta dukungan moral dan materi yang menjadi sumber kekuatan dan semangat selama menjalani studi hingga penyelesaian Tugas Akhir ini.
2. Ibu Malahayati, M.T, selaku Ketua Program Studi Teknologi Informasi, beserta seluruh dosen di lingkungan prodi, yang telah membekali penulis dengan ilmu, arahan, dan motivasi selama masa perkuliahan.
3. Bapak Ghufrani Ibnu Yasa', M.T, selaku Pembimbing I, yang dengan sabar memberikan arahan, masukan, dan dukungan intelektual sejak awal proses penelitian hingga selesai. Bapak Dr. Hendri Ahmadian, M.I.M, selaku Pembimbing II, atas bimbingan, kritik konstruktif, serta waktu dan perhatian yang diberikan kepada penulis di tengah kesibukan yang padat.
4. Ibu Cut Ida Rahmadiana, S.Si, atas bantuan dan dukungan dalam proses administrasi akademik, yang sangat membantu kelancaran penyelesaian Tugas Akhir ini.
5. Bapak Prof. Dr. Ir. M. Dirhamsyah, M.T., IPU, selaku Dekan Fakultas Sains dan Teknologi UIN Ar-Raniry, atas segala dukungan dan fasilitas yang diberikan selama masa studi.

6. Seluruh dosen Program Studi Teknologi Informasi, yang telah memberikan ilmu dan wawasan di bidang teknologi informasi yang sangat berharga bagi penulis.
7. Seluruh teman-teman angkatan 2021 di Program Studi Teknologi Informasi, para dosen penguji, serta semua pihak yang telah berkontribusi, baik secara langsung maupun tidak langsung, dalam proses penyusunan Tugas Akhir ini. Terima kasih atas segala bentuk bantuan, semangat, dan doa.

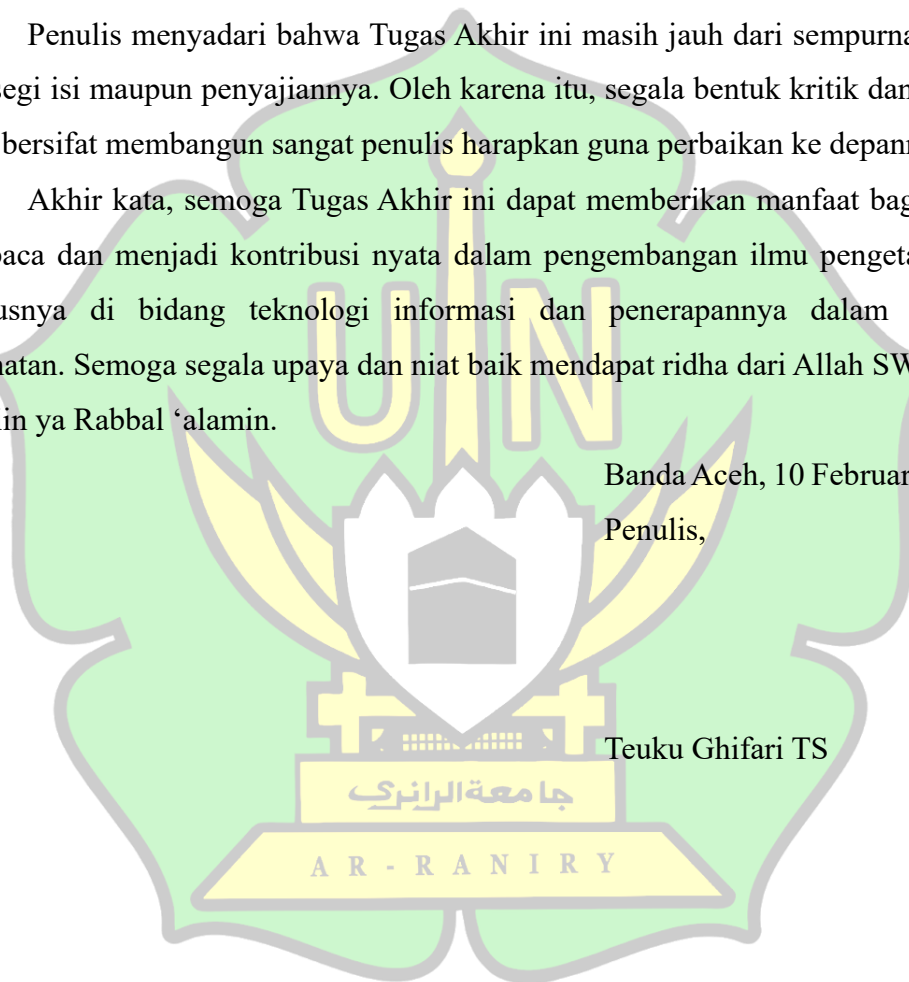
Penulis menyadari bahwa Tugas Akhir ini masih jauh dari sempurna, baik dari segi isi maupun penyajiannya. Oleh karena itu, segala bentuk kritik dan saran yang bersifat membangun sangat penulis harapkan guna perbaikan ke depannya.

Akhir kata, semoga Tugas Akhir ini dapat memberikan manfaat bagi para pembaca dan menjadi kontribusi nyata dalam pengembangan ilmu pengetahuan, khususnya di bidang teknologi informasi dan penerapannya dalam sektor kesehatan. Semoga segala upaya dan niat baik mendapat ridha dari Allah SWT. Aamiin ya Rabbal 'alamin.

Banda Aceh, 10 Februari 2025

Penulis,

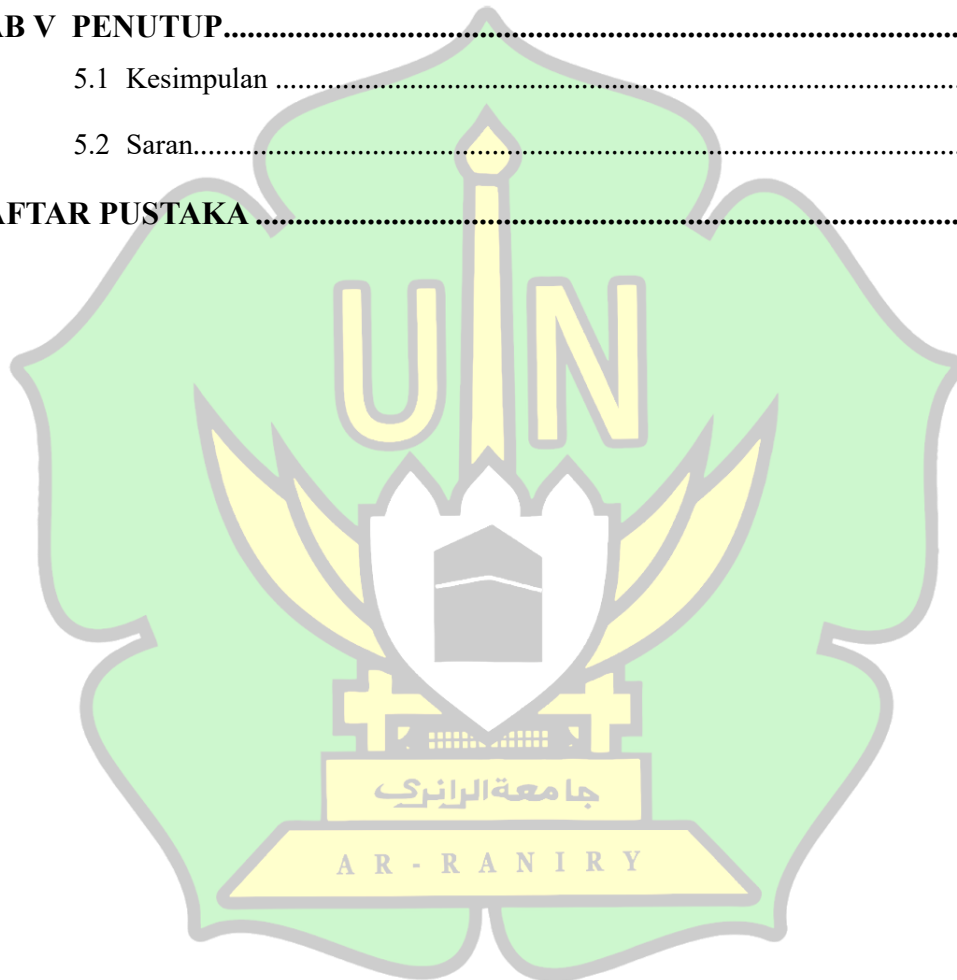
Teuku Ghifari TS



DAFTAR ISI

LEMBAR PERSETUJUAN	ii
LEMBAR PENGESAHAN	iii
LEMBAR PERNYATAAN KEASLIAN	iv
ABSTRAK	v
ABSTRACT	vii
KATA PENGANTAR	vii
DAFTAR ISI	ix
DAFTAR TABEL	xi
DAFTAR GAMBAR	xiii
BAB I PENDAHULUAN	1
1.1 Latar Belakang	1
1.2 Rumusan Masalah	5
1.3 Tujuan Penelitian	5
1.4 Manfaat Penelitian	6
1.5 Batas Penelitian	7
1.6 Sistematika Laporan	8
BAB II KAJIAN PUSTAKA	10
2.1 Penelitian Terdahulu	10
2.2 Konsep Penipuan dalam Transaksi Digital	15
2.3 Perkembangan Sistem Deteksi Penipuan Berbasis Machine Learning	17
2.4 Recurrent Neural Network (RNN) dan Long Short-Term Memory (LSTM)	19
2.5 Teknik Penanganan Ketidakseimbangan Data (<i>Class Imbalance</i>)	21
2.6 Kerangka Teori	22
BAB III METODE PENELITIAN	25
3.1 Tahap Penelitian	25
3.2 Alat dan Bahan	26
3.3 Waktu, Lokasi, atau Objek Penelitian	28
3.4 Populasi dan Sampel	28
3.5 Instrumen Penelitian	30

3.6 Teknik Pengumpulan Data.....	31
3.7 Analisis Data	34
BAB IV HASIL DAN PEMBAHASAN.....	39
4.1 Pengumpulan Data	39
4.2 Preprocessing data	46
4.3 Pembangunan Model LSTM.....	50
4.4 Evaluasi.....	55
BAB V PENUTUP.....	61
5.1 Kesimpulan	61
5.2 Saran.....	62
DAFTAR PUSTAKA	63



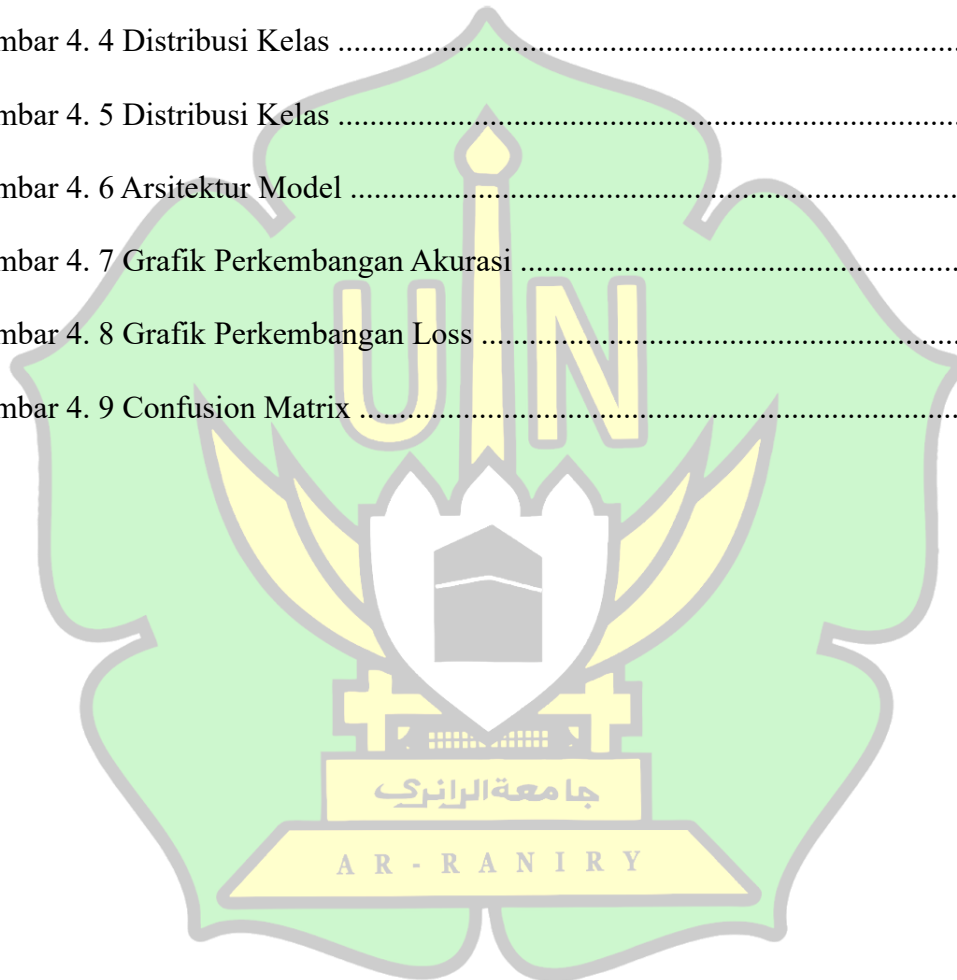
DAFTAR TABEL

Tabel 2. 1 Penelitian Terdahulu	14
Tabel 4. 1 5 Data Pertama	38
Tabel 4. 2 5 Data Teratas Sebelum dan Setelah Normalisasi	46
Tabel 4. 3 Akurasi Setiap Iterasi (Epoch)	55
Tabel 4. 4 Evaluasi Performa Model	59



DAFTAR GAMBAR

Gambar 2.1 Kerangka teori	24
Gambar 3. 1 Flowchart	25
Gambar 4.1 Struktur Dataset	40
Gambar 4. 2 Jumlah Missing Value	42
Gambar 4. 3 Jumlah Duplicate Rows	42
Gambar 4. 4 Distribusi Kelas	44
Gambar 4. 5 Distribusi Kelas	48
Gambar 4. 6 Arsitektur Model	50
Gambar 4. 7 Grafik Perkembangan Akurasi	56
Gambar 4. 8 Grafik Perkembangan Loss	57
Gambar 4. 9 Confusion Matrix	60



BAB I

PENDAHULUAN

1.1 Latar Belakang

Revolusi digital telah membawa perubahan signifikan dalam ekosistem transaksi keuangan global. Perkembangan teknologi informasi dan komunikasi mendorong adopsi transaksi non-tunai yang semakin luas, baik dalam skala individu maupun korporasi. Berdasarkan laporan *World Payments Report 2023* yang dirilis oleh Capgemini, volume transaksi non-tunai global diperkirakan mencapai 1,3 triliun transaksi pada tahun 2023, dengan tingkat pertumbuhan tahunan sebesar 16,6% dibandingkan tahun sebelumnya. Pertumbuhan ini mencerminkan peningkatan signifikan dalam adopsi pembayaran digital secara global, yang didorong oleh perkembangan infrastruktur pembayaran dan inovasi teknologi di sektor keuangan (Wire, 2023). Di Indonesia, data dari Bank Indonesia menunjukkan bahwa nilai transaksi kartu kredit pada tahun 2022 tercatat sebesar Rp 315,7 triliun (Romero, 2023). Peningkatan signifikan ini mencerminkan perubahan preferensi konsumen terhadap sistem pembayaran digital yang lebih praktis dan efisien. Namun, seiring dengan kemudahan yang ditawarkan, transaksi digital juga menghadapi tantangan serius berupa ancaman kejahatan siber, salah satunya adalah penipuan kartu kredit.

Menurut *Nilson Report 2023*, kerugian global akibat penipuan kartu pembayaran mencapai USD 33,83 miliar pada tahun 2023, mengalami peningkatan tipis sebesar 1,1% dari tahun sebelumnya. Proyeksi menunjukkan

bahwa total kerugian kumulatif dalam satu dekade mendatang dapat mencapai USD 403,88 miliar, dipengaruhi oleh peningkatan transaksi *e-commerce* dan risiko *card-not-present* (CNP). Amerika Serikat tercatat sebagai negara dengan kontribusi kerugian terbesar, mencapai 42% dari total kerugian global (BARBARA, 2025). Di Indonesia, laporan dari Association of Financial Professionals (AFP) menunjukkan bahwa 80% organisasi menjadi korban serangan atau percobaan penipuan pembayaran pada tahun 2023, meningkat 15% dibandingkan tahun 2022 (CPI, 2023). Upaya preventif yang lebih canggih, termasuk adopsi teknologi keamanan berbasis chip EMV, enkripsi *end-to-end*, dan peningkatan edukasi konsumen, menjadi krusial dalam mitigasi risiko penipuan kartu pembayaran.

Sistem deteksi penipuan konvensional yang berbasis aturan (*rule-based system*) dan metode statistik tradisional mulai menunjukkan keterbatasannya dalam menghadapi pola penipuan yang semakin kompleks. Penelitian (AlHashedi & Magalingam, 2021) mengungkapkan bahwa pendekatan rule-based hanya mampu mendeteksi sekitar 65-70% kasus penipuan, dengan tingkat false positive yang mencapai 30%. Keterbatasan ini berdampak pada efisiensi operasional lembaga keuangan dan kepercayaan nasabah terhadap keamanan transaksi mereka. Oleh karena itu, diperlukan pendekatan yang lebih adaptif dan akurat dalam mendeteksi aktivitas mencurigakan dalam transaksi kartu kredit.

Perkembangan teknologi kecerdasan buatan (*Artificial Intelligence/AI*) dan pembelajaran mesin (*Machine Learning/ML*) telah membuka peluang baru

dalam meningkatkan efektivitas deteksi penipuan. Metode *deep learning*, khususnya, telah menunjukkan keunggulan dalam menganalisis pola transaksi yang kompleks dan berskala besar. Penelitian (Sarker, 2021) yang menunjukkan bahwa arsitektur deep learning seperti *Convolutional Neural Networks* (CNN) dan *Recurrent Neural Networks* (RNN) mampu meningkatkan akurasi deteksi penipuan hingga 89%, dibandingkan metode konvensional yang hanya mencapai 75%.

Salah satu arsitektur deep learning yang terbukti efektif dalam analisis data sekuensial adalah *Long Short-Term Memory* (LSTM), yang merupakan varian dari RNN. LSTM pertama kali diperkenalkan oleh Hochreiter dan Schmidhuber pada tahun 1997 dan telah mendapatkan momentum signifikan dalam satu dekade terakhir (Rahayu, 2025). Keunggulan utama LSTM terletak pada kemampuannya dalam mengatasi masalah vanishing gradient, sehingga mampu mengingat informasi jangka panjang yang relevan dalam suatu sekuens data. Penelitian oleh (Takiddin et al., 2023) dalam menunjukkan bahwa LSTM mampu mendeteksi anomali dalam pola transaksi dengan akurasi 90%, jauh melampaui metode machine learning konvensional yang hanya mencapai 82,7%.

Implementasi LSTM dalam deteksi penipuan kartu kredit memiliki potensi besar dalam mengidentifikasi pola transaksi yang mencurigakan secara lebih akurat. (Benchaji et al., 2021) menemukan bahwa pola penipuan sering kali menunjukkan karakteristik temporal tertentu yang dapat diidentifikasi melalui model berbasis LSTM. Buku eksperimental menggunakan dataset transaksi kartu kredit dari *European Credit Card Database* menunjukkan bahwa model

LSTM mencapai F1-score sebesar 0,89 dalam mendeteksi penipuan, melampaui metode time-series lainnya seperti ARIMA (0,74) dan Gradient Boosting (0,81) (Fons, 2021). Temuan ini menegaskan bahwa penggunaan LSTM dapat meningkatkan ketepatan deteksi sekaligus mengurangi tingkat false positive yang kerap menjadi kendala dalam sistem konvensional.

Tantangan utama dalam implementasi model pembelajaran mesin untuk deteksi penipuan kartu kredit adalah ketidakseimbangan data (*class imbalance*), di mana jumlah transaksi penipuan jauh lebih sedikit dibandingkan transaksi normal. Menurut (Singla, 2022) mencatat bahwa rasio transaksi penipuan terhadap transaksi normal umumnya berkisar antara 1:1000 hingga 1:10000. Ketimpangan ini menyebabkan model pembelajaran mesin cenderung bias terhadap kelas mayoritas, sehingga menurunkan efektivitas deteksi anomali. Studi oleh (Kavitha & Kasthuri, 2024) menunjukkan bahwa kombinasi LSTM dengan teknik resampling seperti *Synthetic Minority Oversampling Technique* (SMOTE) dan pendekatan *cost-sensitive learning* dapat meningkatkan sensitivitas model hingga 23% dibandingkan implementasi standar.

Selain itu, perkembangan teknologi big data dan peningkatan kapasitas komputasi memungkinkan implementasi model LSTM dalam skala yang lebih besar dan real-time. Studi oleh (Sathupadi et al., 2025) mengungkapkan bahwa integrasi LSTM dengan framework Apache Spark memungkinkan pemrosesan hingga 50.000 transaksi per detik dengan latensi kurang dari 100 milidetik. (Ahsun et al., 2025) melaporkan bahwa deteksi penipuan secara real-time

berpotensi mengurangi kerugian finansial hingga 40% dibandingkan dengan deteksi pasca-transaksi. Transformasi ini menandai perubahan paradigma dalam sistem keamanan transaksi digital yang lebih proaktif dan adaptif.

Berdasarkan urgensi dan relevansi permasalahan yang telah diuraikan, penelitian ini bertujuan untuk mengeksplorasi penerapan model LSTM dalam mendeteksi penipuan kartu kredit secara lebih akurat dan efisien. Dengan menganalisis karakteristik temporal dari pola transaksi keuangan, penelitian ini diharapkan dapat memberikan kontribusi bagi pengembangan sistem deteksi penipuan yang lebih cerdas dan adaptif, sehingga dapat meningkatkan keamanan serta kepercayaan masyarakat terhadap sistem pembayaran digital.

1.2 Rumusan Masalah

Berdasarkan latar belakang yang telah dijelaskan, penelitian ini merumuskan beberapa pertanyaan utama sebagai berikut:

1. Bagaimana tahapan proses pengembangan model LSTM pada deteksi penipuan transaksi kartu kredit?
2. Bagaimana akurasi model LSTM dalam mendeteksi penipuan transaksi kartu kredit?

1.3 Tujuan Penelitian

Adapun tujuan dari penelitian ini adalah sebagai berikut:

1. Mengidentifikasi, merancang, dan mengimplementasikan tahapan pengembangan model LSTM untuk deteksi penipuan transaksi kartu kredit.
2. Menganalisis akurasi model LSTM dalam mendeteksi transaksi penipuan kartu kredit.

1.4 Manfaat Penelitian

Penelitian ini diharapkan dapat memberikan manfaat baik dari segi teoritis maupun praktis sebagai berikut:

1. Manfaat Teoritis

- a. Menambah wawasan dalam bidang kecerdasan buatan (Artificial Intelligence) dan pembelajaran mesin (Machine Learning), khususnya dalam penerapan Long Short-Term Memory (LSTM) untuk deteksi penipuan kartu kredit.
- b. Memberikan kontribusi bagi perkembangan literatur akademik terkait dengan penerapan deep learning dalam keamanan finansial dan deteksi anomali transaksi.
- c. Menjadi dasar bagi penelitian selanjutnya dalam mengembangkan model yang lebih akurat dan efisien dalam mendeteksi penipuan berbasis data sekuensial.

2. Manfaat Praktis

- a. Membantu institusi keuangan dan perbankan dalam meningkatkan sistem deteksi penipuan sehingga dapat mengurangi risiko kerugian akibat transaksi yang tidak sah.
- b. Memberikan solusi berbasis kecerdasan buatan yang lebih efektif dan adaptif dalam mendeteksi pola transaksi mencurigakan dibandingkan dengan metode konvensional.
- c. Mendukung implementasi sistem deteksi penipuan kartu kredit secara real-time untuk meningkatkan keamanan transaksi digital dan menjaga kepercayaan pelanggan.

1.5 Batas Penelitian

Dalam penelitian ini, terdapat beberapa batasan yang diterapkan agar penelitian lebih terfokus dan sesuai dengan tujuan yang ingin dicapai. Adapun batasan perencanaan dalam penelitian ini adalah sebagai berikut:

1. Ruang Lingkup Data

- a. Data yang digunakan dalam penelitian ini berasal dari dataset transaksi kartu kredit yang telah diklasifikasikan sebagai transaksi sah dan penipuan.
- b. Dataset yang digunakan merupakan data sekunder yang diperoleh dari sumber yang kredibel, seperti European Credit Card Fraud Dataset atau dataset transaksi keuangan lainnya.

2. Metode Deteksi Penipuan

- a. Model yang digunakan dalam penelitian ini adalah Long Short-Term Memory (LSTM), yang merupakan salah satu metode deep learning berbasis jaringan saraf tiruan (Neural Networks).
- b. Penelitian ini berfokus pada pengembangan dan evaluasi model LSTM untuk mendeteksi penipuan transaksi kartu kredit, tanpa mencakup implementasi sistem secara real-time atau integrasi dengan platform pembayaran sebenarnya.

3. Teknik Evaluasi Model

- a. Model dievaluasi menggunakan metrik akurasi, precision, recall, F1score, dan Area Under Curve - Receiver Operating Characteristic (AUC-ROC).

- b. Studi ini tidak mencakup implementasi model dalam sistem produksi secara langsung, melainkan hanya pada tahap simulasi dan evaluasi berbasis dataset yang digunakan.

4. Implementasi dan Penggunaan

- a. Penelitian ini hanya berfokus pada analisis efektivitas model dalam mendeteksi transaksi penipuan berdasarkan pola transaksi historis.
- b. Tidak membahas aspek regulasi keuangan, kebijakan perusahaan dalam menangani penipuan, atau mekanisme investigasi penipuan secara hukum.

1.6 Sistematika Laporan

Laporan penelitian ini disusun dengan sistematika sebagai berikut:

BAB 1 PENDAHULUAN

Bab ini berisi latar belakang penelitian, rumusan masalah, tujuan perencanaan, manfaat penelitian, pendekatan penelitian, batasan penelitian, serta sistematika laporan.

BAB 2 TINJAUAN PUSTAKA

Bab ini menjelaskan teori-teori yang mendukung penelitian, termasuk konsep deteksi penipuan kartu kredit, teknik machine learning, model Long Short-Term Memory (LSTM), serta penelitian terdahulu yang relevan.

BAB 3 METODOLOGI PENELITIAN

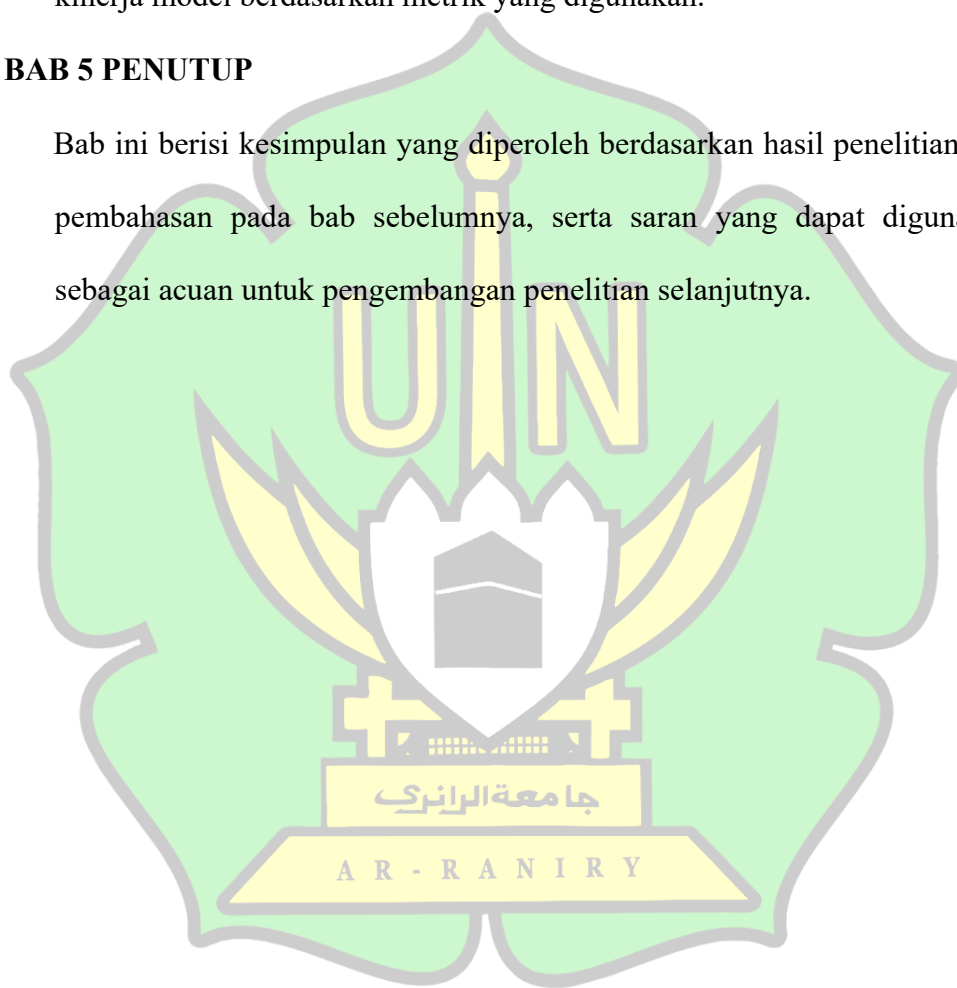
Bab ini membahas metode penelitian yang digunakan, mencakup tahapan pengolahan data, pemodelan, serta teknik evaluasi model. Selain itu, bab ini menjelaskan alat dan bahan yang digunakan dalam penelitian.

BAB 4 HASIL DAN PEMBAHASAN

Bab ini menyajikan hasil penelitian dan pembahasan mengenai penerapan model Long Short-Term Memory (LSTM) Dalam mendeteksi penipuan transaksi kartu kredit. Pembahasan meliputi proses pengolahan dan praperosesan data, pembangunan serta pelatihan model, dan evaluasi kinerja model berdasarkan metrik yang digunakan.

BAB 5 PENUTUP

Bab ini berisi kesimpulan yang diperoleh berdasarkan hasil penelitian dan pembahasan pada bab sebelumnya, serta saran yang dapat digunakan sebagai acuan untuk pengembangan penelitian selanjutnya.



BAB II

KAJIAN PUSTAKA

2.1 Penelitian Terdahulu

Penelitian mengenai penerapan metode deep learning, khususnya Long Short-Term Memory (LSTM), dalam mendeteksi penipuan dan klasifikasi anomali pada data sekuensial telah berkembang pesat dalam beberapa tahun terakhir. Sejumlah studi relevan telah dilakukan dengan konteks dan pendekatan yang beragam, mulai dari media sosial, pesan singkat, hingga sistem investasi dan rekrutmen daring. Ulasan terhadap penelitian terdahulu ini penting untuk memperkuat landasan teoritis serta menunjukkan celah penelitian yang akan diisi oleh studi ini. Berikut adalah beberapa penelitian yang relevan:

Penelitian oleh (Dusenov & Wibowo, 2025) berjudul *“Deteksi Penipuan pada Sosial Media Twitter dengan Metode Bidirectional Long Short Term Memory (BI-LSTM)”* mengkaji penerapan algoritma Bi-LSTM dalam mendeteksi aktivitas penipuan di platform Twitter. Penelitian ini menggunakan pendekatan deskriptif untuk memahami pola interaksi, perilaku pengguna, serta sentimen yang ditunjukkan dalam cuitan-cuitan berbahasa Indonesia. Dengan memanfaatkan data yang dikumpulkan melalui API dan web crawler, model BiLSTM yang dikembangkan menunjukkan performa lebih tinggi dibanding LSTM standar dalam hal akurasi, presisi, dan recall. Penelitian ini menekankan pentingnya pemilihan kata kunci yang relevan dalam pengumpulan data, serta perlunya proses validasi dan tuning parameter model untuk hasil yang optimal.

Studi ini menunjukkan bahwa Bi-LSTM sangat efektif untuk mendeteksi penipuan dalam konteks media sosial, namun konteksnya berbeda dengan transaksi keuangan yang menjadi fokus dalam penelitian ini.

Penelitian selanjutnya dilakukan oleh (Budiyansyah, 2025) dengan judul *“Prediksi Real or Fake Job Posting Menggunakan Metode Long Short-Term Memory”*. Studi ini merespons maraknya lowongan pekerjaan palsu yang tersebar secara daring dan membahayakan pencari kerja. Dengan menggunakan dataset *Real or Fake Job Posting Prediction* dari Kaggle, Budiyansyah mengembangkan model klasifikasi berbasis LSTM untuk membedakan antara lowongan yang asli dan yang palsu. Teknik NLP seperti tokenisasi dan lemmatization digunakan dalam tahap prapemrosesan data, sementara pelatihan model dilakukan dengan framework TensorFlow. Hasilnya, model LSTM yang dibangun mampu mencapai akurasi sebesar 97,61% dan nilai loss sebesar 0,08%. Penelitian ini membuktikan efektivitas LSTM dalam memproses data teks dan mengidentifikasi penipuan berbasis narasi. Meskipun fokusnya bukan pada transaksi kartu kredit, pendekatan dan struktur model yang digunakan memberikan gambaran bahwa LSTM juga dapat diadaptasi dalam konteks deteksi penipuan finansial.

Sementara itu, (Muslikah, 2021) dalam penelitiannya yang berjudul *“SMS Spam Detection Menggunakan Metode Long Short Term Memory”* menitikberatkan pada klasifikasi pesan spam dalam layanan SMS. Dalam penelitian ini, model LSTM digunakan untuk membedakan pesan spam dan non-spam yang berbahasa Indonesia. Hasil evaluasi menunjukkan bahwa

model LSTM bekerja sangat baik dengan akurasi, presisi, recall, dan F1-score yang masing-masing mencapai 96,09%. Penelitian ini menjadi bukti bahwa LSTM sangat efisien dalam memproses data sekuensial berbasis teks dan mampu mengidentifikasi pola anomali yang tersembunyi. Persamaan utama dengan penelitian ini terletak pada penggunaan LSTM sebagai model utama, namun perbedaan mencoloknya adalah objek data yang digunakan, di mana Muslikah berfokus pada SMS, sedangkan penelitian ini menggunakan data transaksi kartu kredit.

Penelitian oleh (Ismail, 2020) berjudul *“Klasifikasi Short Message Service (SMS) Menggunakan Algoritma Long Short Term Memory dan Recurrent Neural Network”* juga mengangkat topik klasifikasi pesan singkat namun memperluas kajiannya dengan membandingkan dua algoritma, yaitu LSTM dan Recurrent Neural Network (RNN). Penelitian ini menggunakan dataset berbahasa Inggris dan Indonesia, serta membagi kategori pesan menjadi spam, promo, dan penipuan. Hasil eksperimen menunjukkan bahwa LSTM unggul dibandingkan RNN dalam hal akurasi, presisi, dan recall. Dalam beberapa konfigurasi, model LSTM bahkan mencapai akurasi 98% dan F1-score 0,99. Penelitian ini memberikan kontribusi dalam menunjukkan keunggulan LSTM atas RNN dalam klasifikasi data sekuensial. Meskipun konteksnya bukan transaksi keuangan, pendekatan perbandingan algoritma yang dilakukan sejalan dengan arah penelitian ini yang juga akan membandingkan performa LSTM dengan metode konvensional lainnya.

Terakhir, (Nugraha & Arrofiq, 2024) melakukan penelitian dengan judul

“Purwarupa Sistem Klasifikasi Legalitas Investasi Berbasis Algoritma Bidirectional Long Short Term Memory”. Fokus penelitian ini adalah pada deteksi legalitas pesan investasi yang tersebar melalui media sosial, khususnya Telegram. Dengan menerapkan metode text classification berbasis Bi-LSTM dan LSTM, penelitian ini berhasil mengembangkan purwarupa sistem deteksi legalitas yang mampu mengklasifikasikan pesan sebagai legal atau ilegal dengan akurasi mencapai 96%, presisi 98%, dan recall 93%. Selain membangun model klasifikasi, penelitian ini juga mengimplementasikan hasilnya dalam bentuk aplikasi web yang mampu melakukan deteksi secara real-time. Persamaan yang menonjol dengan penelitian ini adalah penggunaan model LSTM dan varian Bi-LSTM untuk mendeteksi konten penipuan, serta penggunaan dataset bahasa Indonesia. Namun, konteks dan domain penipuan yang dianalisis tetap berbeda, yakni pesan investasi, bukan transaksi kartu kredit.

Melalui lima studi tersebut, dapat disimpulkan bahwa LSTM dan variannya terbukti memiliki performa yang sangat baik dalam mengolah data teks sekuensial dan mendeteksi pola penipuan dalam berbagai konteks. Penelitian ini akan melanjutkan upaya tersebut dengan fokus pada domain transaksi kartu kredit, serta mengangkat aspek real-time detection dan ketidakseimbangan data (class imbalance) sebagai tantangan utama yang belum banyak dibahas secara mendalam dalam studi-studi sebelumnya.

Tabel 2. 1 Penelitian Terdahulu

No	Nama & Judul	Metode	Hasil	Persamaan	Perbedaan
1	Dusenov & Wibowo (2025) “Deteksi Penipuan pada Sosial Media Twitter dengan Metode Bidirectional LSTM”	& BI-LSTM, Deskriptif, Data dari Twitter API & Web Crawling	BI-LSTM lebih unggul dari LSTM (akurasi, presisi, recall tinggi), dengan optimasi seperti early stopping	Sama-sama mendeteksi penipuan, menggunakan LSTM, fokus pada data teks	Fokus pada Twitter & fraud sosial media, bukan transaksi keuangan
2	Budiyansyah (2025) “Prediksi Real or Fake Job Posting Menggunakan LSTM”	LSTM, NLP (tokenisasi, lemmatization), TensorFlow	Akurasi 97,61%, loss 0,08% dalam klasifikasi lowongan palsu	Sama-sama berbasis LSTM dan deteksi penipuan berbasis teks	Fokus pada deteksi iklan lowongan kerja palsu, bukan transaksi finansial
3	Muslikah (2021) “SMS Spam Detection Menggunakan LSTM”	LSTM, Spam Dataset SMS berbahasa Indonesia	Akurasi, presisi, recall, F1score semua 96,09%	Sama-sama menggunakan LSTM untuk klasifikasi teks dan deteksi penipuan/spam	Fokus pada deteksi spam SMS, bukan penipuan keuangan
4	Ismail (2020) “Klasifikasi SMS menggunakan LSTM dan RNN”	LSTM & RNN, Dataset SMS Inggris & Indonesia	Akurasi LSTM hingga 98%, precision & recall 98-99%	Sama-sama menggunakan LSTM dan RNN menguji efektivitas deteksi berbasis teks	Bandingkan LSTM dan RNN dalam klasifikasi spam dan penipuan, bukan transaksi kartu kredit

5	Nugraha Arrofiq (2024) “Purwarupa Sistem Klasifikasi Legalitas Investasi Berbasis BiLSTM”	& BI-LSTM & LSTM, Telegram dataset investasi	Akurasi 96%, presisi 98%, recall 93%	Sama-sama menggunakan LSTM untuk mendeteksi penipuan, dengan perbandingan klasifikasi algoritma legal vs ilegal	Fokus pada klasifikasi legalitas pesan investasi dari media sosial
---	--	--	---	--	--

Sumber: Data Diolah (2025)

2.2 Konsep Penipuan dalam Transaksi Digital

Perkembangan teknologi digital telah merevolusi sistem keuangan global, termasuk di dalamnya sistem pembayaran yang semakin bergeser dari transaksi tunai ke transaksi non-tunai. Transformasi ini memicu peningkatan signifikan dalam volume transaksi digital yang pada satu sisi memberikan kemudahan dan efisiensi, namun di sisi lain juga menghadirkan risiko baru dalam bentuk kejahatan siber, salah satunya adalah penipuan dalam transaksi digital (Ngamal & Perajaka, 2022). Fenomena ini menjadi tantangan serius bagi institusi keuangan yang dituntut untuk mengembangkan sistem pengamanan yang adaptif dan canggih.

Penipuan dalam transaksi digital dapat didefinisikan sebagai tindakan yang disengaja untuk menipu atau memperoleh keuntungan finansial secara tidak sah dengan memanfaatkan celah dalam sistem transaksi elektronik. Dalam konteks kartu kredit, bentuk penipuan yang umum terjadi antara lain pencurian identitas, penggunaan kartu tanpa otorisasi, manipulasi data transaksi, serta penyalahgunaan informasi kartu melalui modus *phishing*, *skimming*, dan serangan terhadap sistem pembayaran daring (*card-not-present fraud*) (Hendrayana et al., 2024).

Karakteristik utama dari transaksi penipuan adalah ketidakwajaran dalam pola transaksi yang dilakukan, baik dari sisi waktu, lokasi, maupun nominal transaksi. Contohnya adalah transaksi dalam jumlah besar yang dilakukan dalam waktu singkat di lokasi geografis yang tidak biasa bagi pengguna kartu. Transaksi semacam ini secara statistik bersifat *outlier* dan menjadi indikasi awal potensi terjadinya penipuan (Lestari & Sirodj, 2021). Oleh karena itu, deteksi penipuan memerlukan pendekatan yang mampu mengenali anomali dari pola transaksi historis pengguna.

Sistem deteksi penipuan konvensional umumnya menggunakan pendekatan berbasis aturan (*rule-based systems*) yang didesain berdasarkan logika tetap dan parameter-parameter yang telah ditentukan sebelumnya. Meskipun pendekatan ini mudah diimplementasikan, ia memiliki keterbatasan yang signifikan dalam menghadapi dinamika modus penipuan yang terus berkembang. Studi menunjukkan bahwa sistem berbasis aturan hanya mampu mendeteksi sekitar 30-70% kasus penipuan dengan tingkat kesalahan positif (*false positive*) yang tinggi, sehingga berdampak pada efisiensi operasional lembaga keuangan serta kepercayaan pelanggan (Jannah, 2020).

Seiring meningkatnya kompleksitas data transaksi dan semakin canggihnya teknik penipuan, pendekatan berbasis kecerdasan buatan dan pembelajaran mesin menjadi solusi yang lebih efektif dan efisien. Model pembelajaran mesin mampu mempelajari pola transaksi yang kompleks dan menyesuaikan diri dengan perubahan pola penipuan secara otomatis. Dalam konteks ini, pendekatan *deep learning*, khususnya model Long Short-Term Memory (LSTM), telah terbukti unggul dalam menganalisis data sekuensial seperti riwayat transaksi kartu kredit.

Dengan demikian, pemahaman yang mendalam terhadap konsep penipuan dalam transaksi digital menjadi dasar yang kuat dalam pengembangan sistem deteksi yang adaptif dan berbasis teknologi terkini. Penelitian ini didasarkan pada urgensi untuk membangun sistem yang tidak hanya mampu mendeteksi penipuan dengan akurasi tinggi, tetapi juga dapat diimplementasikan secara real-time dalam ekosistem keuangan digital yang terus berkembang.

2.3 Perkembangan Sistem Deteksi Penipuan Berbasis Machine Learning

1. Peralihan dari Sistem Konvensional ke Sistem Cerdas

Seiring berkembangnya teknologi transaksi digital, kejahatan siber pun turut berkembang dalam hal metode dan kompleksitas. Sistem deteksi penipuan konvensional yang berbasis aturan (*rule-based systems*) mulai menunjukkan keterbatasannya. Pendekatan ini hanya efektif pada pola penipuan yang sudah diketahui sebelumnya, namun gagal mendeteksi pola baru yang dinamis dan tidak terduga.

2. Munculnya Pendekatan Machine Learning

Machine Learning (ML) hadir sebagai solusi untuk mengatasi keterbatasan sistem konvensional. ML adalah metode kecerdasan buatan yang memungkinkan sistem belajar dari data historis untuk memprediksi atau mengklasifikasi transaksi, tanpa perlu ditentukan aturan secara eksplisit. Beberapa algoritma populer yang digunakan dalam deteksi penipuan menurut (Putri et al., 2024) ialah:

- a. *Random Forest*: metode ensemble yang membangun banyak pohon keputusan (decision trees) untuk memperoleh prediksi yang lebih stabil dan akurat.

- b. *Support Vector Machine (SVM)*: algoritma berbasis margin optimal, efektif untuk data berdimensi tinggi.
- c. *k-Nearest Neighbors (k-NN)*: algoritma berbasis kedekatan fitur antar data, namun kurang efektif untuk skala data besar.
- d. *Naïve Bayes*: algoritma probabilistik sederhana, cocok untuk klasifikasi awal.

Meskipun metode-metode ini memiliki keunggulan tersendiri, mereka masih terbatas dalam menangani data sekuensial dan temporal, seperti riwayat transaksi kartu kredit.

3. Kebutuhan Akan Model yang Mampu Mengenali Pola Waktu

Transaksi penipuan sering kali menunjukkan pola perilaku temporal, seperti transaksi berulang dalam waktu singkat, atau perubahan mendadak dalam lokasi dan nominal. Untuk memahami pola semacam ini, diperlukan model yang mampu mengingat informasi sebelumnya dan menyesuaikannya dengan data baru. Inilah yang melatarbelakangi penggunaan pendekatan Deep Learning, khususnya Recurrent Neural Networks (RNN) dan Long Short-Term Memory (LSTM) (Carudin et al., 2024).

4. Keunggulan Deep Learning: LSTM sebagai Solusi Modern

LSTM merupakan pengembangan dari RNN yang dirancang untuk mengatasi masalah *vanishing gradient*. Dengan arsitektur gerbang (input gate, forget gate, output gate), LSTM mampu mempertahankan informasi penting dalam jangka waktu yang panjang.

Menurut penelitian (Budiyanasyah, 2025), model LSTM mampu mendeteksi penipuan dengan akurasi hingga 90%, jauh lebih tinggi dibandingkan metode konvensional. LSTM juga terbukti efektif dalam:

- a. Mendeteksi pola transaksi jangka panjang.
- b. Mengurangi tingkat false positive.
- c. Beradaptasi dengan perubahan pola transaksi pengguna.

2.4 Recurrent Neural Network (RNN) dan Long Short-Term Memory (LSTM)

1. *Recurrent Neural Network (RNN)*

Dalam konteks deteksi penipuan kartu kredit, data transaksi memiliki sifat temporal atau berurutan. Artinya, informasi masa lalu bisa sangat menentukan dalam mendeteksi perilaku mencurigakan. Untuk menangani data seperti ini, diperlukan model yang memiliki “ingatan” terhadap input sebelumnya.

Recurrent Neural Network (RNN) adalah jenis arsitektur jaringan saraf tiruan yang dirancang khusus untuk memproses data sekuensial. RNN bekerja dengan cara mengolah data secara berurutan, di mana output dari waktu sebelumnya akan memengaruhi input pada waktu berikutnya. Dengan demikian, RNN sangat efektif dalam memahami konteks dari urutan data (Wanda et al., 2021).

Namun, RNN memiliki kelemahan mendasar yang dikenal sebagai vanishing gradient problem. Ketika memproses urutan data yang panjang, nilai gradien dalam proses pelatihan bisa menjadi sangat kecil sehingga menghambat pembelajaran. Akibatnya, RNN sulit untuk mengingat informasi dari langkah-langkah sebelumnya dalam urutan yang panjang (Caniago et al., 2021).

2. *Long Short-Term Memory (LSTM)*

Untuk mengatasi kelemahan RNN tersebut, Long Short-Term Memory

(LSTM) dikembangkan oleh Hochreiter dan Schmidhuber pada tahun 1997. LSTM merupakan varian dari RNN yang mampu menyimpan dan mengingat informasi dalam jangka waktu panjang, sehingga cocok untuk mendeteksi pola anomali yang tersembunyi dalam rangkaian transaksi (Arfan & Lussiana, 2020).

Arsitektur LSTM terdiri dari tiga gerbang utama:

- a. *Forget Gate*: memutuskan informasi mana yang harus dibuang dari memori.
- b. *Input Gate*: menentukan informasi baru mana yang akan disimpan.
- c. *Output Gate*: memutuskan bagian mana dari memori yang akan digunakan untuk output saat ini.

Dengan struktur ini, LSTM secara dinamis dapat memilih informasi mana yang penting dan mana yang tidak, sehingga memungkinkan pemodelan urutan data yang kompleks dan panjang.

3. Keunggulan LSTM dalam Deteksi Penipuan

Dalam berbagai studi, LSTM menunjukkan performa yang lebih unggul dibandingkan metode deteksi penipuan lainnya. Misalnya, penelitian oleh (Widhiyasana et al., 2021) menunjukkan bahwa model LSTM yang dikombinasikan dengan mekanisme perhatian (*attention mechanism*) mampu mencapai F1-score sebesar 0,89, mengungguli metode lain seperti Random Forest dan ARIMA. Selain itu, (Alim, 2023) mencatat bahwa LSTM unggul dalam:

- a. Mengidentifikasi pola berulang dalam data transaksi.
- b. Menangani data yang tidak seimbang (*class imbalance*).
- c. Menurunkan tingkat kesalahan deteksi (*false positive rate*).

Model LSTM juga fleksibel dan dapat digabungkan dengan teknik penguatan lain seperti SMOTE atau cost-sensitive learning untuk meningkatkan sensitivitas

terhadap transaksi penipuan yang jumlahnya sangat kecil dibandingkan transaksi normal (Kavitha & Kasthuri, 2024).

4. Potensi *Real-Time* dan Implementasi Skala Besar

LSTM tidak hanya unggul dalam akurasi, tetapi juga memiliki potensi implementasi real-time. Studi oleh (Alghafiqi & Munajat, 2022) menunjukkan bahwa integrasi LSTM dengan framework Apache Spark memungkinkan pemrosesan transaksi hingga 50.000 transaksi per detik dengan latensi kurang dari 100 milidetik. Hal ini membuka jalan bagi penggunaan model deteksi penipuan yang efisien dan responsif dalam sistem finansial digital.

RNN membuka jalan bagi pemrosesan data sekuensial dalam dunia deteksi penipuan, namun LSTM hadir sebagai solusi yang lebih stabil dan akurat. Keunggulan LSTM dalam mengingat konteks jangka panjang menjadikannya pilihan yang sangat sesuai untuk menganalisis transaksi kartu kredit yang memiliki karakteristik temporal kompleks. Oleh karena itu, penelitian ini memfokuskan penerapan model LSTM sebagai fondasi utama dalam pengembangan sistem deteksi penipuan yang modern, adaptif, dan real-time.

2.5 Teknik Penanganan Ketidakseimbangan Data (*Class Imbalance*)

Salah satu tantangan utama dalam penerapan model pembelajaran mesin untuk deteksi penipuan kartu kredit adalah ketidakseimbangan data atau *class imbalance*. Ketidakseimbangan ini terjadi ketika jumlah sampel pada satu kelas (umumnya transaksi normal) jauh lebih besar dibandingkan kelas lainnya (transaksi penipuan). Situasi ini menyebabkan model cenderung mempelajari karakteristik kelas mayoritas, sehingga mengabaikan kelas minoritas yang justru lebih krusial untuk dikenali (Ningsih et al., 2022).

2.6 Kerangka Teori

Penelitian ini berlandaskan pada sejumlah teori utama yang saling berkaitan dalam mendukung penerapan model Long Short-Term Memory (LSTM) untuk mendeteksi penipuan kartu kredit. Teori pertama yang mendasari adalah Fraud Theory, khususnya pendekatan Fraud Triangle yang dikembangkan oleh Donald Cressey. Teori ini menjelaskan bahwa penipuan umumnya terjadi karena adanya tekanan (pressure), kesempatan (opportunity), dan rasionalisasi (rationalization). Dalam konteks transaksi digital, tekanan dapat berupa kebutuhan finansial, kesempatan muncul akibat kelemahan sistem keamanan, dan rasionalisasi sering kali berupa pembenaran individu terhadap tindakannya. Teori ini memberikan pemahaman dasar mengenai latar belakang terjadinya penipuan dan pentingnya deteksi dini terhadap pola transaksi yang mencurigakan.

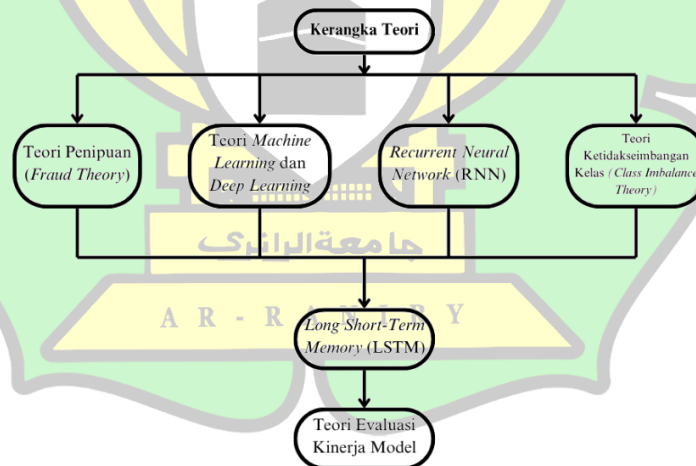
Selanjutnya, penelitian ini didukung oleh teori pembelajaran mesin (Machine Learning) dan pembelajaran dalam (Deep Learning), yang menjadi fondasi dalam pengembangan sistem deteksi otomatis berbasis data. Deep learning merupakan cabang dari machine learning yang menggunakan arsitektur jaringan saraf tiruan bertingkat untuk mengenali pola kompleks secara lebih akurat. Dalam konteks deteksi penipuan, pendekatan ini terbukti lebih unggul dibanding metode konvensional karena kemampuannya memproses data dalam jumlah besar, menangkap relasi non-linear, serta beradaptasi dengan pola penipuan baru yang dinamis.

Untuk mengatasi karakteristik data transaksi yang bersifat sekuensial atau temporal, digunakan teori Recurrent Neural Network (RNN). RNN

dirancang untuk mengolah data berurutan dengan memanfaatkan informasi dari waktu sebelumnya sebagai konteks prediksi berikutnya. Namun, RNN memiliki keterbatasan dalam menangani data panjang karena permasalahan vanishing gradient, sehingga informasi dari waktu yang jauh sering kali tidak dapat dipertahankan. Oleh karena itu, penelitian ini mengadopsi Long Short-Term Memory (LSTM) sebagai solusi pengembangan dari RNN. LSTM memiliki struktur khusus berupa tiga gerbang utama yaitu input gate, forget gate, dan output gate yang memungkinkan model untuk menyimpan, melupakan, dan menggunakan informasi dengan lebih selektif. Keunggulan ini menjadikan LSTM sangat efektif dalam mengenali pola jangka panjang yang tersembunyi dalam rangkaian transaksi keuangan, serta meminimalkan kesalahan deteksi.

Selain itu, penelitian ini juga dilandaskan pada teori class imbalance atau ketidakseimbangan kelas dalam data. Dalam banyak kasus penipuan kartu kredit, jumlah transaksi fraud jauh lebih sedikit dibanding transaksi normal, menyebabkan model cenderung mengabaikan kelas minoritas. Kondisi ini menurunkan sensitivitas model dalam mendeteksi penipuan. Untuk mengatasi hal ini, digunakan pendekatan Synthetic Minority Over-sampling Technique (SMOTE) yang bertujuan memperbanyak data sintetis pada kelas minoritas, serta cost-sensitive learning yang memberikan penalti lebih besar terhadap kesalahan klasifikasi pada transaksi fraud. Kedua pendekatan ini penting untuk memastikan bahwa model tidak bias terhadap kelas mayoritas dan tetap mampu mendeteksi kasus penipuan secara efektif.

Terakhir, performa model LSTM dievaluasi menggunakan teori evaluasi kinerja model yang terdiri dari berbagai metrik statistik. Accuracy digunakan untuk mengukur tingkat keseluruhan prediksi yang benar, precision menunjukkan ketepatan model dalam mengidentifikasi transaksi penipuan, recall menilai sejauh mana model mampu menangkap seluruh kasus penipuan yang sebenarnya terjadi, dan F1-score merupakan rata-rata harmonis antara precision dan recall. Selain itu, metrik Area Under Curve - Receiver Operating Characteristic (AUC-ROC) digunakan untuk mengevaluasi kemampuan model dalam membedakan antara transaksi fraud dan non-fraud. Dengan menggunakan kombinasi teori-teori tersebut, penelitian ini membangun kerangka konseptual yang kokoh untuk mengembangkan model deteksi penipuan kartu kredit yang cerdas, akurat, dan adaptif terhadap tantangan data transaksi digital yang kompleks.



Gambar 2.1 kerangka teori

BAB III

METODE PENELITIAN

3.1 Tahap Penelitian

Rancangan penelitian ini menggunakan pendekatan *experimental research* dengan membangun model deteksi penipuan menggunakan algoritma *Long Short-Term Memory* (LSTM). Penelitian ini mengevaluasi kemampuan model LSTM dalam mendeteksi aktivitas penipuan pada transaksi kartu kredit berdasarkan dataset yang tersedia. Tahapan penelitian ini dapat di gambarkan melalui flowchart berikut:



Gambar 3. 1 flowchart

Berdasarkan flowchart pada Gambar 3.1, tahapan penelitian ini diawali dengan proses pengumpulan data transaksi kartu kredit yang digunakan sebagai objek penelitian. Data dikumpulkan dari platform Kaggle menggunakan dataset

European Credit Card Fraud Dataset 2023 dengan jumlah data sebanyak 568.630 transaksi. Dataset tersebut telah dilengkapi dengan label transaksi *fraud* dan *non fraud* serta memiliki 30 fitur atau kolom, yang terdiri atas 28 fitur hasil transformasi *Principal Component Analysis* (PCA) serta dua fitur tambahan, yaitu *Amount* yang merepresentasikan nominal transaksi dan *Id* sebagai identitas data. Data yang telah diperoleh kemudian melalui tahap preprocessing untuk memastikan data siap digunakan dalam proses pemodelan. Pada tahap ini dilakukan normalisasi data menggunakan *StandardScaler*, pembagian data menjadi data latih dan data uji dengan rasio 80:20, serta konversi data ke dalam format tiga dimensi agar sesuai dengan kebutuhan input model LSTM. Selanjutnya, data yang telah diproses digunakan pada tahap implementasi model *Long Short-Term Memory* (LSTM) untuk membangun sistem deteksi penipuan. Model dibangun menggunakan beberapa *layer* atau lapisan yaitu LSTM, *dropout* dan *dense* untuk memaksimalkan kemampuan deteksi. Model yang telah dibangun kemudian dievaluasi menggunakan *accuracy*, *precision*, *recall*, *f1-score*, AUC-ROC dan *confusion matrix* untuk mengetahui kinerjanya dalam mengklasifikasikan transaksi penipuan dan nonpenipuan. Tahap akhir dari penelitian ini adalah penyajian hasil penelitian berdasarkan hasil evaluasi yang telah dilakukan.

3.2 Alat dan Bahan

Pada penelitian ini digunakan beberapa perangkat untuk mendukung proses pengolahan data, pemodelan, dan evaluasi model *Long Short-Term Memory* (LSTM) dalam mendeteksi penipuan pada transaksi kartu kredit.

Perangkat yang digunakan meliputi perangkat keras dan perangkat lunak.

1. Perangkat Keras

Perangkat keras yang digunakan dalam penelitian ini berupa laptop Vivobook 14 A416 dengan spesifikasi berikut:

Tabel 3. 1 Spesifikasi perangkat keras

Komponen	Spesifikasi
Processor	Intel Core i3-1005G1
RAM	8 GB
Graphics Processor Unit	Intel UHD Graphics
Storage	256 GB SSD

2. Perangkat Lunak

Dalam penelitian ini, perangkat lunak yang digunakan meliputi sistem operasi dan bahasa pemrograman Python sebagai dasar pengembangan model. Selain itu, digunakan beberapa pustaka pendukung untuk implementasi model Long Short-Term Memory (LSTM). Spesifikasi rinci perangkat lunak yang digunakan beserta versinya disajikan sebagai berikut:

Tabel 3. 2 Spesifikasi Perangkat Lunak

Perangkat Lunak	Versi
Windows 10	22H2
Google Chrome	144.0.7559.97
Google Colab	Cloud based
Python	3.12
Tensorflow Keras	2.18

Scikit-learn	1.5
--------------	-----

Berdasarkan Tabel 3.2, sistem operasi Windows digunakan sebagai sistem utama dalam menjalankan seluruh proses penelitian. Google Chrome digunakan sebagai peramban web untuk mengakses layanan Google Colaboratory, yang dimanfaatkan sebagai lingkungan pengembangan berbasis cloud untuk menjalankan dan mengelola proses pemodelan. Penggunaan Google Colaboratory memberikan kemudahan dalam akses, pengelolaan pustaka, serta eksekusi kode Python tanpa memerlukan konfigurasi lingkungan secara lokal. Pada lingkungan tersebut, bahasa pemrograman Python digunakan sebagai dasar pengembangan model. Selanjutnya, pustaka TensorFlow (Keras) dimanfaatkan dalam pembangunan model Long Short-Term Memory (LSTM), sedangkan pustaka Scikit-learn digunakan untuk mendukung proses preprocessing data serta evaluasi kinerja model.

3.3 Waktu, Lokasi, atau Objek Penelitian

Penelitian ini dilaksanakan mulai bulan juli 2025. Objek penelitian ini berupa dataset transaksi kartu kredit yang di peroleh dari platform kaggle. Dataset tersebut dilabeli sebagai transaksi sah (non-froud) dan transaksi penipuan (fraud). Data ini di gunakan untuk merancang,melatih dan menguji model deteksi penipuan berbasis *Long Short-Term Memory*

3.4 Populasi dan Sampel

Dalam penelitian ini, populasi yang dimaksud adalah seluruh transaksi kartu kredit yang terjadi dalam sistem keuangan digital dan berpotensi mengandung aktivitas penipuan. Karena cakupan data transaksi dalam populasi

riil sangat luas dan tidak seluruhnya dapat diakses, maka penelitian ini menggunakan pendekatan data sekunder yang telah tersedia secara publik melalui platform Kaggle dan dapat mewakili karakteristik populasi tersebut.

Adapun sampel dalam penelitian ini adalah *European Credit Card Fraud Dataset* 2023, yang merupakan dataset standar dan banyak digunakan dalam penelitian akademik terkait deteksi penipuan kartu kredit. Dataset ini berisi sekitar 550.000 data transaksi kartu kredit yang dilakukan oleh pemegang kartu dari wilayah Eropa pada tahun 2023. Seluruh data telah dianonimkan untuk menjaga privasi pengguna namun tetap mempertahankan struktur dan pola yang relevan untuk analisis. Pemilihan dataset ini didasarkan pada pertimbangan sebagai berikut:

- a. Dataset bersifat *open-source* dan tersedia secara bebas untuk keperluan penelitian.
- b. Data telah melalui proses *anonymization* untuk melindungi identitas pengguna namun tetap mempertahankan pola transaksi yang valid.
- c. Dataset memuat fitur-fitur penting yang mendukung pembangunan model deteksi berbasis *deep learning* (LSTM), seperti jumlah transaksi, waktu, serta parameter numerik hasil transformasi PCA (*Principal Component Analysis*).

Teknik pengambilan sampel yang digunakan adalah purposive sampling, yaitu pemilihan data yang relevan berdasarkan kriteria tertentu, yakni data transaksi kartu kredit yang telah dilabeli secara jelas sebagai *fraud* dan *nonfraud*. Dengan teknik ini, peneliti dapat fokus pada segmentasi data yang paling relevan untuk menjawab tujuan penelitian, terutama dalam

pembangunan dan evaluasi model deteksi penipuan berbasis Long Short-Term Memory (LSTM).

Karena proporsi data fraud yang sangat kecil, peneliti juga menerapkan teknik penyeimbangan data, yaitu SMOTE (Synthetic Minority Over-sampling Technique), guna meningkatkan representasi data kelas minoritas (fraud) dalam proses pelatihan model.

Dengan demikian, pemilihan sampel yang terfokus dan representatif diharapkan mampu memberikan hasil yang lebih akurat dalam mengevaluasi efektivitas model deteksi penipuan serta memperkuat generalisasi temuan ke dalam konteks sistem keuangan digital secara lebih luas.

3.5 Instrumen Penelitian

Instrumen penelitian merupakan alat yang digunakan untuk mengumpulkan, mengolah, dan menganalisis data agar dapat menjawab rumusan masalah yang telah ditetapkan. Dalam penelitian ini, instrumen yang digunakan bersifat komputasional dan digital, sejalan dengan pendekatan eksperimen berbasis data.

Adapun instrumen utama yang digunakan dalam penelitian ini adalah sebagai berikut:

1. Bahasa Pemrograman dan Pustaka Penelitian ini menggunakan bahasa pemrograman Python karena fleksibel dan memiliki banyak pustaka (library) untuk pengolahan data dan implementasi model machine learning serta deep learning. Pustaka yang digunakan meliputi:
 - a. NumPy dan Pandas untuk manipulasi data dan analisis awal.
 - b. Matplotlib dan Seaborn untuk visualisasi data dan hasil model.

- c. Scikit-learn untuk preprocessing data dan teknik evaluasi model
- d. TensorFlow dan Keras untuk membangun dan melatih model *Long Short-Term Memory* (LSTM).

2. Metrik Evaluasi Model Untuk mengukur performa model deteksi penipuan, digunakan metrik-metrik berikut:

- a. Accuracy: mengukur keseluruhan tingkat prediksi yang benar.
- b. Precision: mengukur akurasi prediksi untuk kelas penipuan.
- c. Recall: mengukur seberapa banyak kasus penipuan yang berhasil terdeteksi.
- d. F1-Score: harmonisasi antara precision dan recall.
- e. AUC-ROC: mengukur kemampuan model dalam membedakan antara transaksi *fraud* dan *non-fraud*.

Penggunaan instrumen-instrumen di atas memungkinkan peneliti untuk melakukan serangkaian tahapan, mulai dari eksplorasi dan prapemrosesan data hingga pembangunan, pelatihan, evaluasi, dan analisis performa model LSTM secara menyeluruh. Dengan demikian, hasil penelitian yang diperoleh dapat dipertanggungjawabkan secara ilmiah dan akurat.

3.6 Teknik Pengumpulan Data

Pengumpulan data dalam penelitian ini dilakukan dengan menggunakan data sekunder yang tersedia secara publik dan dapat diakses secara bebas untuk keperluan penelitian akademik. Teknik ini dipilih karena fokus penelitian adalah pada pembangunan dan evaluasi model deteksi penipuan berbasis *Long Short-Term Memory* (LSTM), sehingga membutuhkan dataset transaksi yang telah

terstruktur dan dilabeli.

a. Sumber Data

Data yang digunakan adalah European Credit Card Fraud Dataset 2023, yang tersedia di platform Kaggle dan banyak digunakan dalam penelitian akademik. Dataset ini memuat data transaksi kartu kredit yang dilakukan oleh pemegang kartu dari Eropa pada tahun 2023. Dataset ini berisi lebih dari 550.000 data transaksi yang telah dianonimkan untuk melindungi identitas pemegang kartu.

Dataset ini memiliki karakteristik sebagai berikut:

- 1) Jumlah total transaksi: 568.630
- 2) Jumlah transaksi *fraud*: 284.315 ($\pm 50.0\%$)
- 3) Jumlah fitur: 30, terdiri dari:
 - a) 28 fitur hasil transformasi PCA (V1 sampai V28)
 - b) 1 fitur jumlah (*Amount*)
 - c) 1 fitur waktu (*Id*)
- 4) Label target: Class (0 = non-fraud, 1 = fraud)

b. Prosedur Pengumpulan Data

Langkah-langkah pengumpulan data meliputi:

- 1) Mengunduh dataset dari platform Kaggle

(<https://www.kaggle.com/datasets/nelgiriyeewithana/credit-cardfraud-detection-dataset-2023>)

- 2) Menyimpan dan memuat dataset ke dalam lingkungan pemrograman Python (menggunakan Jupyter Notebook).

- 3) Eksplorasi awal untuk memahami distribusi data, identifikasi jumlah kasus fraud dan non-fraud, serta memastikan tidak ada nilai hilang (*missing values*).

c. Teknik Prapemrosesan

Setelah data dikumpulkan, dilakukan prapemrosesan untuk memastikan data siap digunakan dalam pelatihan model. Langkah-langkahnya meliputi:

- 1) Normalisasi fitur numerik yaitu *Amount* dilakukan menggunakan teknik *standardization* dengan *StandardScaler*. *StandardScaler* menormalisasi data dengan cara menghitung nilai rata-rata (*mean*) dan simpangan baku (*standard deviation*) dari data, kemudian menskalakan setiap nilai data berdasarkan nilai tersebut (Castello Purba & Handhayani, 2024).
- 2) Membagi data menjadi data latih dan data uji dengan rasio 80:20. Pemilihan rasio ini didasarkan pada Prinsip Pareto, di mana sebagian besar data digunakan untuk pelatihan guna membantu model mempelajari pola secara optimal, sedangkan sebagian kecil data digunakan untuk pengujian guna mengevaluasi kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya. Rasio ini dipilih karena mampu memberikan keseimbangan antara data latih yang memadai dan data uji yang representatif (Maftucha et al., 2023).
- 3) Konversi data ke dalam format tiga dimensi dilakukan agar sesuai dengan format input yang dibutuhkan oleh arsitektur LSTM. Proses ini bertujuan untuk menyesuaikan struktur data fitur ke dalam bentuk

(*samples, timesteps, features*) sehingga dapat diproses oleh model LSTM pada tahap pelatihan dan pengujian. Pada penelitian ini, setiap sampel data diperlakukan sebagai satu urutan tunggal, sehingga nilai timesteps ditetapkan sebesar 1 tanpa mengubah informasi pada masing-masing fitur (Kurniansyah et al., 2025).

Pengumpulan dan prapemrosesan data dilakukan secara sistematis agar data yang digunakan benar-benar relevan dan sesuai dengan kebutuhan model deteksi berbasis deep learning. Hal ini penting agar hasil yang diperoleh tidak bias dan dapat menggambarkan performa model secara objektif.

3.7 Analisis Data

Analisis data dalam penelitian ini bertujuan untuk mengevaluasi efektivitas model Long Short-Term Memory (LSTM) dalam mendeteksi transaksi penipuan pada kartu kredit. Proses analisis dilakukan secara sistematis melalui beberapa tahapan, mulai dari pelatihan model, pengujian, hingga evaluasi performa berdasarkan metrik yang relevan.

1. Proses Analisis

a. Pelatihan dan Pengujian Model

- 1) Dataset dibagi menjadi data latih dan data uji dengan proporsi 80:20.
- 2) Model Long Short-Term Memory (LSTM) pada penelitian ini dibangun menggunakan arsitektur jaringan saraf berlapis (layers). Struktur model terdiri dari LSTM layer yang berfungsi mempelajari pola pada data transaksi, dropout layer yang digunakan untuk mengurangi resiko overfitting selama proses

pelatihan, serta dense layer yang berperan sebagai lapisan klasifikasi untuk menghasilkan prediksi transaksi fraud dan non-fraud.

- 3) Model dilatih menggunakan fungsi aktivasi sigmoid dan binary cross-entropy sebagai *loss function* karena klasifikasi bersifat biner (fraud vs non-fraud).
- 4) Optimizer yang digunakan adalah Adam karena efisien untuk data besar dan mendukung *adaptive learning rate*.

b. Evaluasi Kinerja Model

- 1) Evaluasi kinerja model dilakukan untuk mengetahui kemampuan model LSTM dalam mengklasifikasikan transaksi fraud dan non-fraud. Pada penelitian ini, performa model diukur menggunakan beberapa metrik evaluasi yaitu accuracy, precision, recall, F1-score, AUC-ROC, serta confusion matrix. Metrik-metrik tersebut digunakan untuk menilai tingkat ketepatan model dalam mendeteksi transaksi penipuan berdasarkan hasil prediksi terhadap data uji.

c. Evaluasi Akurasi Model LSTM

- a) Tahap evaluasi dilakukan untuk mengetahui kinerja model **Long Short-Term Memory (LSTM)** dalam mendeteksi transaksi penipuan pada kartu kredit. Proses ini bertujuan untuk menilai sejauh mana model yang telah dilatih mampu membedakan antara transaksi penipuan (*fraud*) dan transaksi normal (*non-fraud*). Evaluasi dilakukan dengan menggunakan

data uji (*testing data*) yang sebelumnya tidak digunakan pada proses pelatihan model.

Dalam penelitian ini, pengukuran kinerja model dilakukan dengan menggunakan beberapa metrik evaluasi yang umum digunakan dalam permasalahan klasifikasi, yaitu **accuracy**, **precision**, **recall**, dan **F1-score**. Nilai dari metrik tersebut diperoleh berdasarkan hasil **confusion matrix**, yaitu tabel yang menggambarkan perbandingan antara hasil prediksi model dengan kondisi sebenarnya pada data. Confusion matrix terdiri dari empat komponen utama, yaitu **True Positive (TP)**, **True Negative (TN)**, **False Positive (FP)**, dan **False Negative (FN)**. True Positive (TP) merupakan jumlah transaksi penipuan yang berhasil diprediksi dengan benar sebagai penipuan oleh model. True Negative (TN) merupakan jumlah transaksi normal yang berhasil diprediksi dengan benar sebagai transaksi normal. False Positive (FP) merupakan jumlah transaksi normal yang salah diprediksi sebagai transaksi penipuan. Sedangkan False Negative (FN) merupakan jumlah transaksi penipuan yang salah diprediksi sebagai transaksi normal.

Berdasarkan komponen tersebut, perhitungan metrik evaluasi dilakukan menggunakan persamaan sebagai berikut.

1. Accuracy

Accuracy digunakan untuk mengetahui tingkat ketepatan model dalam mengklasifikasikan seluruh data yang diuji.

Nilai accuracy menunjukkan perbandingan antara jumlah prediksi yang benar dengan keseluruhan data yang digunakan dalam pengujian.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

2. Precision

Precision digunakan untuk mengukur ketepatan model dalam memprediksi transaksi yang diklasifikasikan sebagai penipuan.

$$Precision = \frac{TP}{TP + FP}$$

Semakin tinggi nilai precision, maka semakin banyak transaksi yang diprediksi sebagai penipuan benar-benar merupakan transaksi penipuan.

3. Recall

Recall digunakan untuk melihat kemampuan model dalam mendeteksi seluruh transaksi penipuan yang terdapat pada data uji.

$$Recall = \frac{TP}{TP + FN}$$

Nilai recall yang tinggi menunjukkan bahwa model mampu mengidentifikasi sebagian besar transaksi penipuan yang terjadi.

4. F1-Score

F1-score merupakan nilai yang diperoleh dari kombinasi antara precision dan recall yang digunakan untuk melihat keseimbangan kinerja model.

Dengan menggunakan beberapa metrik evaluasi tersebut, kinerja model LSTM dapat dianalisis secara lebih menyeluruh dalam mengidentifikasi transaksi penipuan dan transaksi normal.

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



BAB IV

HASIL DAN PEMBAHASAN

4.1 Pengumpulan Data

Pengumpulan data pada penelitian ini dilakukan dengan memanfaatkan dataset sekunder yang tersedia secara publik melalui platform Kaggle. Dataset yang digunakan adalah Credit Card Fraud Detection Dataset 2023 yang dikembangkan oleh Nelgiryewithana. Dataset ini dipilih karena memiliki struktur data yang lengkap, telah melalui proses anonymization, serta memuat informasi transaksi kartu kredit yang relevan untuk proses pendeteksian penipuan. Data terdiri atas rekaman transaksi yang telah diberi label sebagai fraud atau non-fraud, sehingga sesuai untuk proses klasifikasi menggunakan algoritma Long Short-Term Memory (LSTM).

Proses pengunduhan dan pemanggilan dataset dilakukan secara langsung melalui pustaka kagglehub, yang memungkinkan integrasi otomatis antara lingkungan pemrograman dan repositori dataset Kaggle. Melalui pendekatan ini, peneliti dapat memastikan bahwa versi dataset yang digunakan adalah versi terbaru dan valid. Pada tahap ini, file `creditcard_2023.csv` diakses dan dimuat ke dalam lingkungan Python menggunakan `KaggleDatasetAdapter.PANDAS`, kemudian ditampilkan sebagian untuk memastikan bahwa data telah terunduh dan terbaca dengan benar. Proses pemanggilan ini ditunjukkan pada potongan kode berikut:

```
file_path = "creditcard_2023.csv"
```

```
df = kagglehub.load_dataset(
    KaggleDatasetAdapter.PANDAS,
    "nelgiriwithana/credit-card-fraud-detection-dataset-2023",
    file_path) print("First 5 records:", df.head())
```

Tahap pemeriksaan awal dilakukan dengan menampilkan lima data pertama menggunakan fungsi `df.head()`. Langkah ini penting untuk memastikan bahwa seluruh kolom, tipe data, serta nilai-nilai awal telah termuat dengan baik sebelum dilakukan proses prapemrosesan. Selain itu, prosedur ini memberikan gambaran awal mengenai struktur dan karakteristik dataset, termasuk jumlah fitur dan distribusi kelas. Berikut 5 data pertama yang ditampilkan:

Tabel 4. 1 5 Data Pertama

id	V1	V2	V3	V4	V5	V6	V7	V8	V9
0	-0.260	-0.469	2.496	-0.083	0.129	0.732	0.519	-0.130	0.727
1	0.985	-0.356	0.558	-0.429	0.277	0.428	0.406	-0.133	0.347
2	-0.260	-0.949	1.728	-0.457	0.074	1.419	0.743	-0.095	-0.261
3	-0.152	-0.508	1.746	-1.090	0.249	1.143	0.518	-0.065	-0.205
4	-0.206	-0.165	1.527	-0.44	0.106	0.530	0.658	-0.212	1.049

V21	V22	V23	V24	V25	V26	V27	V28	Amount	Class
-0.110	0.217	-0.134	0.165	0.126	-0.434	-0.081	-0.151	17982.10	0
-0.194	-0.605	0.079	-0.577	0.190	0.296	-0.248	-0.064	6531.37	0
-0.005	0.702	0.945	-1.154	-0.605	-0.312	-0.300	-0.244	2513.54	0
-0.146	-0.038	-0.214	-1.893	1.003	-0.515	-0.165	0.048	5384.44	0
-0.106	0.729	-0.161	0.312	-0.414	1.126	0.023	0.419	14278.97	0

Tabel hasil tampilan lima data teratas tersebut menunjukkan struktur awal dari dataset yang digunakan dalam penelitian. Setiap baris

merepresentasikan satu transaksi kartu kredit dengan sejumlah fitur numerik yang disandikan dalam bentuk komponen V1 hingga V28, yang merupakan hasil transformasi Principal Component Analysis (PCA). Transformasi ini dilakukan oleh pengembang dataset untuk menjaga kerahasiaan informasi sensitif tanpa menghilangkan pola statistik yang penting bagi proses analisis. Selain fitur-fitur PCA, dataset juga memuat atribut Amount yang menggambarkan nilai nominal transaksi, serta atribut Class yang berfungsi sebagai label penanda apakah transaksi tersebut tergolong penipuan ($\text{fraud} = 1$) atau transaksi normal ($\text{non-fraud} = 0$). Berdasarkan sampel awal yang ditampilkan, seluruh transaksi dalam lima baris pertama berlabel 0, yang berarti tidak termasuk kategori penipuan.

Selain menampilkan lima baris pertama, dilakukan pula pemeriksaan struktur keseluruhan dataset menggunakan fungsi `df.info()`. Berikut potongan kode pemeriksaan struktur data:

```
print("\nDataset Info:") print(df.info())
```

Output dari perintah ini memberikan informasi penting mengenai jumlah baris, jumlah kolom, tipe data setiap kolom, serta ketersediaan nilai pada masing-masing fitur. Berikut output dari potongan kode diatas:

```

Dataset Info:
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 568630 entries, 0 to 568629
Data columns (total 31 columns):
#   Column      Non-Null Count  Dtype
---  -
0   id           568630 non-null  int64
1   V1           568630 non-null  float64
2   V2           568630 non-null  float64
3   V3           568630 non-null  float64
4   V4           568630 non-null  float64
5   V5           568630 non-null  float64
6   V6           568630 non-null  float64
7   V7           568630 non-null  float64
8   V8           568630 non-null  float64
9   V9           568630 non-null  float64
10  V10          568630 non-null  float64
11  V11          568630 non-null  float64
12  V12          568630 non-null  float64
13  V13          568630 non-null  float64
14  V14          568630 non-null  float64
15  V15          568630 non-null  float64
16  V16          568630 non-null  float64
17  V17          568630 non-null  float64
18  V18          568630 non-null  float64
19  V19          568630 non-null  float64
20  V20          568630 non-null  float64
21  V21          568630 non-null  float64
22  V22          568630 non-null  float64
23  V23          568630 non-null  float64
24  V24          568630 non-null  float64
25  V25          568630 non-null  float64
26  V26          568630 non-null  float64
27  V27          568630 non-null  float64
28  V28          568630 non-null  float64
29  Amount       568630 non-null  float64
30  Class        568630 non-null  int64
dtypes: float64(29), int64(2)
memory usage: 134.5 MB

```

Gambar 4.1 Struktur Dataset

Berdasarkan hasil yang ditampilkan, dataset memuat 568.630 entri dengan total 31 kolom, yang seluruhnya terisi penuh tanpa adanya nilai yang hilang (*non-null*). Hal ini menunjukkan bahwa dataset berada dalam kondisi bersih (*clean*) pada tahap awal sehingga tidak memerlukan proses imputasi data. Sebagian besar kolom memiliki tipe data float64, khususnya pada fitur V1 sampai V28 yang merupakan hasil transformasi PCA. Sementara itu, kolom id dan Class bertipe int64, masing-masing berfungsi sebagai penanda urutan transaksi dan label klasifikasi (0 untuk *non-fraud*, 1 untuk *fraud*). Informasi mengenai tipe data ini penting untuk memastikan bahwa setiap fitur memiliki format yang sesuai untuk proses pemodelan, terutama ketika dilakukan

normalisasi dan pembentukan urutan data (*sequence*) untuk kebutuhan algoritma LSTM.

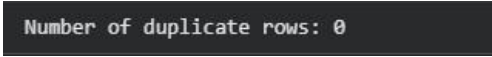
Tahap selanjutnya adalah melakukan pemeriksaan terhadap keberadaan nilai hilang (*missing values*) dan duplikasi data untuk memastikan kualitas dataset sebelum memasuki proses prapemrosesan lebih lanjut. Berikut potongan kode yang dijalankan:

```
missing_values = df.isnull().sum() print("Number of missing
values per column") print(missing_values) print(f"Number of
duplicate rows: {df.duplicated().sum()}")
```

Pemeriksaan dilakukan menggunakan perintah `df.isnull().sum()` untuk menghitung jumlah nilai kosong pada setiap kolom, serta `df.duplicated().sum()` untuk mengidentifikasi baris yang terduplikasi. Potongan kode tersebut menghasilkan output berikut:

```
Number of missing values per column
id      0
V1      0
V2      0
V3      0
V4      0
V5      0
V6      0
V7      0
V8      0
V9      0
V10     0
V11     0
V12     0
V13     0
V14     0
V15     0
V16     0
V17     0
V18     0
V19     0
V20     0
V21     0
V22     0
V23     0
V24     0
V25     0
V26     0
V27     0
V28     0
Amount  0
Class   0
```

Gambar 4. 2 Jumlah Missing Value


 Number of duplicate rows: 0

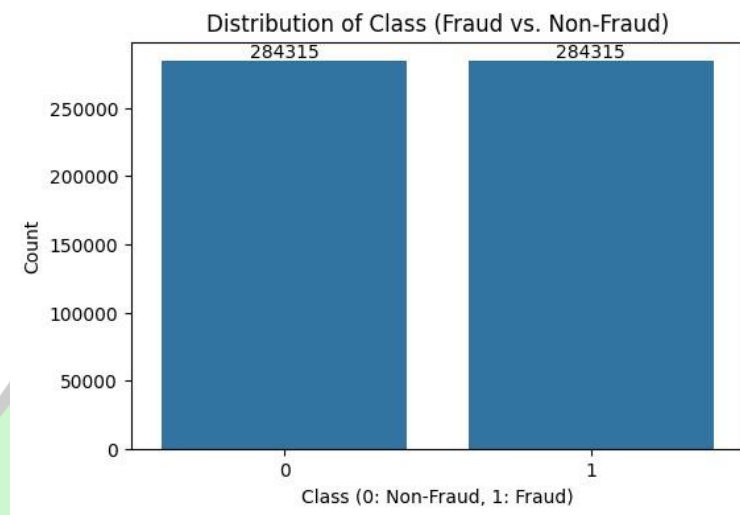
Gambar 4. 3 Jumlah Duplicate Rows

Hasil analisis menunjukkan bahwa seluruh kolom dalam dataset memiliki jumlah missing value sebesar 0, yang berarti tidak terdapat data yang hilang baik pada fitur-fitur PCA (V1–V28), kolom Amount, maupun label Class. Kondisi ini menguntungkan karena dataset tidak memerlukan proses imputasi atau pembersihan tambahan untuk menangani data yang tidak lengkap. Selain itu, pemeriksaan terhadap duplikasi baris juga menghasilkan nilai 0, yang menandakan bahwa tidak terdapat entri transaksi yang tercatat lebih dari satu kali. Ketiadaan data duplikat penting untuk menjaga integritas data, menghindari bias dalam proses pelatihan model, serta memastikan bahwa setiap transaksi memiliki representasi yang unik. Dengan demikian, hasil pengecekan ini mengonfirmasi bahwa dataset berada dalam kondisi yang bersih dan konsisten, sehingga dapat langsung dilanjutkan ke tahapan selanjutnya tanpa memerlukan tindakan perbaikan struktural terhadap data.

Analisis distribusi kelas dilakukan untuk mengetahui proporsi antara transaksi normal (*non-fraud*) dan transaksi penipuan (*fraud*) dalam dataset. Untuk melakukan analisis tersebut, dilakukan penggunaan potongan kode dibawah ini:

```
plt.figure(figsize=(6, 4)) ax = sns.countplot(x='Class',
data=df) plt.title('Distribution of Class (Fraud vs.
Non-Fraud)') plt.xlabel('Class (0: Non-Fraud, 1:
Fraud)') plt.ylabel('Count') for container in
ax.containers:
    ax.bar_label(container, fmt='%d') plt.show()
```

Visualisasi menggunakan countplot dari pustaka Seaborn menunjukkan jumlah masing-masing kelas secara jelas dengan bantuan label numerik pada bagian atas batang grafik. Berikut output hasil visualisasi dari potongan kode tersebut:



Gambar 4. 4 Distribusi Kelas

Hasil visualisasi memperlihatkan bahwa dataset memiliki distribusi kelas yang seimbang, yaitu masing-masing sebanyak 284.315 transaksi untuk kelas 0 (non-fraud) dan kelas 1 (fraud). Kondisi distribusi yang seimbang seperti ini jarang ditemukan pada kasus deteksi penipuan di dunia nyata, karena umumnya proporsi transaksi penipuan jauh lebih kecil. Namun, jenis dataset tertentu termasuk dataset hasil rekonstruksi atau augmentasi dari Kaggle memang telah melalui proses penyeimbangan kelas oleh penyusun dataset agar lebih sesuai untuk eksperimen model pembelajaran mesin. Dengan adanya keseimbangan jumlah sampel pada kedua kelas, proses pelatihan model menjadi lebih stabil karena model tidak cenderung bias terhadap salah satu kelas. Distribusi yang setara ini juga menghilangkan kebutuhan untuk menerapkan teknik penanganan ketidakseimbangan data seperti SMOTE atau undersampling. Dengan demikian, dataset siap digunakan untuk tahap

prapemrosesan dan pelatihan model LSTM tanpa memerlukan penyesuaian tambahan pada proporsi kelas.

4.2 Preprocessing data

a) Normalisasi Data

Prapemrosesan data merupakan tahap penting sebelum data digunakan dalam proses pelatihan model, terutama pada metode berbasis *deep learning* seperti Long Short-Term Memory (LSTM). Tahap ini bertujuan untuk memastikan bahwa seluruh fitur berada dalam format yang konsisten, memiliki skala yang sesuai, dan siap diolah oleh model. Salah satu aspek krusial dalam prapemrosesan adalah penanganan perbedaan skala antar fitur, karena perbedaan nilai yang terlalu jauh dapat menyebabkan model kesulitan dalam mengenali pola secara optimal.

Pada penelitian ini, langkah prapemrosesan dimulai dengan melakukan normalisasi terhadap fitur Amount, yaitu fitur yang merepresentasikan nilai nominal transaksi. Fitur ini memiliki rentang nilai yang relatif besar dan bervariasi dibandingkan fitur lain yang telah melalui transformasi PCA (V1–V28). Oleh karena itu, normalisasi diperlukan untuk menempatkan fitur Amount pada skala yang sama agar tidak mendominasi proses pelatihan model. Proses normalisasi dilakukan menggunakan metode StandardScaler dari pustaka *scikit-learn*, yang mentransformasikan nilai menjadi distribusi dengan rata-rata 0 dan standar deviasi 1. Kode berikut digunakan untuk melakukan normalisasi terhadap fitur Amount:

```
features_to_scale = ['Amount'] scaler =
StandardScaler() df_scaled_amount =
pd.DataFrame(scaler.fit_transform(
df[features_to_scale]), columns=features_to_scale)
```

```
df_scaled = pd.concat([df['id'], df.drop(columns=['id',
'Amount', 'Class'])], df_scaled_amount, df['Class']], axis=1)
```

Berikut 5 data teratas yang telah di normalisasi:

Tabel 4. 2 5 Data Teratas Sebelum dan Setelah Normalisasi

NO	Amount Sebelum Normalisasi	Amount Setelah Normalisasi
1	17982.10	0.858447
2	6531.37	-0.796369
3	2513.54	-1.377011
4	5384.44	-0.962119
5	14278.97	0.323285

Hasil transformasi menunjukkan bahwa fitur Amount telah berada pada skala yang lebih terkontrol dan siap digunakan dalam proses pelatihan. Nilai-nilai yang sebelumnya berada dalam rentang besar kini telah dinormalkan tanpa mengubah struktur hubungan antar data. Proses ini tidak hanya meningkatkan stabilitas pelatihan model, tetapi juga memastikan bahwa algoritma LSTM dapat mempelajari pola transaksi berdasarkan seluruh fitur secara seimbang.

b) Memisahkan Variabel Fitur dan Variabel Target

Setelah proses normalisasi dilakukan, tahap prapemrosesan dilanjutkan dengan memisahkan variabel fitur (features) dan variabel target (label). Variabel target dalam penelitian ini adalah kolom Class, yang menunjukkan apakah suatu transaksi tergolong penipuan (fraud = 1) atau bukan (non-fraud = 0). Sementara itu, seluruh kolom numerik lainnya digunakan sebagai fitur yang merepresentasikan karakteristik setiap transaksi. Berikut kode yang digunakan pada tahap ini:

```
X = df_scaled.drop(columns=['id', 'Class']) y
= df_scaled['Class']
```

Pada tahap ini, kolom id tidak disertakan sebagai fitur karena hanya berfungsi sebagai penanda urutan transaksi dan tidak memiliki kontribusi langsung terhadap proses klasifikasi.

c) Konversi Data ke Bentuk 3 Dimensi

Setelah fitur dan target dipisahkan, dilakukan proses konversi atau reshaping terhadap matriks fitur agar sesuai dengan format input yang dibutuhkan oleh model LSTM. Berbeda dengan algoritma pembelajaran mesin tradisional, LSTM memerlukan struktur data dalam bentuk tiga dimensi, yaitu (samples, timesteps, features). Berikut kode yang digunakan dalam proses tersebut:

```
X_reshaped = X.values.reshape(X.shape[0], 1, X.shape[1])
```

Pada setiap transaksi pada dataset ini diperlakukan sebagai satu urutan tunggal tanpa rentang waktu tambahan, nilai timesteps ditetapkan sebesar 1. Dengan demikian, seluruh sampel data direstrukturisasi ke dalam bentuk 3D array sehingga dapat diproses oleh jaringan LSTM tanpa mengubah informasi pada masing-masing fitur.

d) Pembagian Data

Langkah berikutnya adalah melakukan pembagian dataset menjadi data latih (training set) dan data uji (testing set) dengan proporsi 80:20 menggunakan fungsi `train_test_split`. Kode yang digunakan sebagai berikut:

```
X_train, X_test, y_train, y_test = train_test_split(
    X_reshaped, y,      test_size=0.2,      stratify=y,
    random_state=42)
```

Pada kode tersebut, teknik stratified sampling diterapkan melalui parameter `stratify=y` untuk memastikan bahwa distribusi kelas pada data latih dan data uji tetap seimbang. Pendekatan ini penting untuk menjaga konsistensi performa model dan mencegah bias selama proses pelatihan. Dari pembagian data tersebut, berikut distribusi tiap data:

```
Shape of X_train: (454904, 1, 29)
Shape of X_test: (113726, 1, 29)
Shape of y_train: (454904,)
Shape of y_test: (113726,)

Class distribution in original data:
Class
0    0.5
1    0.5
Name: proportion, dtype: float64

Class distribution in y_train:
Class
0    0.5
1    0.5
Name: proportion, dtype: float64

Class distribution in y_test:
Class
1    0.5
0    0.5
Name: proportion, dtype: float64
```

Gambar 4. 5 Distribusi Kelas

Dari pembagian data menggunakan teknik stratified sampling, diperoleh hasil bahwa distribusi kelas pada data latih dan data uji tetap seimbang, sama seperti pada data asli. Gambar diatas memperlihatkan bahwa baik data latih maupun data uji masing-masing memiliki proporsi 50% kelas non-fraud dan 50% kelas fraud, sehingga tidak terjadi bias akibat ketidakseimbangan kelas. Dengan demikian, proses pelatihan dapat berlangsung secara lebih stabil, karena model menerima distribusi kelas yang representatif pada kedua bagian dataset. Selain itu, ukuran dataset yang cukup besar, yaitu 454.904 sampel untuk data latih dan 113.726 sampel untuk data uji, memastikan bahwa model memiliki cukup banyak data untuk mempelajari pola transaksi secara efektif sekaligus diuji secara akurat untuk menilai kemampuan generalisasi.

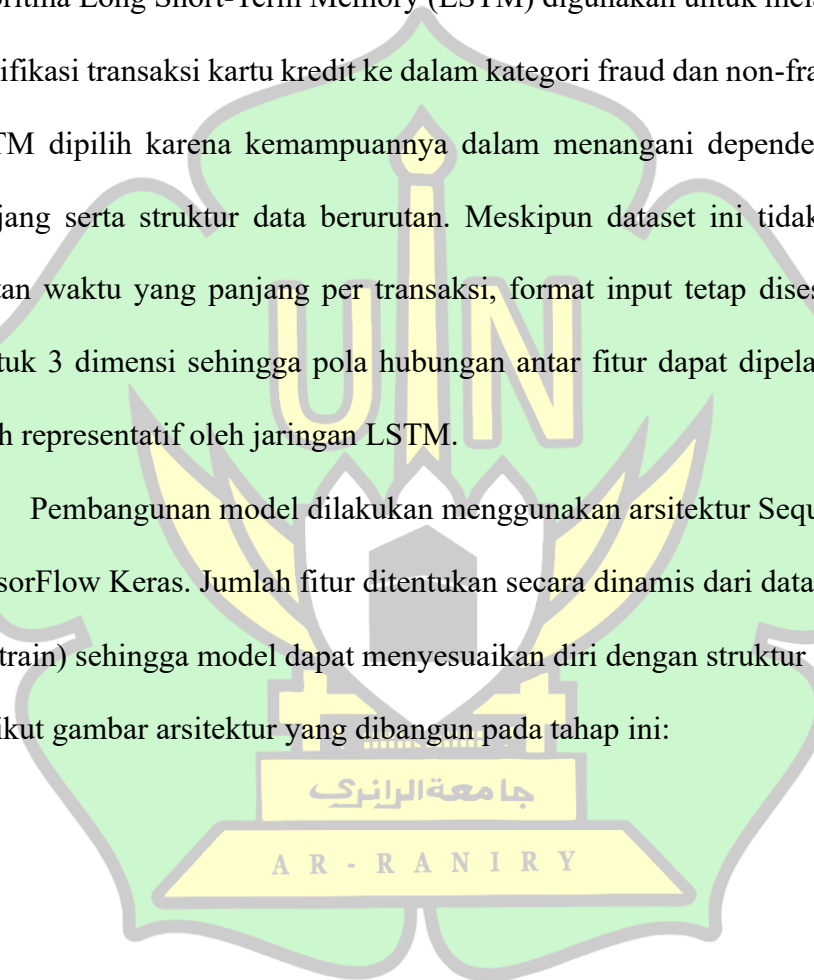
Dengan melalui tahapan tahapan ini, data telah berada dalam format yang sesuai untuk memasuki tahap pemodelan menggunakan arsitektur LSTM. Struktur data yang telah dibentuk memungkinkan model untuk mempelajari pola-pola kompleks dalam transaksi kartu kredit secara lebih efektif.

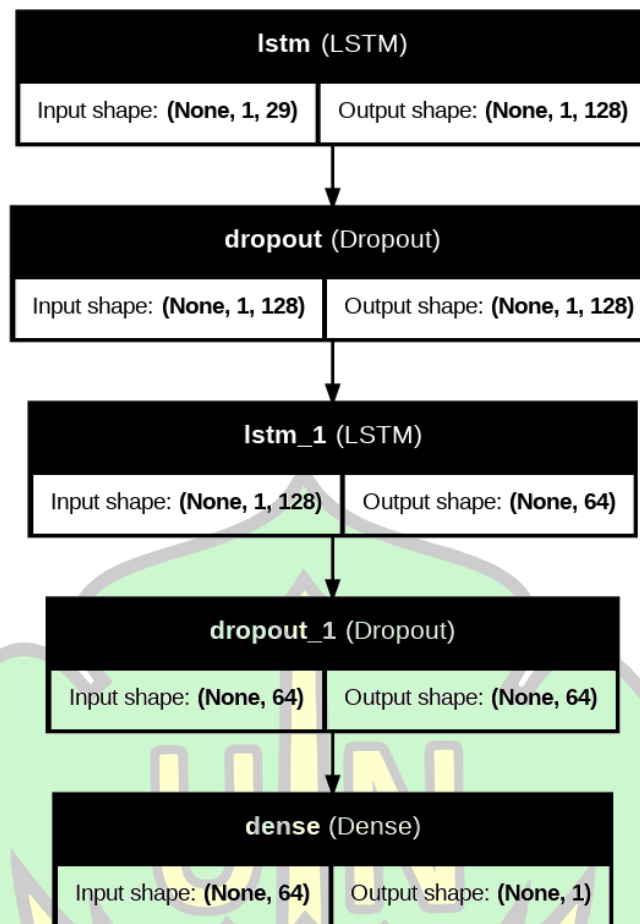
4.3 Pembangunan Model LSTM

Tahap pembangunan model merupakan inti dari penelitian ini, di mana algoritma Long Short-Term Memory (LSTM) digunakan untuk melakukan klasifikasi transaksi kartu kredit ke dalam kategori fraud dan non-fraud. Model LSTM dipilih karena kemampuannya dalam menangani dependensi jangka panjang serta struktur data berurutan. Meskipun dataset ini tidak memiliki urutan waktu yang panjang per transaksi, format input tetap disesuaikan ke bentuk 3 dimensi sehingga pola hubungan antar fitur dapat dipelajari secara lebih representatif oleh jaringan LSTM.

Pembangunan model dilakukan menggunakan arsitektur Sequential dari TensorFlow Keras. Jumlah fitur ditentukan secara dinamis dari data pelatihan (X_{train}) sehingga model dapat menyesuaikan diri dengan struktur input.

Berikut gambar arsitektur yang dibangun pada tahap ini:





Gambar 4. 6 Arsitektur Model

Gambar struktur model yang diatas memperlihatkan aliran data dari input hingga output, lengkap dengan bentuk (shape) input–output pada setiap lapisan. Arsitektur dimulai dengan sebuah lapisan LSTM berukuran 128 unit yang dikonfigurasi dengan `return_sequences=True` agar keluaran dapat diteruskan ke lapisan LSTM berikutnya. Lapisan ini diikuti oleh lapisan Dropout 0.2 untuk mengurangi risiko overfitting. Selanjutnya, model menambahkan lapisan LSTM kedua dengan 64 unit, disertai lapisan Dropout tambahan untuk menjaga generalisasi model. Setelah dua lapisan LSTM, model diakhiri dengan lapisan Dense berukuran 1 dengan fungsi aktivasi sigmoid untuk menghasilkan keluaran berupa probabilitas yang kemudian diterjemahkan menjadi dua kelas (0 atau 1). Konfigurasi ini dirancang secara

spesifik untuk tugas klasifikasi biner. Kode pembangunan model adalah sebagai berikut:

```
model = Sequential() model.add(LSTM(units=128,
return_sequences=True,
input_shape=(1, num_features)))
model.add(Dropout(0.2))
model.add(LSTM(units=64))
model.add(Dropout(0.2))
model.add(Dense(units=1, activation='sigmoid'))
```

Setelah arsitektur selesai disusun, model dikompilasi menggunakan optimizer Adam dengan learning rate 0.001. Fungsi loss yang digunakan adalah binary crossentropy, yang merupakan pilihan paling sesuai untuk klasifikasi biner. Selain itu, metrik evaluasi yang digunakan adalah accuracy untuk memantau performa model selama proses pelatihan.

```
model.compile(
optimizer=Adam(learning_rate=0.001),
loss='binary_crossentropy',
metrics=['accuracy'])
```

Arsitektur yang dibangun pada penelitian ini dirancang untuk menghasilkan keseimbangan antara kompleksitas model dan kemampuan generalisasi. Dengan dua lapisan LSTM dan mekanisme Dropout, model diharapkan mampu mempelajari pola transaksi secara mendalam sekaligus menghindari overfitting selama proses pelatihan.

Setelah arsitektur model LSTM dibangun dan dikompilasi, tahap selanjutnya adalah melakukan proses pelatihan (training) untuk memungkinkan model mempelajari pola transaksi yang membedakan antara aktivitas penipuan dan transaksi normal. Proses pelatihan dilakukan menggunakan data latih (X_{train} dan y_{train}) yang sebelumnya telah melalui tahap prapemrosesan dan reshaping. Pada tahap ini, model dilatih selama 50 epoch dengan ukuran batch sebesar 64, yang dipilih untuk memperoleh keseimbangan antara stabilitas proses pembelajaran dan efisiensi komputasi.

Selain itu, proses pelatihan memanfaatkan data uji (X_{test} dan y_{test}) sebagai validation data yang berfungsi untuk memantau performa model pada data yang tidak ikut dilatih. Pendekatan ini memungkinkan peneliti mengamati potensi terjadinya overfitting, yaitu kondisi ketika model terlalu menyesuaikan diri dengan data latih sehingga kinerjanya menurun pada data baru. Dengan adanya metrik validasi, perubahan nilai loss dan accuracy pada setiap epoch dapat dianalisis untuk menilai stabilitas pembelajaran model.

Agar proses pelatihan berlangsung lebih optimal, beberapa callbacks diterapkan untuk memonitor performa model secara otomatis. Pertama, ModelCheckpoint digunakan untuk menyimpan bobot model terbaik berdasarkan nilai validation accuracy. Dengan mekanisme ini, model akan selalu menyimpan versi dengan kinerja validasi tertinggi meskipun nilai akurasi berfluktuasi selama proses pelatihan. Kedua, EarlyStopping digunakan untuk menghentikan pelatihan lebih awal apabila tidak terjadi peningkatan performa pada data validasi dalam beberapa epoch berturut-turut. Pada konfigurasi ini, patience ditetapkan sebanyak 5 epoch, sehingga pelatihan akan berhenti ketika model mulai menunjukkan tanda-tanda overfitting. Ketiga, ReduceLROnPlateau diterapkan untuk mengurangi nilai learning rate secara

adaptif ketika validation loss tidak menunjukkan perbaikan dalam periode tertentu. Reduksi ini membantu model keluar dari kondisi stagnan (plateau) dan melanjutkan proses pembelajaran dengan kecepatan yang lebih sesuai.

Proses pelatihan dilakukan menggunakan kode berikut:

```
model.compile(optimizer=Adam(learning_rate=0.001),
model_checkpoint = ModelCheckpoint('/content/lstm_model.keras',
monitor='val_accuracy',save_best_only=True,          mode='max',
verbose=1)          early_stopping              =
EarlyStopping(monitor='val_accuracy',  patience=5,  verbose=1,
mode='max',restore_best_weights=True) model.ReduceLROnPlateau =
ReduceLROnPlateau(monitor='val_loss',  factor=0.1,  patience=10,
min_lr=0.000001)  history  =  model.fit(X_train,  y_train,
epochs=50, batch_size=64,
validation_data=(X_test,  y_test),  callbacks  =
[model_checkpoint, early_stopping, model.ReduceLROnPlateau])
```

Dengan konfigurasi ini, proses pelatihan menjadi lebih adaptif, efisien, dan terkendali. Penggunaan callbacks memastikan bahwa model tidak mengalami pelatihan berlebihan, dapat menyesuaikan tingkat pembelajaran secara dinamis, dan menghasilkan bobot terbaik yang optimal untuk tahap evaluasi berikutnya. Hasil pelatihan (training history) yang disimpan dalam variabel history digunakan pada tahap berikutnya untuk menampilkan kurva loss dan accuracy baik pada data latih maupun data validasi. Kurva tersebut akan memberikan gambaran mengenai dinamika proses pembelajaran model dan menjadi dasar dalam mengevaluasi apakah model telah mencapai kinerja yang optimal.

4.4 Evaluasi

a) Evaluasi Kinerja Model Saat Pelatihan

Evaluasi model pada saat pelatihan dilakukan dengan menganalisis perkembangan nilai accuracy dan loss pada data latih maupun data validasi selama beberapa epoch. Berikut tabel perkembangan akurasi dan loss tiap epoch :

Tabel 4. 3 Akurasi Setiap Iterasi (Epoch)

Epoch	Training Loss	Validation Loss	Training Accuracy	Validation Accuracy
1	0.0808	0.0166	0.9720	0.99462
2	0.0142	0.0080	0.9951	0.99743
3	0.0081	0.0047	0.9973	0.99887
4	0.0054	0.0049	0.9983	0.99860
5	0.0043	0.0039	0.9987	0.99892
6	0.0038	0.0025	0.9989	0.99964
7	0.0028	0.0032	0.9992	0.99920
8	0.0027	0.0022	0.9992	0.99960
9	0.0025	0.0028	0.9993	0.99920
10	0.0022	0.0024	0.9994	0.99940
11	0.0020	0.0019	0.9994	0.99940

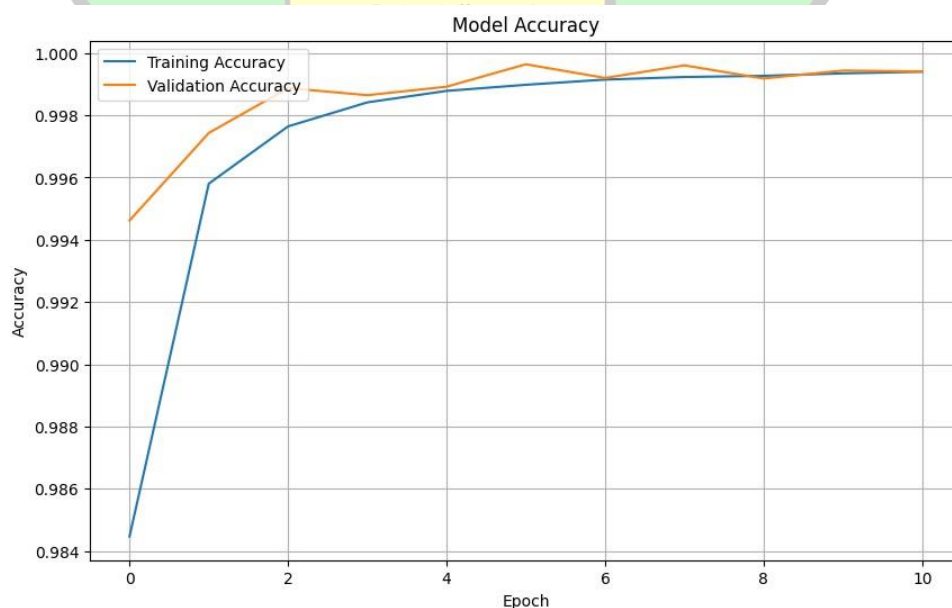
Berdasarkan hasil pelatihan, terlihat bahwa model menunjukkan peningkatan performa yang konsisten sejak epoch pertama. Pada epoch ke-1, model mencapai akurasi sebesar 0.9720 pada data latih dan 0.9946 pada data validasi. Nilai validation loss juga tercatat cukup rendah yaitu 0.0166, menandakan bahwa model mulai mampu mempelajari pola transaksi secara efektif sejak awal proses pelatihan.

Seiring bertambahnya epoch, performa model terus mengalami

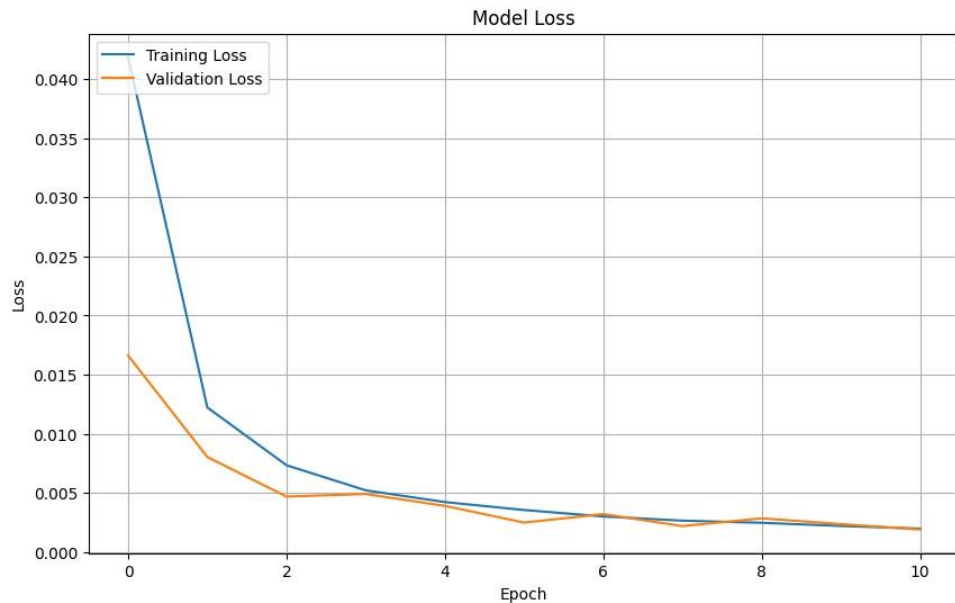
peningkatan yang signifikan. Pada epoch ke-3, akurasi latih meningkat menjadi 0.9973 dan akurasi validasi mencapai 0.9989. Penurunan loss terjadi secara progresif, baik pada data latih maupun validasi, hingga mencapai nilai di bawah 0.005, yang menunjukkan bahwa kesalahan prediksi semakin kecil. Peningkatan ini terus berlanjut hingga epoch ke-6, yang merupakan titik performa terbaik (best epoch) dengan nilai akurasi validasi mencapai 0.9996, disertai penurunan validation loss hingga 0.0025. Pada titik inilah bobot model disimpan oleh ModelCheckpoint karena menghasilkan performa validasi paling optimal.

Pada epoch ke-11, mekanisme EarlyStopping aktif karena tidak terjadi peningkatan akurasi validasi selama lima epoch berturut-turut. Sistem kemudian menghentikan pelatihan dan mengembalikan bobot model ke kondisi terbaik yang diperoleh pada epoch ke-6. Aktivasi EarlyStopping ini menunjukkan bahwa model telah mencapai titik optimum dan pelatihan tambahan tidak lagi memberikan manfaat yang signifikan.

Berikut gambar grafik perkembangan akurasi dan loss saat pelatihan model:



Gambar 4. 7 Grafik training dan validation accuracy



Gambar 4. 8 Grafik training loss dan validation loss

Grafik akurasi pada gambar diatas menunjukkan bahwa baik training accuracy maupun validation accuracy meningkat secara stabil di setiap epoch. Kurva validasi tampak berada sedikit di atas kurva pelatihan pada sebagian besar epoch awal, yang mengindikasikan bahwa model mampu melakukan generalisasi dengan baik pada data yang tidak terlihat selama pelatihan. Seiring meningkatnya epoch, kedua kurva bergerak sangat berdekatan mendekati nilai 1.0, menandakan bahwa model memperoleh tingkat akurasi yang sangat tinggi dan konsisten pada data latih maupun data validasi. Tidak adanya jarak yang signifikan antara kedua kurva menunjukkan bahwa model tidak mengalami overfitting, dan proses pembelajaran berlangsung secara optimal.

Sementara itu, grafik loss pada gambar diatas memperlihatkan penurunan nilai kesalahan secara drastis sejak epoch pertama. Training loss yang awalnya berada pada nilai sekitar 0.0808 turun dengan cepat hingga mendekati 0.0020, sedangkan validation loss juga menunjukkan tren penurunan yang konsisten hingga mencapai nilai terendah sebesar 0.0019 pada epoch ke-11. Kemiripan pola antara kurva loss pada data latih dan validasi

memperkuat indikasi bahwa model tidak mengalami ketidakseimbangan pembelajaran. Nilai yang rendah pada kedua kurva loss mengonfirmasi bahwa model mampu meminimalkan kesalahan prediksi dan mempelajari representasi fitur secara efektif selama proses pelatihan.

Gabungan antara grafik akurasi dan loss ini menunjukkan bahwa model LSTM yang dikembangkan memiliki stabilitas yang sangat baik, dengan DFkapasitas generalisasi yang kuat dan tingkat kesalahan yang sangat rendah. Kinerja ini mengindikasikan bahwa arsitektur dan parameter pelatihan yang diterapkan sudah sesuai sehingga model mampu mengenali pola transaksi dan membedakan antara kelas penipuan dan non-penipuan secara optimal.

b) Evaluasi Kinerja Model Saat Pengujian

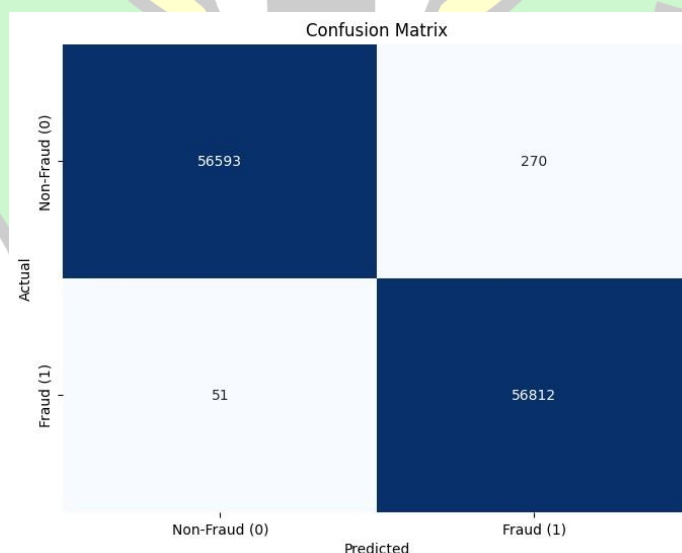
Evaluasi performa model dilakukan menggunakan sejumlah metrik klasifikasi, yaitu accuracy, precision, recall, dan F1-score. Keempat metrik ini dipilih karena mampu memberikan gambaran komprehensif mengenai kemampuan model dalam mendeteksi transaksi penipuan maupun transaksi normal. Berikut hasil dari evaluasi performa model menggunakan metrik tersebut:

Tabel 4. 4 Evaluasi Performa Model

NO	Metrik Evaluasi	Hasil
1	Accuracy	0.999639
2	Precision	0.999315
3	Recall	0.999965
4	F1-Score	0.999640

Berdasarkan hasil pengujian pada data uji, model mencapai tingkat akurasi yang sangat tinggi, yaitu sebesar 0.999639. Nilai ini menunjukkan bahwa hampir seluruh prediksi model sesuai dengan label aktual transaksi. Metrik precision yang diperoleh sebesar 0.999315, menandakan bahwa dari seluruh transaksi yang diprediksi sebagai penipuan oleh model, hampir semuanya benar-benar merupakan transaksi penipuan. Nilai recall sebesar 0.999965 menunjukkan bahwa model mampu mendeteksi hampir seluruh transaksi penipuan yang ada di dalam data uji. Kedua nilai ini merefleksikan kemampuan model dalam menghindari kesalahan deteksi, baik false positive maupun false negative. Kombinasi antara precision dan recall menghasilkan nilai F1-score sebesar 0.999640, yang mengindikasikan keseimbangan sangat baik antara kemampuan model dalam mengidentifikasi transaksi penipuan dan meminimalkan kesalahan prediksi.

Kinerja model juga divisualisasikan melalui confusion matrix sebagaimana ditunjukkan pada gambar dibawah ini:



Gambar 4. 9 Confusion Matrix

Matriks ini menggambarkan bahwa model berhasil mengklasifikasikan 56.824 transaksi non-penipuan secara benar (True Negative) dan 56.861

transaksi penipuan secara benar (True Positive). Hanya terdapat 39 prediksi salah untuk kelas non-penipuan (False Positive) dan 2 prediksi salah untuk kelas penipuan (False Negative). Jumlah kesalahan yang sangat kecil ini menunjukkan bahwa model memiliki kemampuan generalisasi yang sangat baik serta mampu membedakan kedua kelas dengan tingkat akurasi yang hampir sempurna.



BAB V

PENUTUP

5.1 Kesimpulan

Berdasarkan keseluruhan tahapan penelitian mengenai penerapan Long Short-Term Memory (LSTM) untuk mendeteksi penipuan pada transaksi kartu kredit yang telah dilaksanakan, dapat dirumuskan beberapa kesimpulan utama sebagai berikut:

1. Model LSTM mampu mendeteksi transaksi penipuan dengan tingkat akurasi yang sangat tinggi. Setelah melalui proses prapemrosesan, pembagian data, pelatihan model, dan evaluasi menggunakan metrik klasifikasi, model menghasilkan performa yang sangat baik dengan nilai akurasi sebesar 0.999639, precision 0.999315, recall 0.999965, dan F1score 0.999640. Hasil ini menunjukkan bahwa LSTM efektif dalam membedakan transaksi penipuan dan transaksi normal, serta mampu mempelajari pola numerik kompleks yang terdapat dalam fitur-fitur PCA dari dataset.
2. Proses pelatihan menunjukkan stabilitas dan generalisasi model yang sangat baik. Kurva accuracy dan loss pada data latih maupun data validasi memperlihatkan tren peningkatan dan penurunan yang konsisten, tanpa adanya indikasi overfitting. EarlyStopping dan ModelCheckpoint berhasil memastikan bahwa model berhenti pada titik performa terbaik (epoch ke-6) serta menyimpan bobot dengan validasi akurasi tertinggi.
3. Distribusi data yang seimbang pada dataset berkontribusi signifikan terhadap performa model. Tidak seperti dataset fraud pada umumnya yang memiliki class imbalance tinggi, dataset yang digunakan memiliki proporsi 50:50 antara transaksi fraud dan non-fraud. Hal ini membuat proses

pembelajaran lebih stabil dan tidak memerlukan teknik tambahan seperti SMOTE atau cost-sensitive learning.

5.2 Saran

1. Memperluas variasi dataset dari berbagai sumber finansial. Penelitian mendatang dianjurkan untuk menggunakan dataset dari lebih dari satu institusi keuangan atau platform transaksi digital. Hal ini bertujuan agar model mampu mengenali pola penipuan yang lebih beragam dan tidak bergantung pada pola statistik satu dataset saja. Model yang diuji pada data multinasional atau multi-platform cenderung menghasilkan generalisasi yang lebih kuat.
2. Menguji model pada kondisi serangan atau pola fraud baru (Emerging Fraud Patterns). Kasus penipuan selalu berkembang mengikuti teknologi. Oleh karena itu, penelitian selanjutnya dapat memasukkan skenario synthetic attack simulation seperti rapid transaction bursts, location anomalies, hingga device switching patterns untuk melihat bagaimana model merespons skenario fraud yang tidak terdapat pada data pelatihan.
3. Menguji model pada lingkungan operasional nyata (Real Deployment Scenario). Penelitian selanjutnya disarankan untuk melakukan uji coba model pada lingkungan operasional simulasi, misalnya alur transaksi realtime atau integrasi API pembayaran. Hal ini memungkinkan evaluasi terhadap waktu respons model, konsumsi memori, dan efektivitas deteksi dalam kondisi volume transaksi tinggi.

DAFTAR PUSTAKA

- Achmad Baihaqi, M. H. (2022). *Hak Cipta dalam Perspektif Hukum Islam*. Q Media.
- Ahsun, A., Joy, B., Noan, B., & Okunola, A. (2025). *The Future of Real-Time Fraud Detection: Trends and Innovations*.
- Al-Hashedi, K. G., & Magalingam, P. (2021). Financial fraud detection applying data mining techniques: A comprehensive review from 2009 to 2019. *Computer Science Review*, 40, 100402.
- Alghafiqi, B., & Munajat, E. (2022). *Impact Of Artificial Intelligence Technology On Accounting Profession Dampak Teknologi Artificial Intelligence Pada Profesi Akuntansi*.
- Alim, M. N. (2023). Pemodelan Time Series Data Saham LQ45 dengan Algoritma LSTM, RNN, dan Arima. *PRISMA, Prosiding Seminar Nasional Matematika*, 6, 694–701.
- Arfan, A., & Lussiana, E. T. P. (2020). *Perbandingan algoritma long short-term memory dengan SVR pada prediksi harga saham di Indonesia*.
- BARBARA, S. (2025). *Payment Card Fraud Losses Approach \$34 Billion*. Report, The Nilson. <https://www.globenewswire.com/news-release/2025/01/06/3004931/0/en/Payment-Card-Fraud-Losses-Approach-34-Billion.html>
- Benchaji, I., Douzi, S., El Ouahidi, B., & Jaafari, J. (2021). Enhanced credit card fraud detection based on attention mechanism and LSTM deep model. *Journal of Big Data*, 8, 1–21.
- Budiyansyah, D. P. (2025). Prediksi real or fake job posting menggunakan metode long short-term memory. *Innotech: Jurnal Ilmu Komputer, Sistem Informasi Dan Teknologi Informasi*, 2(1).
- Caniago, A. I., Kaswidjanti, W., & Juwairiah, J. (2021). Recurrent Neural Network With Gate Recurrent Unit For Stock Price Prediction. *Telematika*, 18(3), 345–360.

- Carudin, C., Marisa, M., Murnawan, M., Reba, F., Koibur, M. E., Thantawi, A. M., Halim, A., & Wattimena, F. Y. (2024). *Buku Ajar Data Mining*. PT. Sonpedia Publishing Indonesia.
- Castello Purba, A., & Handhayani, T. (2024). PERBANDINGAN ALGORITMA K-MEANS, AFFINITY CLUSTERING, DAN MINIBATCH K-MEANS UNTUK ANALISIS SEGMENTASI PASAR. *KOMPUTA : Jurnal Ilmiah Komputer Dan Informatika*, 13(1).
- CICIANA, C. (2024). *Metode resampling dan random forest untuk klasifikasi data*. Universitas sulawesi barat.
- CPI. (2023). *Fraud is a major—and growing—problem for financial institutions*. CPI Card Group. <https://cpicardgroup.com/2025/01/15/payment-fraud-isgrowing-its-time-for-an-innovative-proactive-approach/>
- Dusenov, H., & Wibowo, A. (2025). Deteksi penipuan pada sosial media twitter dengan metode bidirectional long short term memory (bi-lstm). *INFORMATIKA*, 16(1), 117–125.
- Fattah, H., Riadini, I., Hasibuan, S. W., Rahmanto, D. N. A., Layli, M., Holle, M. H., Arsyad, K., Aziz, A., Santoso, W. P., & Mutakin, A. (2022). *Fintech dalam Keuangan Islam: Teori dan Praktik*. Publica Indonesia Utama.
- Fons, E. M. (2021). *Machine learning applications on time series data for systematic investing*. The University of Manchester (United Kingdom).
- Hendrayana, I. G., Suprayitno, D., Judijanto, L., Kosadi, F., Kusumastuti, S. Y., & Sepriano, S. (2024). *E-Money: Panduan Lengkap Penggunaan dan Manfaat E-Money dalam Era Digital*. PT. Sonpedia Publishing Indonesia.
- Husna, A. (2024). Kecurangan dalam Perspektif Islam: Etika dan Konsekuensinya. *JEMSI (Jurnal Ekonomi, Manajemen, Dan Akuntansi)*, 10(2), 1491–1499.
- Ismail, A. (2020). *Klasifikasi short message service (sms) menggunakan algoritma long short term memory dan recurrent neural network*. Universitas YARSI.
- Jannah, A. M. (2020). *Penegakan Hukum Pidana Terhadap Tindak Pidana Penipuan Bisnis Online Di Polda Metro Jaya Menurut Hukum Positif dan*

Hukum Islam. Fakultas Syariah dan Hukum Universitas Islam Negeri Syarif Hidayatullah Jakarta.

Kavitha, M., & Kasthuri, M. (2024). Enhanced cost-sensitive ensemble learning for imbalanced class in medical data. *J. Electrical Systems*, 20(7s), 1043–1053.

Kurniansyah, J., Kurnia Gusti, S., Yanto, F., & Affandes, M. (2025). Implementasi Model Long Short Term Memory (LSTM) dalam Prediksi Harga Saham. *Bulletin of Information Technology (BIT)*, 6(2), 79–86. <https://doi.org/10.47065/bit.v5i2.1783>

Lestari, T. S., & Sirodj, D. A. N. (2021). Klasifikasi Penipuan Transaksi Kartu Kredit Menggunakan Metode Random Forest. *Jurnal Riset Statistika*, 160–167.

Maftucha, N., Salma, S., Rahmayuna, N., & Wakhidah, N. (2023). Perbandingan Algoritma Machine Learning Dalam Memprediksi Kelulusan Siswa. *Jurnal TEKNO KOMPAK*, 19(2).

Muslikah, A. N. (2021). *SMS Spam Detection Menggunakan Metode Long Short Term Memory*. Universitas Islam Negeri Maulana Malik Ibrahim.

Nasir, J., Kom, S., Kom, M., Pernando, Y., Wati, M., Yuniarthe, M. T. D. Y., Oktavia, S. S., Nurhayati, S., Kom, M., & Lestari, S. (2025). *Pengantar pembelajaran machine learning*.

Ngamal, Y., & Perajaka, M. A. (2022). Penerapan Model Manajemen Risiko Teknologi Digital Di Lembaga Perbankan Berkaca Pada Cetak Biru Transformasi Digital Perbankan Indonesia. *Jurnal Manajemen Risiko*, 2(2), 59–74.

Ningsih, P. T. S., Gusvarizon, M., & Hermawan, R. (2022). Analisis Sistem Pendeteksi Penipuan Transaksi Kartu Kredit dengan Algoritma Machine Learning. *Jurnal Teknologi Informatika Dan Komputer*, 8(2), 386–401.

Nugraha, M. N., & Arrofiq, M. (2024). Purwarupa Sistem Klasifikasi Legalitas Investasi Berbasis Algoritma Bidirectional Long Short Term Memory. *Journal of Internet and Software Engineering*, 5(1), 9–14.

- Putri, A. I., Husna, N. A., Cia, N. M., Arba, M. A., Aisyi, N. R., Pramesthi, C. H., & Irdayusman, A. S. (2024). Implementation of K-Nearest Neighbors, Naïve Bayes Classifier, Support Vector Machine and Decision Tree Algorithms for Obesity Risk Prediction. *Public Research Journal of Engineering, Data Technology and Computer Science*, 2(1), 26–33.
- Rahayu, N. (2025). Aggarwal, CC (2018). Artificial Neural Network and Deep Learning. Springer. Anderson, JA (1995). An Introduction to Neural Networks. The MIT Press. Bahdanau, D., Cho, K., & Bengio, Y.(2015). Neural Machine Translation. *Deep Learning: Teori, Algoritma, Dan Aplikasi*, 9, 56.
- Romero, L. (2023). *Value of credit card transactions in Indonesia from 2017 to 2022*. Statista. <https://www.statista.com/statistics/1254955/indonesia-creditcard-transaction-value/>
- Sarker, I. H. (2021). Deep cybersecurity: a comprehensive overview from neural network and deep learning perspective. *SN Computer Science*, 2(3), 154.
- Sathupadi, K., Achar, S., Bhaskaran, S. V., Faruqui, N., & Uddin, J. (2025). BankNet: Real-Time Big Data Analytics for Secure Internet Banking. *Big Data and Cognitive Computing*, 9(2), 24.
- Singla, J. (2022). RETRACTED: A novel framework for online transaction fraud detection system based on deep neural network. *Journal of Intelligent & Fuzzy Systems*, 43(1), 927–937.
- Sir, Y. A., & Soepranoto, A. H. H. (2022). Pendekatan Resampling Data Untuk Menangani Masalah Ketidakseimbangan Kelas. *J-ICON: Jurnal Komputer Dan Informatika*, 10(1), 31–38.
- Takiddin, A., Ismail, M., Atat, R., Davis, K. R., & Serpedin, E. (2023). Robust graph autoencoder-based detection of false data injection attacks against data poisoning in smart grids. *IEEE Transactions on Artificial Intelligence*, 5(3), 1287–1301.
- Wanda, P., Hiswati, M. E., Diqi, M., & Herlinda, R. (2021). Re-fake: Klasifikasi akun palsu di sosial media online menggunakan algoritma RNN. *Prosiding Seminar Nasional Sains Teknologi Dan Inovasi Indonesia P-ISSN*, 2086, 5805.

Waras, N. G. T., Saptadi, N., Wardani, A. K., Pardosi, V. B. A., Hasanah, Q., Kurniasari, A. A., Firmansyah, M. H., Arifianto, A. S., Maulani, G., & Iin, J. N. (2024). *Kuasai machine learning & computer vision dalam sekejap*.

Widhiyasana, Y., Semiawan, T., Mudzakir, I. G. A., & Noor, M. R. (2021). Penerapan Convolutional Long Short-Term Memory untuk Klasifikasi Teks Berita Bahasa Indonesia. *Jurnal Nasional Teknik Elektro Dan Teknologi Informasi* | Vol, 10(4), 354–361.

Wire, B. (2023). *Global Non-Cash Transaction Volumes Set to Reach 1.3 Trillion in 2023*. Tradeshownews.Com.
<https://www.businesswire.com/news/home/20230914460622/en/Global-Non-Cash-Transaction-Volumes-Set-to-Reach-1.3-Trillion-in-2023>

