

PENILAIAN KOMPARATIF METODE KLASIFIKASI NEURAL NETWORK DAN RANDOM FOREST UNTUK KNOWLEDGE DISCOVERY PADA PENYAKIT DIABETES

Muhammad Aqil Zidane, Rihan Naufaldihanif, Aisha Nuraini Kusuma,
Bryan Hanggara, Adley Clark Peter Wijaya, Ken Ditha Tania, Winda Kurnia Sari

Sistem Informasi, Universitas Sriwijaya
Jl. Raya Palembang - Prabumulih No.KM. 32, Ogan Ilir, Indonesia
09031182227006@student.unsri.ac.id

ABSTRAK

Penelitian ini membahas penerapan data mining dalam prediksi diabetes, yang menjadi isu penting dalam bidang kesehatan dan teknologi informasi. Permasalahan utama yang diangkat adalah tingginya angka penderita diabetes yang sering terlambat terdiagnosis, yang berdampak pada meningkatnya risiko komplikasi serius dan biaya perawatan yang tinggi. Oleh karena itu, penelitian ini bertujuan untuk mengembangkan model prediksi diabetes yang lebih akurat menggunakan teknik *data mining*. Metode yang digunakan dalam penelitian ini adalah pendekatan kuantitatif dengan teknik pengumpulan data melalui analisis dataset diabetes menggunakan algoritma klasifikasi seperti *Random Forest*, dan *Neural Network*. Hasil penelitian menunjukkan bahwa algoritma *Random Forest* memiliki tingkat akurasi tertinggi sebesar 96,88% dibandingkan metode *Neural Network* dengan akurasi sebesar 89,23%, yang mengindikasikan bahwa metode *Random Forest* ini efektif dalam mendeteksi potensi *pre-diabetes* lebih dini. Kesimpulan dari penelitian ini menegaskan bahwa pemanfaatan data mining dapat meningkatkan akurasi prediksi *pre-diabetes* serta memberikan rekomendasi bagi tenaga medis dalam pengambilan keputusan diagnostik.

Kata kunci : Diabetes, Data Mining, Neural Network, Random Forest, Knowledge Discovery

1. PENDAHULUAN

Diabetes menjadi salah satu jenis penyakit yang termasuk paling banyak penderita di Indonesia [4]. Salah satu jenis diabetes yaitu Diabetes Melitus telah menjadi salah satu penyakit kronis dengan pertumbuhan tercepat di dunia yang menimbulkan tantangan serius bagi sistem kesehatan global. Menurut Internasional Diabetes Federation (IDF), pada tahun 2021, 10.5% dari populasi orang dewasa atau sekitar 537 juta orang dewasa berusia 20-79 tahun hidup dengan diabetes, dan angka ini diproyeksikan meningkat sebesar 46% atau 783 juta orang dewasa pada tahun 2045 [14]. Peningkatan Diabetes dari tahun ke tahun ini disebabkan oleh berbagai faktor, seperti pola makan dan gaya hidup yang tidak teratur, kadar glukosa yang tinggi dalam darah yang dikarenakan kekurangan insulin pada organ pankreas [4] [5] [8]

[11] [12] [13] [16]. Jika diabetes tidak ditangani dengan benar, diabetes sering kali akan berujung pada komplikasi serius seperti kerusakan organ tubuh yaitu penyakit kardiovaskular pada jantung, gagal ginjal, obesitas, kanker, dan penyakit neuropati pada saraf yang dimana dapat menyebabkan kegagalan fungsi jaringan dan organ dalam waktu jangka panjang dan dapat menurunkan kualitas hidup penderitanya dan meningkatkan beban ekonomi pada sistem kesehatan [3][4][7][10][11][12][14].

Dalam upaya mengatasi permasalahan tersebut, maka dilakukan penerapan topik pemodelan dalam data mining berdasarkan faktor-faktor kesehatan menjadi sangat penting.

Pemodelan yang digunakan dalam penelitian ini yaitu *Neural Network* dan *Random Forest*. Kedua

model ini termasuk ke dalam tipe *supervised learning* yang dimana membutuhkan data pelatihan untuk membangun suatu model klasifikasinya [10]. *Random Forest* dipilih karena kemampuannya dalam menangani data dengan variabel yang kompleks, mengurangi *overfitting*, serta memberikan interpretasi yang lebih baik dibandingkan model lain yang cenderung kurang transparan. Penelitian sebelumnya menunjukkan bahwa *Random Forest* mampu mencapai akurasi di angka 97,88% dalam prediksi diabetes, menjadikannya metode yang lebih andal dibandingkan model lain seperti SVM atau *Naïve Bayes* yang memiliki akurasi lebih rendah, sekitar 91-95% [19]. Melalui kedua model ini, klasifikasi menggunakan model *Neural Network* dan *Random Forest* dapat secara otomatis mendeteksi pola dan kumpulan data-data yang direpresentasikan oleh kumpulan kata dengan banyak fitur yang beragam sehingga dapat mengolah data-data tersebut dengan baik [17].

Metode yang digunakan dalam analisis mendalam terhadap berbagai variabel kesehatan, seperti kadar glukosa darah, tekanan darah, indeks massa tubuh (BMI), dan riwayat keluarga, untuk mengidentifikasi individu dengan risiko tinggi secara lebih akurat [16]. Penelitian yang diterbitkan dalam *Journal of Applied Informatics and Computing* menunjukkan bahwa penggunaan algoritma neural network dalam klasifikasi data mencapai akurasi sebesar 80,36% yang dimana temuan ini menegaskan bahwa teknik data mining dapat meningkatkan ketepatan diagnosis dan memungkinkan intervensi dini yang lebih efektif[20].

Dengan demikian, penerapan pemodelan topik

ini diperlukan untuk mendapatkan wawasan melalui *knowledge discovery* tentang pola yang muncul. Lalu, penelitian ini akan menghasilkan output berupa *knowledge discovery* yang dimana akan menjadi sebuah penelitian penting untuk mendukung pengambilan keputusan [17]. Selain itu, penelitian ini akan melakukan klasifikasi dan mendapatkan akurasi yang tinggi sehingga kedepannya penelitian ini dapat membantu dalam personalisasi perawatan, mempermudah tenaga medis dalam merancang strategi pengobatan yang disesuaikan dengan profil risiko masing-masing pasien. Hal tersebut tidak hanya meningkatkan efektivitas terapi, tetapi juga meminimalkan terjadinya komplikasi penyakit di masa mendatang. Selain itu, pemahaman yang mendalam tentang pola dan determinan tentang penyakit diabetes melalui analisis data berkontribusi pada pengembangan kebijakan kesehatan yang lebih tepat sasaran dan efisien. Oleh karena itu, penerapan teknik klasifikasi data mining dalam penanganan diabetes merupakan langkah yang strategis untuk menekan laju prevalensi dan meningkatkan kualitas hidup penderita diabetes kedepannya.

2. TINJAUAN PUSTAKA

Penelitian sebelumnya membahas implementasi *Data Mining* dengan Algoritma Naïve Bayes untuk mengklasifikasi kelayakan penerima bantuan sembako. Masalah utamanya dalam ketepatan sasaran bantuan sembako. Metode yang digunakan yaitu menggunakan algoritma Naïve Bayes. Hasil utama penelitiannya didapatkan dengan akurasi sebesar 86%, recall 85%, dan presisi 88%. Untuk penelitian kedepannya, disarankan untuk menggunakan variasi data training, data testing, dan atribut untuk mendapatkan model dengan akurasi terbaik. Future Works dapat mencakup peningkatan akurasi dengan variasi data dan atribut yang lebih luas [1].

Paper sebelumnya membahas masalah yang diidentifikasi adalah gangguan psikologis yang memengaruhi kehidupan sehari-hari. Metode yang digunakan adalah algoritma C4.5 untuk mengklasifikasi gangguan psikologis. Hasil utama menunjukkan akurasi data training sebesar 57.50% dan data testing sebesar 72.67%. Future Works dapat mencakup peningkatan akurasi diagnosa gangguan psikologis hingga 70%. Penelitian ini penting untuk deteksi dini gangguan mental [2].

Penelitian sebelumnya mengkaji penggunaan algoritma C4.5 dalam klasifikasi penderita diabetes. Penelitian ini menunjukkan bahwa algoritma C4.5 mampu menghasilkan akurasi sebesar 85% dalam mengklasifikasikan pasien diabetes berdasarkan atribut seperti glukosa, BMI, dan tekanan darah. Metode ini menunjukkan potensi yang tinggi untuk digunakan dalam diagnosis dini diabetes, yang dapat mengurangi risiko komplikasi di masa depan [3].

Selain itu, penelitian lain oleh menggunakan algoritma k-NN untuk prediksi penyakit diabetes

dengan normalisasi data. Hasilnya menunjukkan bahwa proses normalisasi dapat meningkatkan akurasi model, dengan akurasi tertinggi mencapai 74% ketika menggunakan metode normalisasi Min-Max [4].

Pada paper lainnya dengan membahas masalah respon masyarakat terhadap kabar harian Covid-19 dari Twitter Kementerian Kesehatan yang tidak selaras. Metode yang digunakan adalah klasifikasi *Naive Bayes Classifier*. Hasil utamanya adalah klasifikasi sentimen menjadi positif, negatif, dan netral. Future works dapat mencakup pengembangan kamus sendiri, pengujian metode klasifikasi lain, dan peningkatan akurasi [5].

Dalam paper yang membahas analisis tentang aplikasi Pluang menggunakan algoritma *Naive Bayes* and *Support Vector Machine* (SVM). Masalah yang diidentifikasi adalah analisis sentimen terhadap ulasan aplikasi Pluang. Metode yang digunakan adalah Naive Bayes dan SVM. Hasil penelitian menunjukkan bahwa model SVM lebih baik daripada Naive Bayes dengan akurasi 99.50%, presisi 99.67%, recall 99.33%, dan F1 score 99.50%. Future works dapat mencakup peningkatan akurasi dengan menambah jumlah data penelitian dan penambahan parameter pada pengujian SVM seperti parameter literasi, gamma, dan lambda pada kernel. SVM efektif untuk menganalisis sentimen ulasan aplikasi Pluang[6].

Penelitian sebelumnya tentang Klasifikasi, Gejala, Diagnosis, Pencegahan, dan Pengobatan Diabetes Melitus. Masalah yang diidentifikasi adalah peningkatan prevalensi diabetes yang berdampak pada kesejahteraan masyarakat. Metode yang digunakan adalah *review* literatur untuk menggambarkan klasifikasi, gejala, diagnosis, pencegahan, dan pengobatan diabetes. Hasil utama dari paper ini adalah pentingnya pengenalan dini *pre-diabetes* dan peran gaya hidup sehat dalam pencegahan diabetes. Future works dapat mencakup penelitian lebih lanjut tentang terapi obat-obatan dan intervensi gaya hidup untuk mengurangi risiko diabetes [7].

Peneliti lain melakukan klasifikasi data menggunakan algoritma C4.5, yang mengembangkan ID3 dengan menangani data kosong dan pemangkasan pohon keputusan. Metode ini efektif dalam menyederhanakan pengambilan keputusan dengan kompleksitas $O(m.n^2)$ [8].

Di sisi lain, peneliti meneliti klasifikasi lima jenis kayu menggunakan *Local Binary Pattern* (LBP) dan *Support Vector Machine* (SVM). Dari 750 citra kayu, metode ini mencapai akurasi 91,3% pada sigma 0,3 dengan error 8,7%, membuktikan efektivitasnya dalam otomatisasi klasifikasi kayu [9].

Pada tahun 2020, telah dilakukan penelitian terkait dengan judul Penerapan Metode Klasifikasi K-Nearest Neighbor pada Dataset Penderita Penyakit Diabetes. Masalah yang diangkat adalah pengukuran performa metode klasifikasi dalam

mengklasifikasikan dataset penderita diabetes, khususnya dengan menggunakan algoritma *K-Nearest Neighbor* (KNN). Dalam algoritma yang digunakan, performa diukur berdasarkan akurasi, presisi, *recall*, dan *F-Measure* dengan nilai K yang berbeda (K=3, 4, dan 5). Hasil utama menunjukkan bahwa akurasi tertinggi adalah 39% pada K=3, presisi tertinggi 65% pada K=3 dan K=5, *recall* tertinggi 36% pada K=3, dan *F-Measure* tertinggi 46% pada K=3. Namun, hasil ini dinilai belum optimal karena jumlah data yang digunakan relatif kecil. Maka kedepannya dapat dilakukan penelitian terbaru dengan penambahan jumlah data dan menerapkan *cross-validation* agar performa model meningkat [10].

Penelitian lain juga, membahas implementasi algoritma Naïve Bayes untuk klasifikasi penyakit diabetes pada perempuan. Dataset yang digunakan berasal dari Kaggle, dengan total 9 atribut. RapidMiner digunakan sebagai tools untuk menguji akurasi, *precision*, *recall*, dan AUC. Hasil penelitian menunjukkan bahwa metode Naïve Bayes memberikan akurasi sebesar 78.50%, *precision* 85.24%, *recall* 83.64%, dan AUC 0.855, yang menunjukkan performa yang cukup baik dalam mendeteksi diabetes pada perempuan [11].

Penelitian oleh Zudyanti Dwi Rahma Sari menggunakan algoritma C4.5 untuk membantu para medis dalam mengklasifikasi pasien diabetes berdasarkan gejala yang ada. Dengan menggunakan RapidMiner, penelitian ini berhasil mencapai akurasi 95.90% pada pengujian *10-Fold Cross Validation*, menunjukkan efektivitas metode ini dalam mendeteksi diabetes secara dini [12].

3. METODE PENELITIAN

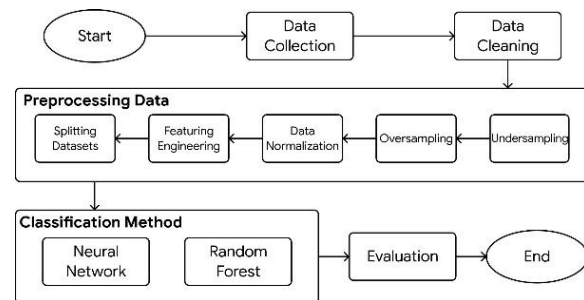
3.1. Data Collection

Tabel 1. Data Feature Mentah

Feature		
HighBP	HighChol	CholCheck
BMI	Smoker	Stroke
HeartDiseaseorAttack	PhysActivity	Fruits
Veggies	HvyAlcoholConsump	AnyHealthcare
NoDocbcCost	GenHlth	MenHlth
PhysHlth	DiffWalk	Sex
Age	Education	Income

Data yang digunakan dalam penelitian ini diperoleh dari “Diabetes Health Indicators Dataset” (<https://www.kaggle.com/datasets/alexteboul/diabetes-health-indicators-dataset>), yang tersedia di platform Kaggle. Dataset ini berasal dari Behavioral Risk Factor Surveillance System (BRFSS), sebuah survei kesehatan tahunan yang dilakukan oleh Centers for Disease Control and Prevention (CDC) di Amerika Serikat. BRFSS bertujuan untuk mengumpulkan informasi terkait faktor risiko kesehatan, kebiasaan gaya hidup, serta kondisi medis

yang dapat memengaruhi kesehatan populasi dewasa. Data dalam dataset ini mencakup berbagai indikator kesehatan yang berhubungan dengan diabetes, seperti riwayat medis, kebiasaan merokok, serta status tekanan darah dan kolesterol. Dengan cakupan data yang luas dan metodologi pengumpulan yang sistematis, dataset ini menjadi sumber yang relevan dan terpercaya untuk menganalisis faktor-faktor yang berkontribusi terhadap diabetes dalam populasi.



Gambar 1. Tahapan penelitian

3.2. Data Cleaning

Pada tahap data cleaning, telah dipastikan bahwa dataset yang digunakan berada dalam kondisi bersih dan siap untuk dianalisis. Evaluasi menyeluruh terhadap kualitas data telah dilakukan, mencakup pengecekan *missing values* (data yang hilang), *duplicate entries* (entri duplikat), dan inkonsistensi data. Hasil evaluasi menunjukkan bahwa tidak ditemukan adanya *missing values* dalam dataset, yang mengindikasikan bahwa data setiap entri data telah lengkap dan tidak memerlukan proses imputasi atau penghapusan data lebih lanjut. Selain itu, pemeriksaan terhadap kemungkinan adanya duplikasi data juga tidak menemukan entri yang terduplikasi, baik secara identik maupun yang mirip namun tidak relevan. Hal ini memastikan bahwa dataset terbebas dari redundansi yang dapat memengaruhi hasil analisis. Selanjutnya, pengecekan terhadap format dan konsistensi data, seperti keseragaman kategori, dan tipe data yang sesuai juga menunjukkan bahwa dataset terbebas dari anomali atau inkonsistensi yang dapat mengganggu interpretasi. Dengan demikian, dapat disimpulkan bahwa dataset yang digunakan telah melalui validasi kualitas yang ketat dan berada dalam kondisi optimal untuk tahap analisis lebih lanjut.

3.3. Preprocessing Data

Pada tahap *pre-processing* dalam penelitian ini, dilakukan beberapa langkah kritis untuk mempersiapkan data sebelum proses analisis atau pemodelan lebih lanjut. Langkah-langkah tersebut meliputi *oversampling* data, normalisasi data, *feature engineering*, dan pembagian *dataset*.

Pertama-tama, setelah dilakukan peninjauan data, ditemukan bahwa dalam *dataset* yang digunakan, terdapat ketidakseimbangan kelas yang signifikan, di mana jumlah sampel kategori sehat jauh lebih besar dibandingkan dengan kategori diabetes

dan *pre-diabetes*. Ketimpangan ini dapat menyebabkan model prediktif lebih cenderung mengklasifikasikan data ke kategori sehat, sehingga mengurangi akurasi prediksi terhadap individu yang termasuk dalam kategori diabetes dan *pre-diabetes*.

Oleh karena itu, diterapkan *oversampling* dengan metode Borderline-SMOTE (*Synthetic Minority Over-sampling Technique*) untuk mensintesis sampel baru pada kategori diabetes dan *pre-diabetes*. Borderline-SMOTE bekerja dengan menggunakan algoritma *k-Nearest Neighbors* (*k-NN*), di mana untuk setiap sampel dalam kelas minoritas, dipilih beberapa sampel tetangga terdekat, lalu data sintetis dibuat dengan interpolasi antara sampel asli dan tetangga terpilih. Dengan cara ini, jumlah sampel pada kategori minoritas bertambah tanpa sekadar menduplikasi data yang sudah ada.

Dengan metode ini, distribusi kelas dalam dataset menjadi lebih seimbang, sehingga model dapat belajar dengan lebih optimal tanpa bias terhadap kategori mayoritas.

Setelah proses penyeimbangan data, selanjutnya akan dilakukan normalisasi data dengan menggunakan *Standard Scaler*. Normalisasi ini bertujuan untuk menyamakan skala dari setiap fitur dalam dataset dengan mentransformasi data sehingga memiliki mean (rata-rata) 0 dan standar deviasi 1. Dengan normalisasi, proses pelatihan model menjadi lebih stabil dan hasil prediksi lebih akurat.

```
scaler = StandardScaler()
scaler.fit(X_train)

X_train_scaled = scaler.transform(X_train)
X_test_scaled = scaler.transform(X_test)
```

Gambar 2. Proses scaling data

Kedua, dilakukan *feature engineering* untuk mengeliminasi fitur-fitur yang kurang relevan atau tidak signifikan terhadap prediksi. Beberapa fitur dihilangkan berdasarkan analisis korelasi dan pengaruh antara *feature* dan *class*. Fitur-fitur yang dieliminasi meliputi: *AnyHealthcare* (akses ke layanan kesehatan), *Sex* (jenis kelamin), dan *Smoker* (perokok). Penghapusan fitur-fitur ini bertujuan untuk menyederhanakan model dan meningkatkan efisiensi serta akurasi prediksi.

Tabel 2. Penyeleksian Feature

Feature	Class
HighBp	Diabetes
BMI	
HighChol	
Stroke	
HeartDiseaseorAttack	
PhyscActivity	
Fruits	
Veggies	

Feature	Class
HvyAlcoholConsump	
DiffWalk	
CholCheck	
Age	
NoDocbcCost	
Income	
Education	
PhysHlth	
MenHlth	
GenHlth	

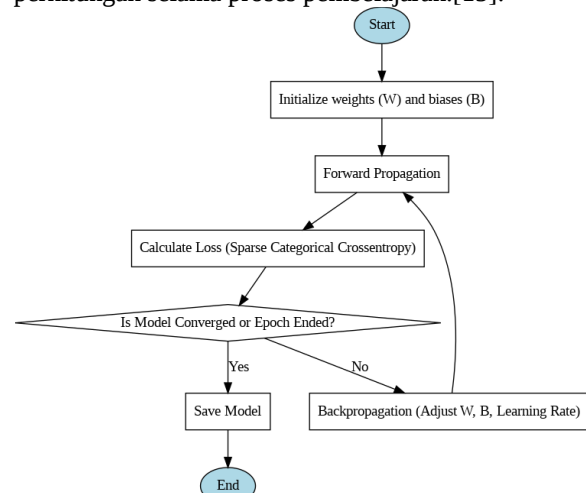
Ketiga, pada model *Neural Network*, dataset dibagi menjadi tiga bagian dengan komposisi sebagai berikut: *training set* (70%), *validation set* (25%), dan *test set* (5%). Pada model *Random Forest*, dataset dibagi menjadi dua bagian dengan komposisi sebagai berikut: *training set* (80%) dan *test set* (20%). Pembagian ini dilakukan untuk memastikan evaluasi model yang komprehensif dan menghindari *overfitting*.

Dengan melakukan langkah-langkah *pre-processing* tersebut, dataset menjadi lebih siap untuk digunakan dalam proses pemodelan, sehingga diharapkan dapat menghasilkan prediksi yang lebih akurat dan andal.

3.4. Classification Method

Setelah melakukan tahap *pre-processing*, tahap selanjutnya yaitu melakukan pemodelan. Ada dua model yang digunakan untuk melakukan pemodelan, yaitu *Neural Network* dan *Random Forest*.

Neural network merupakan sistem pengolahan data yang menyerupai jaringan saraf biologis yang ada di manusia. *Neural network* merupakan model buatan yang mereplikasi fungsi otak manusia dengan menggunakan komputer untuk melakukan berbagai perhitungan selama proses pembelajaran.[15].



Gambar 3. Flowchart Neural Network

```
# Stops training if 'val_loss' doesn't improve after 10 epochs and
# restores the best weights
early_stopping = EarlyStopping(monitor='val_loss',
                                patience=10,
                                restore_best_weights=True)
```

Gambar 4. Early stopping callback

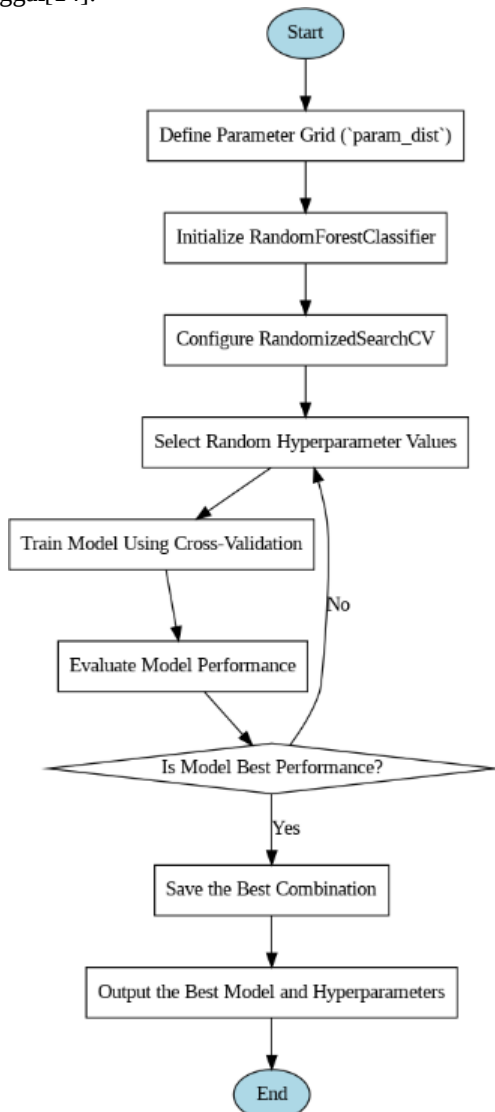
```
# Reduce learning rate when a metric has stopped improving.
lr_scheduler = keras.callbacks.ReduceLROnPlateau
(monitor='val_loss', factor=0.5, patience=5, min_lr=1e-6)
```

Gambar 5. Exponential decay learning rate schedule

```
# Compile the model with Adam optimizer and sparse categorical
cross-entropy loss for multi-class classification
model.compile(optimizer=tf.keras.optimizers.Adam(learning_rate=0
.001),
              loss='sparse_categorical_crossentropy',
              metrics=['accuracy'])
```

Gambar 6. Model compilation

Random Forest merupakan algoritma *machine learning* berbasis *ensemble* yang menggabungkan hasil dari beberapa *decision tree* untuk meningkatkan akurasi prediksi dan mengurangi risiko *overfitting*. Algoritma ini bekerja dengan membangun banyak pohon keputusan secara acak, baik dari segi data sampel maupun fitur yang dipilih, sehingga menghasilkan model yang lebih stabil dan akurat dibandingkan dengan satu pohon keputusan tunggal[14].



Gambar 7. Flowchart Random Forest

```
# Define parameter grid
param_dist = {
    'n_estimators': [10, 50, 100],
    'max_depth': [10, 20, 30, None],
    'min_samples_split': [2, 5, 10],
    'min_samples_leaf': [1, 2, 4],
    'bootstrap': [True, False]
}

# Model Initialization
rdF = RandomForestClassifier(random_state=42, n_jobs=-1)

# RandomizedSearchCV
random_search = RandomizedSearchCV(
    estimator=rdF,
    param_distributions=param_dist,
    n_iter=10,
    cv=3,
    verbose=2,
    random_state=42,
    n_jobs=-1
)
```

Gambar 8. Random Forest Parameter setup

```
Best Parameters: {'n_estimators': 50, 'min_samples_split': 10,
'min_samples_leaf': 2, 'max_depth': None, 'bootstrap': False}
```

Gambar 9. Best Random Forest parameter Result

3.5. Evaluation

Pada tahap ini, dilakukan evaluasi kinerja model dengan menguji data pada setiap kelas. Tujuan dari pengujian ini adalah untuk menentukan tingkat akurasi terbaik dari masing-masing kelas. Akurasi dihitung sebagai persentase dari total data yang berhasil diidentifikasi dengan benar. Penilaian ini dilakukan menggunakan *Accuracy*, *Confusion Matrix*, *precision*, *recall*, dan *f1-score* [18].

Ukuran penilaian tersebut adalah sebagai berikut :

3.6. Confusion Matrix

Confusion Matrix merupakan sebuah metode untuk evaluasi yang menggunakan tabel matriks seperti tabel 2 berikut:

Tabel 3. Confusion Matrix

Aktual	Prediksi		
	0	1	2
0	TP 0	FN 1	FN 2
1	FP 0	TP 1	FN 2
2	FP 0	FP 1	TP 2

3.7. Accuracy (A)

Accuracy adalah rasio prediksi model yang benar dari total prediksi model yang dilakukan. *Accuracy* merepresentasikan seberapa sering model memprediksi kelas yang benar termasuk kelas positif maupun negatif.

$$A = \frac{\sum_i(TP_i)}{\sum_i(TP_i + FP_i + FN_i)} \dots (1)$$

3.8. Precision (P)

Precision adalah rasio prediksi model kelas positif yang benar dari total prediksi model kelas positif. *Precision* mempresentasikan seberapa sering model memprediksi kelas positif dengan benar diantara total kelas positif.

$$P = \frac{TP_i}{TP_i + FP_i} \dots (2)$$

3.9. Recall (R)

Recall adalah rasio prediksi model kelas positif yang benar dari total sampel data kelas positive. *Recall* akan menunjukkan kemampuan model kita untuk mengenali *review* positif yang sebenarnya.

$$R = \frac{TP_i}{TP_i + FN_i} \dots (3)$$

3.10. F1-score (F1)

F1-score adalah matrik evaluasi yang menggabungkan *Precision* dan *Recall* dalam memperoleh evaluasi performa yang lebih seimbang.

$$F1 = 2 \cdot \frac{P_i \cdot R_i}{P_i + R_i} \dots (4)$$

4. HASIL DAN PEMBAHASAN

4.1. Neural Network Confusion Matrix

Tabel 4. Confusion Matrix Model NN

Prediction/ Actual	0	1	2	Total
0	8204	234	941	9379
1	128	9371	33	9532
2	1516	312	7770	9598
Total	9848	9917	8744	28509

Pada *Confusion Matrix Model Neural Network*, terdapat 8.204 sampel yang diprediksi benar adanya sehat (kelas 0), 9.371 sampel yang diprediksi benar adanya pre-diabetes (kelas 1), dan 7.770 sampel yang diprediksi benar adanya diabetes (kelas 2). Namun masih terdapat kesalahan klasifikasinya, terdapat 1.516 sampel yang seharusnya berlabel sehat (kelas 0) dan 941 sampel yang diprediksi sehat, namun pada kondisi aktualnya diabetes (kelas 0).

4.2. Random Forest Confusion Matrix

Tabel 5. Confusion Matrix Model Random Forest

Prediction/ Actual	0	1	2	Total
0	35698	34	2051	37783
1	850	37066	342	38258
2	5006	172	32814	37992
Total	41554	37272	35207	114033

Pada *Confusion Matrix Model Random Forest*, terdapat 35.698 sampel yang diprediksi benar adanya sehat (kelas 0), 37.066 sampel yang diprediksi benar adanya pre-diabetes (kelas 1), dan 32.814 sampel yang diprediksi benar adanya diabetes (kelas 2). Namun masih terdapat kesalahan klasifikasinya, terdapat 2.051 sampel diprediksi diabetes tapi kenyataannya sehat (kelas 0) dan 5.006 sampel diprediksi sehat tapi kenyataannya diabetes (kelas 2).

4.3. Neural Network Classification Report

Tabel 6. Classification Report Model Neural Network

Class	Precision	Recall	F1 Score
0	0.83	0.87	0.85
1	0.94	0.98	0.96
2	0.89	0.81	0.85
Train Accuracy			89.23%
Validation Accuracy			89.16%
Testing Accuracy			89%

Pada model *Neural Network*, Hasil *classification report* menunjukkan bahwa model memiliki performa yang baik secara keseluruhan, dengan akurasi pelatihan 89.23%, validasi 89.16%, dan pengujian 89%. Model paling baik dalam mengenali kelas pre-diabetes (*Recall* 98%), yang berarti sangat sedikit kasus pre-diabetes yang salah diklasifikasikan. Namun, model sedikit kesulitan dalam mengklasifikasikan kelas diabetes (*Recall* 81%), yang menunjukkan bahwa beberapa kasus diabetes masih sering diklasifikasikan ke kelas lain. Dibandingkan dengan *Random Forest*, performa model ini sedikit lebih rendah dalam hal *recall* untuk kelas diabetes, sehingga mungkin diperlukan teknik tambahan seperti *balancing data* atau *fine-tuning parameter* untuk meningkatkan akurasi lebih lanjut.

4.4. Random Forest Classification Report

Tabel 7. Classification Report Model Random Forest

Class	Precision	Recall	F1 Score
0	0.86	0.94	0.90
1	0.99	0.97	0.98
2	0.93	0.86	0.90
Training Accuracy			96.88%
Validation Accuracy			92.58%

Hasil *classification report* menunjukkan bahwa model *Random Forest* memiliki performa yang sangat baik dalam mendeteksi Pre-Diabetes dengan 0.96 *Accuracy*, *Precision* 0.99 dan *Recall* 0.97, menghasilkan *F1 Score* 0.98. Model juga cukup baik dalam mengenali Non-Diabetes dengan *Recall* 0.94, meskipun *Precision*-nya lebih rendah di 0.86. Sementara itu, untuk kelas Diabetes, *Recall* lebih rendah di 0.86, menunjukkan bahwa masih ada kasus Diabetes yang tidak terdeteksi dengan baik. Secara keseluruhan, model memiliki akurasi pelatihan 96.88% dan akurasi validasi 92.58%, menunjukkan kemampuan generalisasi yang cukup baik, meskipun performanya masih bisa ditingkatkan dalam mendeteksi kasus Diabetes.

Pada kedua model, hasil penelitian menunjukkan bahwa baik model *Neural Network* dan *Random Forest* memiliki performa yang lebih tinggi dalam mendeteksi kondisi *pre-diabetes* dibandingkan dengan diabetes. Hal ini terlihat dari nilai *Precision* sebesar 0.99 dan *Recall* 0.97 pada kelas *pre-diabetes* (kelas 1), yang lebih tinggi dibandingkan dengan kelas diabetes (kelas 2) yang memiliki *Recall* 0.86.

Artinya, model lebih mampu mengenali individu yang berada dalam tahap *pre-diabetes* dengan akurasi yang lebih baik, sementara kemampuan untuk mengidentifikasi kasus diabetes masih kurang optimal. Hal ini mengindikasikan bahwa pola dalam data *pre-diabetes* lebih mudah dikenali oleh model, kemungkinan karena karakteristiknya yang lebih jelas dibandingkan dengan variasi kompleks yang terjadi pada diabetes. Oleh karena itu, diperlukan pendekatan tambahan seperti analisis data *real-time* atau tren longitudinal untuk meningkatkan sensitivitas model dalam mendeteksi kasus diabetes secara lebih akurat.

5. KESIMPULAN DAN SARAN

Dari penerapan 2 metode klasifikasi pada klasifikasi penyakit diabetes, yaitu *Neural Network* dan *Random Forest* dapat disimpulkan bahwa Model *Neural Network* dan *Random Forest* menunjukkan performa yang cukup baik dalam klasifikasi kondisi kesehatan berdasarkan *Confusion Matrix* dan *Classification Report*. Pada *Neural Network*, terdapat 8.204 sampel yang diklasifikasikan benar sebagai sehat (kelas 0), 9.371 sampel benar sebagai *pre-diabetes* (kelas 1), dan 7.770 sampel benar sebagai diabetes (kelas 2). Namun, masih terdapat kesalahan klasifikasi, seperti 1.516 sampel diabetes yang diklasifikasikan sebagai sehat. Dari hasil *classification report*, model ini memiliki akurasi validasi 89.16%, dengan *Precision* 0.94, *Recall* 0.98, dan *F1 Score* 0.96 untuk *pre-diabetes*, tetapi *Recall* lebih rendah untuk diabetes (0.81), menunjukkan bahwa beberapa kasus masih salah diklasifikasikan. Sementara itu, pada *Random Forest*, terdapat 35.698 sampel yang benar diklasifikasikan sebagai sehat, 37.066 sebagai *pre-diabetes*, dan 32.814 sebagai diabetes. Namun, masih terdapat 5.006 sampel yang seharusnya diabetes tetapi diklasifikasikan sebagai sehat. Model ini memiliki akurasi validasi lebih tinggi, yaitu 92.58%, dengan *Precision* 0.99, *Recall* 0.97, dan *F1 Score* 0.98 untuk *pre-diabetes*. Model juga cukup baik dalam mengenali kasus sehat (*Recall* 0.94), tetapi masih memiliki *Recall* lebih rendah pada diabetes (0.86). Secara keseluruhan, *Random Forest* memiliki performa lebih unggul dalam hal akurasi dan generalisasi, tetapi masih perlu penyempurnaan dalam mendeteksi kasus diabetes agar tidak terjadi *underdiagnosis*. Dengan demikian, berdasarkan hasil penelitian ini, terdapat beberapa saran untuk penelitian selanjutnya, seperti menerapkan metode klasifikasi lain dengan menggabungkan teknik optimasi dan menggunakan lebih banyak data untuk meningkatkan kualitas topik dan skor akurasi dari kedua model tersebut.

DAFTAR PUSTAKA

[1] International Diabetes Federation, "Diabetes Facts and Figures," International Diabetes Federation. Accessed: Mar. 03, 2025. [Online].

Available: <https://idf.org/about-diabetes/diabetes-facts-figures/>

- [2] A. Damuri, U. Riyanto, H. Rusdianto, and M. Aminudin, "Implementasi Data Mining dengan Algoritma Naïve Bayes Untuk Klasifikasi Kelayakan Penerima Bantuan Sembako," *JURIKOM (Jurnal Riset Komputer)*, vol. 8, no. 6, p. 219, Dec. 2021, doi: 10.30865/jurikom.v8i6.3655.
- [3] S. Febriani and H. Sulistiani, "Analisis Data Hasil Diagnosa Untuk Klasifikasi Gangguan Kepribadian Menggunakan Algoritma C4.5," *Jurnal Teknologi dan Sistem Informasi (JTSI)*, vol. 2, no. 4, 2021, doi: <https://doi.org/10.33365/jtsi.v2i4.1373>.
- [4] B. A. R. P. Wahyu, A. F. Faroz, C. P. Mahendra, and R. K. Hapsari, "Klasifikasi Penderita Penyakit Diabetes Berdasarkan Decision Tree Menggunakan Algoritma C4.5," *INTEGER: Journal of Information Technology*, vol. 8, no. 1, Mar. 2023, doi: <https://doi.org/10.31284/j.integer.2023.v8i1.4423>.
- [5] M. Sholeh, D. Andayati, and Rr. Y. Rachmawati, "Data Mining Model Klasifikasi Menggunakan Algoritma K-Nearest Neighbor Dengan Normalisasi Untuk Prediksi Penyakit Diabetes," *TelKa*, vol. 12, no. 02, pp. 77–87, Oct. 2022, doi: 10.36342/teika.v12i02.2911.
- [6] E. T. Handayani and A. Sulistiyawati, "Analisis Sentimen Respon Masyarakat Terhadap Kabar Harian Covid-19 Pada Twitter Kementerian Kesehatan Dengan Metode Klasifikasi Naive Bayes," *Jurnal Teknologi dan Sistem Informasi*, vol. 2, no. 3, Sep. 2021, doi: <https://doi.org/10.33365/jtsi.v2i3.906>.
- [7] B. A. Maulana, M. J. Fahmi, A. M. Imran, and N. Hidayati, "Analisis Sentimen Terhadap Aplikasi Pluang Menggunakan Algoritma Naive Bayes dan Support Vector Machine (SVM)," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 2, pp. 375–384, Feb. 2024, doi: 10.57152/malcom.v4i2.1206.
- [8] D. Hardianto, "TELAH KOMPREHENSIF DIABETES MELITUS: KLASIFIKASI, GEJALA, DIAGNOSIS, PENCEGAHAN, DAN PENGOBATAN," *Jurnal Bioteknologi & Biosains Indonesia (JBBi)*, vol. 7, no. 2, pp. 304–317, Jan. 2021, doi: 10.29122/jbbi.v7i2.4209.
- [9] P. B. N. Setio, D. R. S. Saputro, and B. Winarno, "Klasifikasi engan Pohon Keputusan Berbasis Algoritme C4.5," *PRISMA, Prosiding Seminar Nasional Matematika*, vol. 3, pp. 64–71, Feb. 2020.
- [10] N. Neneng, N. U. Putri, and E. R. Susanto,

- “Klasifikasi jenis kayu menggunakan support vector machine berdasarkan ciri tekstur local binary pattern,” *CYBERNETICS*, vol. 4, no. 02, Jan. 2021, doi: 10.29406/cbn.v4i02.2324.
- [11] A. M. Argina, “Penerapan Metode Klasifikasi K-Nearest Neighbor pada Dataset Penderita Penyakit Diabetes,” *Indonesian Journal of Data and Science*, vol. 1, no. 2, pp. 29–33, Jul. 2020, doi: 10.33096/ijodas.v1i2.11.
- [12] A. Veronica Agustin and A. Voutama, “Implementasi Data Mining Klasifikasi Penyakit Diabetes Pada Perempuan Menggunakan Naïve Bayes,” *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 7, no. 2, pp. 1002–1007, Aug. 2023, doi: 10.36040/jati.v7i2.6808.
- [13] Z. D. R. Sari, Jasmir Jasmir, and Y. Arvita, “Penerapan Data Mining untuk prediksi penyakit diabetes menggunakan algoritma C4.5 zudyanti dwi rahma sari1, ja,” *Jurnal Informatika Dan Rekayasa Komputer (JAKAKOM)*, vol. 4, no. 1, pp. 827–834, Apr. 2024, doi: 10.33998/jakakom.2024.4.1.1624.
- [14] A. P. Siregar, D. P. Purba, J. P. Pasaribu, and K. R. Bakara, “Implementasi Algoritma Random Forest Dalam Klasifikasi Diagnosis Penyakit Stroke,” *Jurnal Nasional Komputasi dan Teknologi Informasi (JNKTI)*, vol. 2, no. 4, pp. 155–164, Nov. 2023, doi: 10.55606/juprit.v2i4.3039.
- [15] N. A. Santoso, N. Fadilah, R. D. Kurniawan, and A. Supratman, “Penerapan Neural Network Method dengan Struktur Backpropagation dalam Menentukan Prediksi Stock Barang,” *Remik: Riset dan E-Jurnal Manajemen Informatika Komputer*, vol. 7, no. 3, pp. 1585–1592, Aug. 2023.
- [16] A. N. P. Haryadi, W. Setiawan, and D. A. Fatah, “Penerapan Data Mining untuk Klasifikasi Penyakit Diabetes Menggunakan Metode Decision Tree,” *Jurnal Nasional Komputasi dan Teknologi Informasi (JNKTI)*, vol. 7, no. 6, pp. 1484–1495, Dec. 2024.
- [17] M. D. Pratiwi and K. D. Tania, “Knowledge Discovery Through Topic Modeling on GoPartner User Reviews Using BERTopic, LDA, and NMF,” *Journal of Applied Informatics and Computing*, vol. 9, no. 1, pp. 1–7, Jan. 2025, doi: 10.30871/jaic.v9i1.8782.
- [18] R. Indransyah, Y. H. Chrisnanto, and P. N. Sabrina, “Klasifikasi sentimen pergelaran MotoGP di Indonesia menggunakan algoritma Correlated Naïve Bayes Classifier,” *INFOTECH Journal*, vol. 8, no. 2, pp. 60–66, Dec. 2022. DOI: 10.31949/infotech.v8i2.3103.
- [19] W. Apriliah, I. Kurniawan, M. Baydhowi, and T. Haryati, “Prediksi Kemungkinan Diabetes pada Tahap Awal Menggunakan Algoritma Klasifikasi Random Forest,” *SISTEMASI: Jurnal Sistem Informasi*, vol. 10, no. 1, pp. 163–171, Jan. 2021.
- [20] A. Peryanto, A. Yudhana, and R. Umar, “Klasifikasi Citra Menggunakan Convolutional Neural Network dan K Fold Cross Validation,” *Journal of Applied Informatics and Computing*, vol. 4, no. 1, pp. 45–51, May 2020, doi: 10.30871/jaic.v4i1.2017