

Implementasi Algoritma Neural Network untuk Deteksi Penipuan Transaksi Kartu Kredit

Kurnia Prayogi¹, Dicky Octaviano^{2*}, Zulfati Dinul Fatiha³

¹Fakultas Teknologi Informasi, Universitas Nusa Mandiri, Jakarta, Indonesia

^{2,3}Facultas Ekonomi dan Bisnis, Universitas Bina Sarana Informatika, Jakarta, Indonesia

Email: ¹14220039@nusamandiri.ac.id, ²dicky.doc.@bsi.ac.id, ³zulfati.zdf@bsi.ac.id

Penulis Korespondensi*

(received: 20-02-26, revised: 18-04-26, accepted: 10-06-26)

Abstrak

Penelitian ini mengevaluasi performa lima algoritma supervised learning untuk deteksi penipuan kartu kredit menggunakan dataset 690 data dari Kaggle dengan teknik Random Oversampling (ROS). Model seperti k-Nearest Neighbor (k-NN), Support Vector Machine (SVM), Logistic Regression, Neural Network, dan Ensemble menunjukkan tingkat akurasi rata-rata antara 80% hingga 90% dalam mendeteksi penipuan. Kontribusi penelitian ini adalah menyediakan perbandingan sistematis beberapa algoritma klasifikasi pada dataset yang sama dengan teknik penyeimbangan data. Hasil uji coba dengan teknik random oversampling menunjukkan bahwa Neural Network (aktivasi SELU dan RELU), mencapai kinerja terbaik dengan accuracy 90%, precision 86%, recall 94%, dan nilai f1-score 90%. Pendekatan Neural Network dengan random oversampling terbukti efektif dalam meningkatkan ketepatan prediksi terhadap penipuan dalam transaksi finansial dibandingkan dengan pendekatan tanpa penggunaan sampling. Keterbatasan penelitian ini adalah ukuran dataset yang kecil (690 data) yang dapat mempengaruhi kemampuan generalisasi model.

Kata Kunci: Prediksi Penipuan, Neural Network, Klasifikasi, Random Oversampling

Abstract

This study evaluates the performance of five supervised learning algorithms for credit card fraud detection using a dataset of 690 records from Kaggle with Random Oversampling (ROS) to address class imbalance. Models such as k-Nearest Neighbor (k-NN), Support Vector Machine (SVM), Logistic Regression, Neural Network, and Ensemble exhibit average accuracy rates between 80% and 90% in detecting fraud. The aim of this study is to enhance fraud detection classification quality in future transactions by understanding fraud patterns and contexts, contributing a comparative analysis of multiple classifiers under a balanced sampling strategy. Testing results with random oversampling indicate that Neural Network (SELU + RELU activation) achieves the best performance with 90% accuracy, 86% precision, 94% recall, and 90% f1-score. The Neural Network approach with random oversampling proves effective in enhancing fraud prediction accuracy in financial transactions compared to approaches without sampling. However, the small dataset size (690 records) represents a limitation that may affect generalizability.

Keyword: Fraud Predict, Neural Network, Classification, Random Oversampling

1. PENDAHULUAN

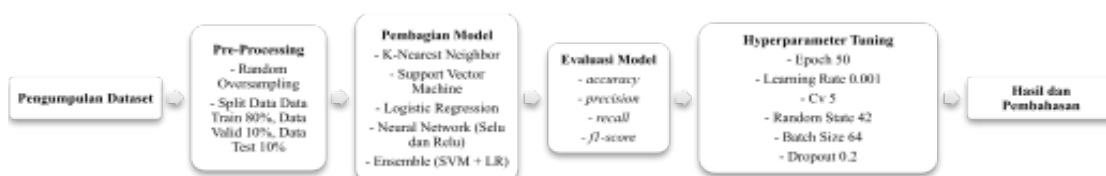
Dalam era digital yang sudah berkembang pesat dimana penggunaan kartu kredit kini semakin meluas, tetapi di sisi lain juga dapat meningkatkan risiko penipuan. Data pada Bank Indonesia dan Asosiasi Kartu Kredit Indonesia menunjukkan pertumbuhan yang signifikan dalam jumlah kartu kredit dan volume transaksi di Indonesia sejak 2023, menyoroti pentingnya perlindungan terhadap keamanan transaksi menggunakan kartu kredit tidaklah luput dari tindak kejahatan seperti transaksi ilegal yang dilakukan oleh pihak tak dikenal dengan memanfaatkan kebocoran data pribadi pemilik kartu kredit [1]. Penelitian ini bertujuan untuk mengevaluasi dan menerapkan algoritma deep learning dalam mendeteksi penipuan kartu kredit dengan akurasi tinggi serta membandingkannya dengan metode klasifikasi lainnya. Penelitian ini menggunakan dataset dari *kaggle*, dimana penelitian ini akan menguji berbagai arsitektur model deep learning, mengoptimalkan parameter, dan melakukan validasi untuk memastikan efektivitas model dalam sistem deteksi penipuan kartu kredit.

Meskipun penelitian-penelitian sebelumnya telah menunjukkan keberhasilan berbagai algoritma klasifikasi dalam deteksi penipuan, terdapat beberapa research gap yang perlu diisi. Pertama, sebagian besar penelitian terdahulu menggunakan dataset berukuran besar (puluhan ribu hingga ratusan ribu data), sehingga performa algoritma pada dataset kecil dengan karakteristik kelas tidak seimbang belum banyak dieksplorasi secara komparatif. Kedua, belum ada studi yang secara sistematis membandingkan lima algoritma (k-NN, SVM, Logistic Regression, Neural Network, dan Ensemble) secara bersamaan pada dataset yang sama dengan teknik Random Oversampling. Ketiga, pengaruh ukuran dataset yang sangat terbatas (690 data) terhadap kemampuan generalisasi model deep learning belum didiskusikan secara memadai dalam literatur. Penelitian ini hadir untuk menjawab gap tersebut dengan melakukan evaluasi komparatif yang sistematis dan mendiskusikan implikasi ukuran dataset terhadap validitas model yang dihasilkan.

Penelitian ini juga mengambil dari beberapa penelitian terdahulu yang menggunakan beberapa algoritma klasifikasi seperti pada penelitian ini berjudul "*Analisa Fraud Pada Transaksi Kartu Kredit Menggunakan Algoritma Random Forest*" oleh Ravina Armiani, dan Eka Puji Agustini memaparkan bahwa dari hasil pengujian tersebut tanpa SMOTE menunjukkan kinerja sebesar 97,98%, sedangkan dengan penggunaan SMOTE kinerja klasifikasi turun menjadi 90,68%. Penggunaan algoritma *random forest* dan SMOTE menghasilkan akurasi sebesar 90,68%, presisi sebesar 94,74%, recall sebesar 86,14%, dan F1-Score sebesar 90,23%[2]. Penelitian selanjutnya yang berjudul "*Penerapan Deep Learning dalam Deteksi Penipuan Transaksi Keuangan Secara Elektronik*" oleh Faried Zamachsari, dan Niken Puspitasari, dimana pada penelitian tersebut menggunakan pengujian tanpa SMOTE dan dengan SMOTE. Pada penggunaan SMOTE mendapatkan akurasi dengan *Deep Neural Network (Selu dan Relu)* sebesar 99.602%, *Random Forest* 99.537%, *Extra Tress* 99.530%, *Adaboost* 99.443%, *Stochastic Gradient Boosting* 99.455%, dan *Bagged Decision Trees* 99.522%. Selanjutnya dengan penggunaan SMOTE mendapatkan *Deep Neural Network (Selu dan Relu)* sebesar 99.217%, *Random Forest* 99.802%, *Extra Tress* 99.835%, *Adaboost* 97.116%, *Stochastic Gradient Boosting* 97.196%, dan *Bagged Decision Trees* 99.781%[3]. Kemudian penelitian berikutnya dengan judul "*Klasifikasi Penipuan Transaksi Kartu Kredit Menggunakan Metode Random Forest*" oleh Tiara Suci Lestari, dan Dwi Agustin Nuriani Sirodj yang melakukan pengujian dengan mendeskripsikan data dan melakukan pembagian data dengan proporsi 75% untuk data pelatihan (37.656 data) dan 25% untuk data pengujian (12.552 data), sehingga total data berjumlah 50.208 transaksi, dapat disimpulkan bahwa metode *random forest* efektif diterapkan pada klasifikasi penipuan transaksi kartu kredit. Metode ini menghasilkan akurasi 97,275%, sensitivitas 98,795%, presisi 97,976%, *F-Measure* 98,384%, dan nilai AUC 94,065%. Hasil ini menunjukkan kinerja yang sangat baik, karena semua matriks klasifikasi berada dalam rentang 90-100% [4].

2. METODOLOGI PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif yang didasarkan pada filsafat *positivisme*, dengan fokus pada populasi atau sampel tertentu. Data akan dikumpulkan melalui instrumen penelitian dan dianalisis secara kuantitatif [5]. Penelitian ini akan menggunakan metode kuantitatif yang memanfaatkan fungsi dari klasifikasi *k-Nearest Neighbor* (k-NN) untuk mendeteksi penipuan kartu kredit yang bertujuan untuk mengevaluasi nilai *accuracy*, *recall*, *f1-score*, dan *precision*, sehingga dapat disimpulkan efektivitas k-NN dalam mendeteksi transaksi penipuan dan diharapkan dapat memberikan hasil deteksi yang optimal dalam mengidentifikasi transaksi penipuan dan mengurangi kerugian finansial yang diakibatkan oleh penipuan kartu kredit. Penelitian ini melibatkan beberapa tahapan, yang dijelaskan pada Gambar 1 berikut.



Gambar 1. Tahapan Penelitian

2.1. Pengumpulan Dataset

Penelitian ini menggunakan dataset dari *kaggle* yang disediakan oleh Reza Semyari dengan judul "*Credit-Card-Fraud-Detection*". Dataset ini terdiri dari 690 baris data yang mencakup informasi mengenai customer yang menggunakan kartu kredit serta berbagai fitur terkait. Tujuan dari dataset ini adalah untuk mengidentifikasi penipuan berdasarkan variabel input numerik yang telah ditransformasi menggunakan *Principal Component Analysis* (PCA) demi menjaga privasi data, fitur asli dan informasi latar belakang lainnya.

Fitur pada dataset ini diberi label seperti *customer_id*, A_1 hingga A_14 serta label *class* yang menunjukkan apakah transaksi tersebut merupakan penipuan (1) atau bukan (0).

Tabel 1. Dataset Transaksi Kartu Kredit

Customer_id	A_1	A_2	A_3	A_4	A_5	A_6	...	class
15776156	1	22.08	11.46	2	4	4	...	0
15739548	0	22.67	7.0	2	8	4	...	0
15662854	0	29.58	1.75	1	4	4	...	0
15687688	0	21.67	11.5	1	5	3	...	1
15715750	1	20.17	8.17	2	6	4	...	1
15571121	0	15.83	0.585	2	8	8	...	1
...	...	...	...	...	...	...	...	...
15592412	1	41.0	0.04	2	10	4	...	1

2.2. Preprocessing

Preprocessing data adalah tahapan yang penting dalam mempersiapkan data untuk di analisa. Proses ini bertujuan mengubah data mentah menjadi format yang siap digunakan untuk pengujian, sehingga memastikan hasil analisa yang akurat dan relevan[6]. Berdasarkan data yang telah dikumpulkan, *pre-processing* melibatkan serangkaian tahapan, antara lain seperti split data dan fitur skaling dimana untuk split data bertujuan memisahkan data menjadi data latih dan data uji untuk mempermudah dalam pengujian dan memvalidasi kinerja model deteksi penipuan kartu kredit, sedangkan fitur skaling menggunakan model *random oversampling* yang bertujuan untuk memastikan setiap fitur memiliki skala yang serupa. *Random Oversampling* (ROS) adalah metode yang menambahkan data dari kelas minoritas ke dalam set data pelatihan secara acak. Proses ini diulang hingga jumlah data di kelas minoritas seimbang dengan jumlah data di kelas mayoritas [6].

2.3. Deep Learning

Deep Learning adalah sebuah teknik berbasis jaringan saraf tiruan yang telah menjadi sangat populer dalam beberapa tahun terakhir sebagai metode implementasi *Machine Learning*. Menurut berbagai artikel, *Deep Learning* tidak terbatas pada bidang tertentu, tetapi telah diakui sebagai bentuk pembelajaran universal yang mampu menyelesaikan beragam masalah di berbagai sektor [7].

2.4. Random Over Sampler

Random Oversampling (ROS) bekerja dengan menambahkan data dari kelas minoritas ke set data pelatihan secara acak. Proses ini dilakukan berulang kali hingga jumlah data dari kelas minoritas menjadi setara dengan jumlah data dari kelas mayoritas. Langkah pertama yang dilakukan adalah menentukan selisih antara jumlah data kelas mayoritas dan minoritas. Setelah itu, data dari kelas minoritas dipilih secara acak dan ditambahkan ke set pelatihan, proses ini diulang hingga selisih tersebut terpenuhi [6].

2.5. K-Nearest Neighbor (k-NN)

K-Nearest Neighbor (K-NN) adalah algoritma *supervised learning* yang dilatih sebelumnya agar dapat melakukan prediksi dan klasifikasi pada data yang akan digunakan. Algoritma ini populer di kalangan peneliti karena memiliki beberapa keunggulan, seperti akurasi yang tinggi, insensitivitas terhadap outlier, dan tidak memerlukan asumsi tertentu terhadap data. Namun, K-NN juga memiliki beberapa kelemahan, termasuk perlunya menentukan nilai “k” yang optimal, biaya komputasi yang tinggi, dan kebutuhan memori yang besar[8].

2.6. Ensemble Learning

Ensemble Learning merupakan metode yang memanfaatkan berbagai algoritma pembelajaran untuk mencapai prediksi yang lebih baik daripada yang bisa diperoleh dari satu algoritma. Berbeda dengan *Ensemble* dalam mekanika statistik yang biasanya tidak terbatas[3], dalam *Machine Learning*, ketika sebuah model belum mampu memberikan hasil yang akurat, diperlukan pendekatan tambahan untuk meningkatkan akurasi. Salah satu cara yang efektif adalah dengan menggabungkan *Ensemble* dengan pendekatan *hybrid*, yang bertujuan untuk meningkatkan akurasi pada dataset dengan kelas yang tidak seimbang [9].

2.7. Logistic Regression

Logistic Regression adalah jenis regresi yang menghubungkan satu atau lebih variabel independen dengan variabel dependen yang berbentuk kategori, seperti 0 dan 1, benar atau salah, besar atau kecil. Bentuk variabel independen yang kategorikal ini membedakan regresi logistik dari regresi berganda atau regresi linier lainnya [10].

2.8. Support Vector Machine

Support Vector Machine (SVM) adalah model pembelajaran terawasi yang menganalisis data untuk keperluan klasifikasi dan regresi. Selain klasifikasi linier, SVM juga mampu menangani klasifikasi non-linier secara efisien melalui penggunaan trik kernel, yang secara implisit memetakan data input ke dalam ruang fitur berdimensi tinggi. Intinya, SVM menetapkan margin antara kelas – kelas, di mana margin ini ditarik sedemikian rupa untuk memaksimalkan jarak antara margin dan kelas, sehingga meminimalkan kesalahan klasifikasi [1].

2.9. Neural Network

Neural Network adalah model matematis yang meniru fungsi jaringan saraf biologis. Jaringan saraf tiruan ini terdiri dari neuron buatan yang saling terhubung dan bekerja bersama untuk memproses informasi. Secara umum, model ini merupakan sistem adaptif yang menyesuaikan strukturnya berdasarkan informasi eksternal atau internal yang diproses selama fase pembelajaran, dimana mampu memodelkan hubungan yang kompleks antara input dan output, serta mengidentifikasi pola dalam data supaya efektif dalam berbagai bidang seperti estimasi, pengenalan tulisan tangan, penilaian harga, dan banyak lagi [12].

Dalam penelitian ini, arsitektur Neural Network yang digunakan terdiri dari dua hidden layer dengan fungsi aktivasi SELU pada layer pertama dan RELU pada layer kedua. Model dikompilasi menggunakan optimizer Adam dengan learning rate default (0.001) dan fungsi loss binary crossentropy. Detail spesifikasi arsitektur disajikan pada Tabel 2 berikut.

Tabel 2. Spesifikasi Arsitektur Model Neural Network

Komponen	Spesifikasi
Input Layer	14 fitur (A_1 hingga A_14)
Hidden Layer 1	64 neuron, aktivasi SELU
Hidden Layer 2	32 neuron, aktivasi RELU
Output Layer	1 neuron, aktivasi Sigmoid
Optimizer	Adam (learning rate = 0.001)
Loss Function	Binary Crossentropy
Epochs	100

2.10. Evaluasi Model

Dalam penelitian ini, dilakukan juga evaluasi terhadap model dengan tujuan untuk melihat kesesuaian klasifikasi menggunakan *confusion matrix*. Dengan penyesuaian ini untuk menilai hasil kinerja model, *confusion matrix* juga bisa menunjukkan jumlah kesalahan prediksi baik pada kelas minoritas maupun kelas mayoritas [11].

1) Accuracy

Accuracy adalah ukuran yang menunjukkan sejauh mana model klasifikasi dapat memprediksi kelas yang benar dari data yang diberikan [12].

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (1)$$

2) Recall

Recall adalah metrik yang menunjukkan berapa banyak data yang benar-benar diklasifikasikan dalam suatu kategori dibandingkan dengan total jumlah data yang sebenarnya termasuk dalam kategori tersebut [12].

$$Recall = \frac{TP}{TP+FN} \quad (2)$$

3) Precision

Precision adalah rasio antara jumlah label yang diidentifikasi dengan benar dan total jumlah label yang diidentifikasi [13].

$$Precision = \frac{TP}{TP+FP} \times 100\% \quad (3)$$

4) *F1-Score*

F1-Score adalah kombinasi dari nilai precision dan recall yang hasilnya digunakan sebagai nilai pengukuran [14].

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (4)$$

3. HASIL DAN PEMBAHASA

3.1. Pengumpulan Dataset

Dataset di Table 3 ini merepresentasikan variabel input numerik yang telah ditransformasi menggunakan *Principal Component Analysis* (PCA) untuk mengidentifikasi penipuan.

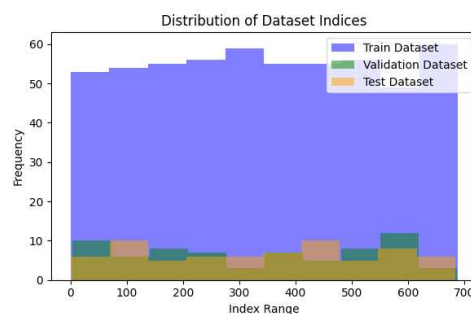
Tabel 3. Contoh Data dari Dataset Transaksi Kartu Kredit

No	Customer ID	A 1	A 2	A 3	A 4	A 5	A 6	A 7	A 8	A 9	A 10	A 11	A 12	A 13	A 14	Class
1	15776156	1	22.08	11.46	2	4	4	1.585	0	0	0	1	2	100	1213	0
2	15739548	0	22.67	7.00	2	8	4	0.165	0	0	0	0	2	160	1	0
3	15662854	0	29.58	1.75	1	4	4	0.000	0	0	0	1	2	280	1	0
4	15687688	0	21.67	11.50	1	5	3	0.000	1	1	11	1	2	0	1	1
5	15715750	1	20.17	8.17	2	6	4	1.960	1	1	14	0	2	60	159	1

Sumber: Dataset Credit Card Fraud Detection (Reza Semyari, Kaggle, 2023)

3.2. Preprocessing

Pada tahap ini, data yang telah dikumpulkan akan dibagi menjadi tiga bagian utama yaitu data train, data validasi, dan data uji. Sebelum pembagian ini dilakukan, peneliti menggunakan proses sampling menggunakan *random oversampling*. Tujuan dari *random oversampling* adalah untuk memastikan distribusi yang seimbang antara kelas yang berbeda dalam data yang digunakan untuk melatih model. Proses pembagian dataset ini melibatkan pengambilan semua kolom sebagai representasi fitur dari dataset kecuali satu kolom terakhir, sementara satu kolom terakhir dari dataset dijadikan sebagai target atau label. Hasil dari pembagian dataset ini menghasilkan 552 data train, 69 data test, dan 69 data validasi.



Gambar 3. Pembagian Dataset

3.3. Pelatihan Model

Pada pelatihan model, peneliti menggunakan beberapa model algoritma untuk klasifikasi deteksi penipuan, model tersebut antara lain *k-Nearest Neighbor* (k-NN), *Support Vector Machine* (SVM), *Logistic Regression*, *Neural Network*, dan *Ensemble Learning*. Pada penggunaan model *Ensemble Learning* melakukan gabungan beberapa model pembelajaran mesin untuk meningkatkan kinerja klasifikasi, disini peneliti menggunakan gabungan antara SVM dan *Logistic Regression*. Teknik pada *Ensemble Learning* yang digunakan adalah *hard voting* dimana setiap model dalam ensemble yang memberikan suara untuk kelas tertentu, dan kelas yang menerima suara terbanyak akan dipilih sebagai prediksi akhir.

Spesifikasi parameter setiap model yang digunakan dalam penelitian ini disajikan pada Tabel 4 berikut. Penyajian dalam bentuk tabel ini menggantikan tampilan potongan kode program (Gambar 4–8) agar lebih sesuai dengan standar penulisan artikel ilmiah.

Tabel 4. Spesifikasi Parameter Model Klasifikasi

Model	Parameter Utama	Justifikasi Pemilihan
<i>k-Nearest Neighbor</i>	$k=5$, $\text{metric}=\text{euclidean}$	Sederhana, efektif pada dataset kecil, tidak memerlukan asumsi distribusi data
<i>Support Vector Machine</i>	$\text{kernel}=\text{RBF}$, $C=1.0$, $\text{gamma}=\text{'scale'}$	Mampu menangani data berdimensi tinggi dan kelas tidak seimbang
<i>Logistic Regression</i>	$\text{solver}=\text{lbfgs}$, $\text{max_iter}=1000$, $C=1.0$	Baseline klasifikasi biner yang interpretatif dan konvergensi stabil
<i>Neural Network</i>	2 hidden layer (64-SELU, 32-RELU), $\text{optimizer}=\text{Adam}$, $\text{epochs}=100$	Mampu memodelkan hubungan non-linear kompleks antara fitur dan label
<i>Ensemble Learning</i>	Hard voting: SVM + Logistic Regression	Menggabungkan kekuatan dua model untuk prediksi yang lebih robust

3.4. Evaluasi Model

Pada langkah selanjutnya, dilakukan penilaian terhadap model yang telah dibuat sebelumnya. Evaluasi model ini memanfaatkan *confusion matrix*, yang berfungsi untuk mengukur kinerja suatu model atau algoritma. Dengan menggunakan *confusion matrix*, kita dapat mengidentifikasi tingkat kesalahan dari berbagai algoritma dengan melihat matrik seperti *accuracy*, *precision*, *recall*, dan *f1-Score*.

Tabel 5. Confusion Matrix Seluruh Model Klasifikasi

Model Klasifikasi	TN (Tidak Penipuan yang diprediksi benar)	FP (Tidak Penipuan yang diprediksi Penipuan)	FN (Penipuan yang diprediksi Tidak Penipuan)	TP (Penipuan yang diprediksi benar)
k-Nearest Neighbor (k-NN)	24	9	5	31
Support Vector Machine (SVM)	27	6	2	34
Logistic Regression	29	4	4	32
Neural Network (SELU + RELU)	31	2	5	31
Ensemble Learning (SVM + Logistic Regression)	29	4	4	32

Berdasarkan Tabel 5 di atas hasil pengujian beberapa model, terlihat bahwa penggunaan *Neural Network* dengan *Selu* dan *Relu* menunjukkan akurasi yang memadai. Penerapan *Random Over Sampling* (ROS) juga memberikan dampak positif pada dataset. Pada pengujian dengan *Neural Network* terdapat 31 prediksi yang benar untuk transaksi yang tidak melakukan penipuan, namun terdapat 2 kesalahan prediksi. Sedangkan untuk kelas yang melakukan penipuan, model ini berhasil memprediksi 5 transaksi dengan benar dan 31 transaksi dengan kesalahan prediksi. Pada Tabel 6 akan menampilkan nilai *accuracy*, *precision*, *recall*, dan *f1-score* dari metode *confusion matrix* tersebut.

Tabel 6. Evaluasi Klasifikasi Tidak Penipuan

Model Klasifikasi	<i>Accuracy</i>	<i>Recall</i>	<i>F1-Score</i>	<i>Precision</i>
k-Nearest Neighbor	80%	72%	77%	83%
Support Vector Machine	88%	82%	87%	93%
Logistic Regression	88%	88%	88%	88%
Neural Network (Selu + Relu)	90%	94%	90%	86%
Ensemble Learning (SVM + Logistic Regression)	88%	88%	88%	88%

3.5. Analisis Hasil dan Klasifikasi Model

Setelah melakukan prediksi dan evaluasi pada dataset menggunakan metode *k-Nearest Neighbor* (k-NN), *Support Vector Machine* (SVM), *Logistic Regression*, *Neural Network*, dan *Ensemble* yang menggabungkan SVM dan *Logistic Regression*, langkah selanjutnya adalah melakukan analisis. Dari pengujian yang sudah dilakukan, dapat diamati bahwa untuk kelas yang tidak melakukan penipuan dengan model unggul dengan prediksi benar terbanyak dari penggunaan *Neural Network* sebanyak 31 prediksi, diikuti oleh *Ensemble* (SVM + *Logistic Regression*) dan *Logistic Regression* 29 prediksi, SVM 27 prediksi, dan k-NN 24 prediksi, selain itu dengan jumlah kesalahan prediksi terendah dimiliki oleh *Neural Network* sebanyak 2 prediksi, diikuti oleh *Ensemble* (SVM + *Logistic Regression*) dan *Logistic Regression* 4 prediksi, SVM 6 prediksi, dan k-NN 9 prediksi. Sementara itu, untuk kelas yang melakukan deteksi penipuan menggunakan model *Neural Network* memiliki performa relatif lebih kecil dari model lainnya dalam kesalahan prediksi, sedangkan *k-Nearest Neighbor* gagal memprediksi secara benar dalam kelas ini. Analisis ini menunjukkan bahwa model *Neural Network* cukup menonjol dengan akurasi tinggi dan kesalahan prediksi yang rendah, meskipun dalam praktiknya dari faktor lain seperti waktu komputasi, kebutuhan sumber daya, dan interpretasi model juga perlu dipertimbangkan sesuai dengan kebutuhan aplikasi atau tujuan bisnis tertentu. Hasil pengujian menggunakan teknik *random oversampling* menunjukkan bahwa beberapa algoritma, terutama *Neural Network* dengan *Selu* dan *Relu*, mencapai kinerja terbaik. Dengan *accuracy* mencapai 90%, *precision* 86%, *recall* 94%, dan nilai *f1-score* 90%, *Neural Network* dengan *random oversampling* telah terbukti efektif dibandingkan dengan pengujian tanpa penggunaan sampling dalam meningkatkan ketepatan prediksi terhadap kasus penipuan dalam transaksi pada kartu kredit.

4. KESIMPULAN

Berdasarkan hasil dari penelitian yang telah dilakukan, dapat disimpulkan bahwa penerapan deep learning dalam prediksi penipuan menggunakan berbagai model memberikan hasil yang beragam dalam memprediksi dan mengklasifikasikan deteksi penipuan. Pada model yang digunakan seperti k-Nearest Neighbor (k-NN), Support Vector Machine (SVM), Logistic Regression, Neural Network, dan Ensemble menunjukkan tingkat akurasi rata-rata antara 80% hingga 90% dalam mendeteksi penipuan, yang mengungguli kinerja baseline tanpa sampling. Model Neural Network dengan aktivasi SELU dan RELU mencapai kinerja terbaik dengan akurasi 90%, precision 86%, recall 94%, dan f1-score 90%. Analisis dalam penelitian ini berkontribusi dalam memberikan perbandingan sistematis berbagai algoritma klasifikasi dengan teknik penyeimbangan data pada kasus deteksi penipuan kartu kredit. Tujuannya adalah meningkatkan keamanan dan kenyamanan dari sektor finansial yang penuh dengan aktifitas transaksi, dengan mengembangkan sistem pengklasifikasian yang lebih efisien dan responsif terhadap kebutuhan pengguna.

Penelitian ini memiliki beberapa keterbatasan yang perlu diperhatikan. Pertama, dataset yang digunakan hanya terdiri dari 690 baris data, yang merupakan ukuran yang relatif kecil untuk pelatihan model deep learning dan dapat mempengaruhi kemampuan generalisasi model pada data nyata berskala besar. Kedua, fitur-fitur dalam dataset telah melalui transformasi PCA sehingga interpretabilitas fitur asli tidak dapat dilakukan. Ketiga, penelitian ini belum membandingkan performa model pada skenario tanpa teknik sampling sebagai baseline yang komprehensif. Untuk penelitian selanjutnya, disarankan menggunakan dataset yang lebih besar, mengeksplorasi arsitektur Neural Network yang lebih dalam, serta mempertimbangkan teknik oversampling lain seperti SMOTE untuk meningkatkan performa deteksi penipuan.

DAFTAR PUSTAKA

- [1] P. T. S. Ningsih, M. Gusvarizon, and R. Hermawan, "Analisis Sistem Pendeteksi Penipuan Transaksi Kartu Kredit dengan Algoritma Machine Learning," *J. Teknol. Inform. dan Komput.*, vol. 8, no. 2, pp. 386–401, 2022, doi: 10.37012/jtik.v8i2.1306.
- [2] R. Armiani and E. P. Agustini, "Analisa Fraud Pada Transaksi Kartu Kredit Menggunakan Algoritma Random Forest," *J. Teknol. Inf. dan Terap.*, vol. 9, no. 2, pp. 118–126, 2022, doi: 10.25047/jtit.v9i2.297.
- [3] F. Zamachsari and N. Puspitasari, "Penerapan Deep Learning dalam Deteksi Penipuan Transaksi Keuangan Secara Elektronik," *J. RESTI (Rekayasa Sist. dan Teknol. Informasi)*, vol. 5, no. 2, pp. 203–212, 2021, doi: 10.29207/resti.v5i2.2952.
- [4] T. S. Lestari and D. A. N. Sirodj, "Klasifikasi Penipuan Transaksi Kartu Kredit Menggunakan Metode Random Forest," *J. Ris. Stat.*, vol. 1, no. 2, pp. 160–167, 2022, doi: 10.29313/jrs.v1i2.525.
- [5] A. F. Riany and G. Testiana, "Penerapan Data Mining untuk Klasifikasi Penyakit Stroke Menggunakan Algoritma Naïve Bayes," *J. SAINTEKOM*, vol. 13, no. 1, pp. 42–54, 2023, doi:

- 10.33020/saintekom.v13i1.352.
- [6] S. Diantika, "Penerapan Teknik Random Oversampling Untuk Mengatasi Imbalance Class Dalam Klasifikasi Website Phishing Menggunakan Algoritma Lightgbm," *JATI (Jurnal Mhs. Tek. Inform.,* vol. 7, no. 1, pp. 19–25, 2023, doi: 10.36040/jati.v7i1.6006.
- [7] Muhammad Haris Diponegoro, Sri Suning Kusumawardani, and Indriana Hidayah, "Tinjauan Pustaka Sistematis: Implementasi Metode Deep Learning pada Prediksi Kinerja Murid," *J. Nas. Tek. Elektro dan Teknol. Inf.,* vol. 10, no. 2, pp. 131–138, 2021, doi: 10.22146/jnteti.v10i2.1417.
- [8] E. R. Alfyyah, R. Andreswari, and E. Sutoyo, "Analisis dan deteksi fraud pada data panggilan menggunakan algoritma k-nearest neighbor (studi kasus: pt xyz)," *e-Proceeding Eng.,* vol. 7, no. 2, pp. 6640–6646, 2020.
- [9] A. Nugroho, M. A. Soeleman, R. A. Pramunendar, A. Affandy, and A. Nurhindarto, "Peningkatan Performa Ensemble Learning pada Segmentasi Semantik Gambar dengan Teknik Oversampling untuk Class Imbalance," *J. Teknol. Inf. dan Ilmu Komput.,* vol. 10, no. 4, pp. 899–908, 2023, doi: 10.25126/jtiik.20241046831.
- [10] M. R. Romadhon and F. Kurniawan, "A Comparison of Naive Bayes Methods, Logistic Regression and KNN for Predicting Healing of Covid-19 Patients in Indonesia," *3rd 2021 East Indones. Conf. Comput. Inf. Technol. EIconCIT 2021,* pp. 41–44, 2021, doi: 10.1109/EIconCIT50028.2021.9431845.
- [11] W. I. Sabilla and C. Bella Vista, "Implementasi SMOTE dan Under Sampling pada Imbalanced Dataset untuk Prediksi Kebangkrutan Perusahaan," *J. Komput. Terap.,* vol. 7, no. 2, pp. 329–339, 2021, doi: 10.35143/jkt.v7i2.5027.
- [12] A. Kusuma and H. Nurramdhani Irmanda, "Analisis Sentimen Pada Ulasan Aplikasi Indodax di Google Play Store Menggunakan Metode Support Vector Machine," 2022.
- [13] S. A. P. Perdana, T. Bharata Aji, and R. Ferdiana, "Aspect Category Classification dengan Pendekatan Machine Learning Menggunakan Dataset Bahasa Indonesia (Aspect Category Classification with Machine Learning Approach Using Indonesian Language Dataset)," *J. Nas. Tek. Elektro dan Teknol. Inf.,* vol. 10, no. 3, pp. 229–235, 2021.
- [14] S. Clara, D. Laksmi Prianto, R. Al Habsi, E. Friscila Lumbantobing, and N. Chamidah, "Implementasi Seleksi Fitur Pada Algoritma Klasifikasi Machine Learning Untuk Prediksi Penghasilan Pada Adult Income Dataset," *Semin. Nas. Mhs. Ilmu Komput. dan Apl. Jakarta-Indonesia,* vol. 2, no. 1, pp. 741–747, 2021.
- [15] Credit-Card-Fraud-Detection: Dataset. Retrieved Juni 26, 2023, from <https://www.kaggle.com/datasets/rezasemyari/credit-card-fraud-detection>. Kaggle (2023, Juni 26).
- [16] Y. Sahin and E. Duman, "Detecting Credit Card Fraud by ANN and Logistic Regression," in *Proc. Int. Symp. Innovations Intell. Syst. Appl.,* 2011, pp. 315–319, doi: 10.1109/INISTA.2011.5946108.
- [17] A. Dal Pozzolo, O. Caelen, R. A. Johnson, and G. Bontempi, "Calibrating Probability with Undersampling for Unbalanced Classification," in *Proc. IEEE Symp. Comput. Intell. Data Mining,* 2015, pp. 1–8, doi: 10.1109/CIDM.2015.7415158.
- [18] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic Minority Over-sampling Technique," *J. Artif. Intell. Res.,* vol. 16, pp. 321–357, 2002, doi: 10.1613/jair.953.
- [19] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. Cambridge, MA: MIT Press, 2016. [Online]. Available: <https://www.deeplearningbook.org>
- [20] S. Bhattacharyya, S. Jha, K. Tharakunnel, and J. C. Westland, "Data Mining for Credit Card Fraud: A Comparative Study," *Decis. Support Syst.,* vol. 50, no. 3, pp. 602–613, 2011, doi: 10.1016/j.dss.2010.08.008.
- [21] G. Ke et al., "LightGBM: A Highly Efficient Gradient Boosting Decision Tree," in *Advances in Neural Information Processing Systems,* vol. 30, 2017, pp. 3146–3154. [Online]. Available: <https://proceedings.neurips.cc/paper/2017/hash/6449f44a102fde848669bdd9eb6b76fa-Abstract.html>