

Pemanfaatan Teknologi Machine Learning pada Klasifikasi Jenis Hipertensi Berdasarkan Fitur Pribadi

Pramesti Dewi^{*1}, Purwono Purwono², Safar Dwi Kurniawan³

¹Program Studi Keperawatan Universitas Harapan Bangsa, Purwokerto

²Program Studi Informatika Universitas Harapan Bangsa, Purwokerto

³Program Studi D3 Teknik Komputer Politeknik Harapan Bersama, Tegal

E-mail: ^{*1}pramestidewi@uhb.ac.id, ²purwono@uhb.ac.id, ³safar.kurniawan45@gmail.com

Abstrak

Hipertensi tampaknya menjadi faktor utama dalam perkembangan penyakit seperti stroke, gagal jantung, infark miokard, fibrilasi atrium, penyakit arteri perifer, dan diseksi aorta. Prediksi dini jenis hipertensi dari riwayat kesehatan merupakan hal yang penting agar kita dapat mengetahui penyakit yang disebabkan olehnya. Prediksi ini dapat diperoleh dengan memanfaatkan teknologi machine learning untuk menemukan pengetahuan baru dari data dasar sehingga menemukan pola yang valid, berguna, dan mudah dipelajari. Model klasifikasi random forest diusulkan dalam penelitian ini. Kontribusi kami dalam penelitian ini adalah membuat model klasifikasi random forest dengan teknik baru yaitu perbaikan data untuk melakukan tuning hyperparameter. Kami melihat peneliti sebelumnya hanya mengejar nilai akurasi yang tinggi semata. Berbeda dengan penelitian sebelumnya, kami menggunakan teknik optimasi hyperparameter gridsearch cv pada model klasifikasi random forest. Parameter terbaik untuk model random forest yaitu max_depth = 80, max_features = 3, min_samples_leaf = 3, min_samples_split = 8, dan n_estimators = 1000 yang direkomendasikan dari teknik gridsearch cv. Akurasi sebelum optimasi adalah 72,3%, sedangkan setelah optimasi adalah 86,1%. Hal ini menunjukkan peningkatan akurasi sebesar 13,7% setelah menerapkan metode grid search cv pada klasifikasi jenis hipertensi menggunakan model random forest.

Kata Kunci—hipertensi, klasifikasi, machine learning, teknologi

1. PENDAHULUAN

Kesehatan merupakan hal yang paling didambakan manusia dalam bentuk kesejahteraan baik secara fisik, mental dan sosial [1]. Tekanan darah ialah kekuatan yang diterapkan dengan mengedarkan darah menuju dinding arteri tubuh [2]. Terdapat dua jenis tekanan darah yaitu tekanan darah sistolik dan tekanan darah diastolik [2]. Tekanan darah sistolik ialah tekanan dalam pembuluh darah saat jantung berdetak sedangkan tekanan diastolik ialah tekanan yang terjadi diantara ketukan denyut jantung. Hipertensi didiagnosis jika nilai tekanan sistolik sama atau lebih besar daripada 140 mmHg [3] dan nilai tekanan diastolik lebih besar atau sama dengan 90 mmHg [4]. Hipertensi rupanya menjadi salah satu faktor utama terhadap terjadinya penyakit seperti stroke, gagal jantung, infark miokard, atrial fibrilasi, penyakit arteri perifer, dan diseksi aorta [5]. Hipertensi juga didefinisikan sebagai sindrom klinis dengan peningkatan tekanan vaskular [2].

Berdasarkan data yang dirujuk dari National High Blood Pressure Education Program, terdapat beberapa jenis tekanan darah yang terjadi pada manusia dewasa [2]. Tekanan darah disebut normal jika angka sistolik <120 mmHg dan diastolik <80 mmHg. Tekanan darah disebut *prehypertension* jika angka sistolik antara 120-139 mmHg atau diastolik antara 80-99 mmHg. Tekanan darah disebut *stage-1 hypertension* jika angka sistolik berkisar antara 140-159 mmHg atau diastolik antara 90-99 mmHg. Tekanan darah disebut *stage-2 hypertension* jika angka sistolik lebih dari 160 mmHg atau lebih dari sama dengan 100 mmHg. Prehipertensi merupakan tahapan awal seseorang sebelum menderita hipertensi yang mana tidak diikuti dengan gejala yang

dirasakan oleh individu yang mengalaminya, namun beresiko lebih tinggi untuk menjadi penyakit hipertensi hingga penyakit kardiovaskule [6]. Kejadian hipertensi di dunia mencapai angka 1,13 milyar orang, dimana 31% diantaranya beresiko terhadap orang dewasa dan terus meningkat sebesar 5,1% dibandingkan dengan prevalensi global pada tahun 2000-2010 [7]. Dilansir dari data Kementerian Kesehatan, hipertensi terjadi pada kelompok usia 31-44 tahun dengan persentase 31,6%, pada umur 45-54 tahun sebesar 45,3% dan 55-64 tahun sebesar 55,2% [8].

Perkembangan teknologi terus meningkat khususnya dibidang kesehatan [9]. Salah satu teknologi dalam dunia kesehatan diantaranya adalah pemanfaatan *machine learning* pada berbagai aplikasi keperawatan dan medis namun tidak terbatas pada resiko kardiovaskular dan identifikasi penyakit jantung [10]. *Machine learning* ialah sebuah aplikasi komputer dan algoritma matematika yang memanfaatkan data sebagai media pembelajaran dan menghasilkan prediksi di masa yang akan datang [11]. Prediksi dini sebuah penyakit menjadi masalah penting dalam bidang kesehatan dan *machine learning* hadir menjadi sebuah solusi untuk mengatasi masalah tersebut [12]. Dalam prediksi diagnostik *machine learning* telah dibuktikan menjadi metode yang menjanjikan [13]. Metode ini dibuat dengan serangkaian alat dan teknik yang dapat mengeksplorasi kumpulan data berjumlah besar, dan membantu menemukan pola-pola tertentu yang bermanfaat bagi ilmu pengetahuan [14].

Beberapa penelitian terdahulu yang memanfaatkan *machine learning* telah dilakukan oleh Santoni [12] dimana pengolahan data menggunakan alat yang disebut dengan KNIME. Algoritma *machine learning* yaitu *artificial neural network* berhasil untuk memprediksi penyakit hipertensi dengan nilai akurasi data hingga 94.7%. Penelitian yang telah dilakukan oleh Chamidah [15], memanfaatkan algoritma *machine learning* dengan membandingkan model yang menggunakan dan tanpa *oversampling*. Hasil penelitian tersebut menghasilkan nilai akurasi dalam prediksi penyakit hipertensi sebesar 91%. Penelitian yang telah dilakukan oleh Lase [16] memanfaatkan algoritma *machine learning* yaitu LVQ pada klasifikasi jenis hipertensi dengan pedoman ESH dan menghasilkan nilai ketepatan akurasi sistem prediksi sebesar 94,67%. Penelitian yang telah dilakukan oleh Rios [17] memanfaatkan *machine learning* dengan mengoptimalkan fitur yang diekstraksi dari sinyal PPG melalui transformasi hamburan *wavelet* dan klinis seperti usia, indeks masa tubuh dan detak jantung dalam memprediksi normotensi dan prehipertensi dan menghasilkan nilai akurasi sebesar 71,42%.

Berdasarkan penelitian terdahulu kami secara khusus berkontribusi pada penggunaan salah satu model *machine learning* klasifikasi yaitu *random forest*. Penelitian-penelitian terdahulu sudah menggunakan berbagai jenis teknik klasifikasi pada data hipertensi namun kami belum melihat adanya optimasi dengan memanfaatkan *tuning hyperparameter*. Kami tidak hanya mencari nilai akurasi dengan cara-cara yang sudah dilakukan sebelumnya, namun lebih banyak melakukan optimasi sehingga nilai akurasi dari model klasifikasi yang dibuat bisa menjadi hasil yang lebih baik.

2. METODE PENELITIAN

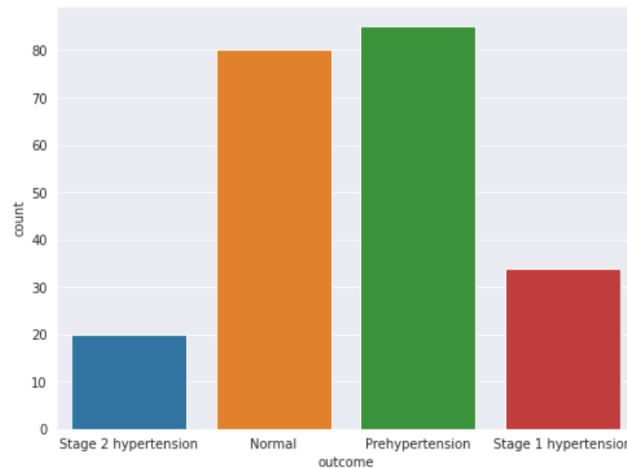
2.1. Dataset

Penelitian ini menggunakan dataset dari PPG-BP (*photoplethysmography-blood pressure*) yang dapat digunakan untuk mengetahui kinerja dari model klasifikasi *machine learning* [18]. Dataset ini memiliki beberapa fitur yang digunakan dalam prediksi hipertensi antara lain adalah jenis kelamin, usia, tinggi (cm), berat badan (kg), tekanan darah sistolik (mmHg), tekanan darah diastolik (mmHg), denyut jantung (bpm), dan BMI (kg/m). Kelas utama hipertensi yang terdapat pada dataset tersebut ialah normal (sehat), prehipertensi, hipertensi stadium-1, dan hipertensi stadium-2. Dataset PPG-BP dikumpulkan dari 219 orang dewasa dengan rentang umur

21-86 tahun. Jumlah masing-masing fitur pada dataset ini dapat dilihat pada Tabel 1 dan grafik perbandingan jumlahnya dapat dilihat pada Gambar 1.

Tabel 1. Jumlah Data Fitur

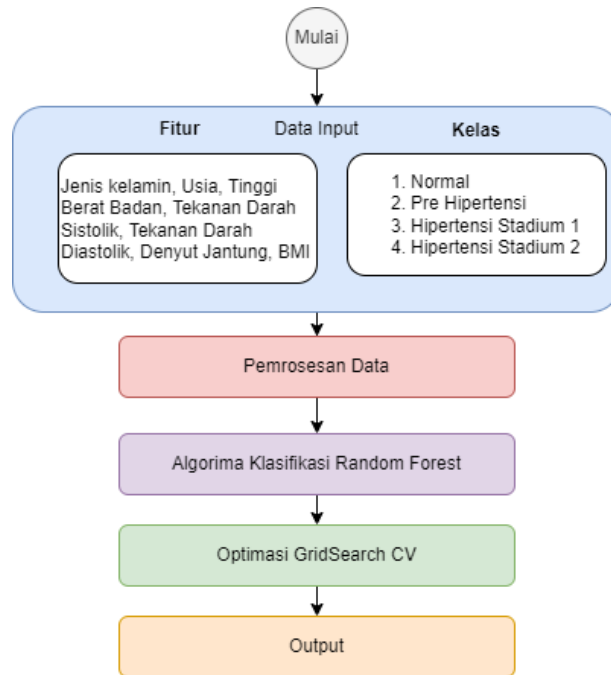
Fitur	Jumlah
Prehypertension	85
Normal	80
Stage 1 hypertension	34
Stage 2 hypertension	20



Gambar 1. Perbandingan Jumlah Fitur Dataset Hipertensi

2.2. Metode

Pemanfaatan salah satu algoritma machine learning digunakan dalam mengklasifikasi jenis hipertensi berdasarkan data pribadi. Pembelajaran mesin dengan teknik diawasi (*supervised learning*) digunakan dalam penelitian ini. Pembelajaran mesin diawasi yaitu melatih mesin dengan memberikan label, artinya beberapa data sudah diberikan data yang benar [19]. Sebagai data inputan pada pembelajaran mesin ini, kami menggunakan delapan fitur dari data pribadi yang telah dikumpulkan dalam bentuk *dataset*. Algoritma yang digunakan ialah *random forest*. Algoritma ini sering digunakan pada teknik-teknik klasifikasi pada pembelajaran mesin [20]. Klasifikasi pada hipertensi ini bersifat *multiclass* karena memiliki empat jenis pada kategori penyakitnya dan *random forest* sangat cocok untuk menangani hal ini [2]. Diagram blok yang diusulkan dalam penelitian ini dapat dilihat pada Gambar 2.



Gambar 2. Diagram Blok Metode Penelitian

Berdasarkan Gambar 2, ada beberapa tahapan penelitian yang dilakukan. Penelitian dimulai dengan memberikan data input yaitu fitur-fitur dan kelas pada *dataset* hipertensi. *Dataset* tersebut tentunya sudah melewati tahap pemrosesan data (*pre-processing*). Tahap ini memberikan efek perbaikan jika terdapat data yang bersifat *nullable* atau kosong yang kemudian dapat kita hapus atau ganti dengan data rata-rata (*mean*) jika data berupa angka [21]. Data-data fitur yang bersifat string selanjutnya diubah ke dalam bentuk numerik agar algoritma dapat memprosesnya dengan baik [22]. Sebagai contoh kelas “normal” diubah menjadi nilai 1, prehipertensi menjadi 2, hipertensi stadium-1 menjadi 3 dan hipertensi stadium-2 menjadi 4. Hal tersebut dibuat dengan teknik *label encoding* yang mampu merubah data string menjadi data numerik. Data yang sudah bersifat numerik kemudian dilakukan penskalaan fitur yang akan menyatukan tiap variabel mandiri dalam format data baru [23]. Teknik normalisasi juga digunakan untuk menormalkan kolom fitur yang dimiliki pada kisaran [0,1] dengan metode *min-max scaling*. Selanjutnya pemilihan fitur memanfaatkan *principal component analysis* (PCA) untuk menangani masalah *overfitting* akibat redundansi data. PCA berfungsi untuk mengurangi jumlah fitur dari kumpulan data dengan mempertahankan variabel sebanyak mungkin [22]. Data-data yang sudah melewati pemrosesan data selanjutnya dibagi menjadi data latih dan data uji dengan perbandingan 70% berbanding 30% [24]. Data yang sudah matang selanjutnya diterapkan pada model klasifikasi *machine learning* yaitu *random forest*. Hasil pembelajaran ini kemudian di optimasi parameternya hingga menghasilkan parameter yang terbaik dalam pembelajarannya. Hasil optimasi di uji coba kembali dengan teknik evaluasi yang umum digunakan pada model klasifikasi *machine learning*.

2.3. Random Forest

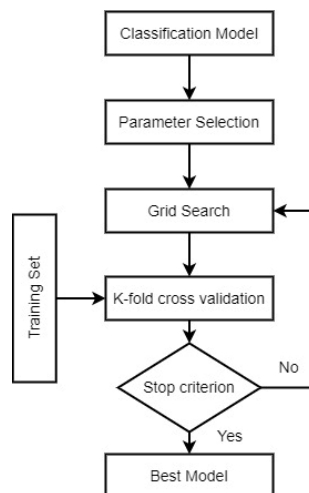
Random forest adalah sebuah algoritma *machine learning* tipe *ensemble* pohon berstruktur untuk klasifikasi dan regresi yang pertama kali diusulkan oleh Breiman pada tahun 2001 [20]. Karena struktur yang sederhana dan memiliki kinerja luar biasa di dibandingkan algoritma *machine learning* lainnya, sehingga banyak digunakan di berbagai bidang. Algoritma ini merupakan kombinasi pohon keputusan dengan dua program prosedur pemilihan acak yang dapat mengurangi varians dan meninggalkan bias secara tetap. Prosedur acak pertama adalah

agregasi *bootstrap*, juga dikenal sebagai *bagging*, yang menghasilkan sejumlah set pelatihan baru dengan pengambilan sampel secara acak. Pergantian dari set pelatihan asli dan penggabungan hasil pohon dipasang oleh set pelatihan baru yang diperoleh di atas. Hal tersebut merupakan prosedur acak kedua dimana pemilihan subruang acak akan memilih subset fitur secara acak di setiap simpul pohon. Alasan untuk melakukan ini adalah untuk mengurangi korelasi pohon di *random forest*.

Terdapat tiga aspek utama dalam menjalankan algoritma *random forest* diantaranya adalah [22] (1) Melakukan *bootstrap* sampling untuk membangun pohon prediksi. (2) Setiap pohon keputusan memprediksi dengan prediktor acak. (3) *Random forest* melakukan prediksi dengan mengombinasikan hasil dari setiap pohon keputusan dengan cara majority vote untuk melakukan klasifikasi.

2.4. GridSearch CV

Peningkatan nilai akurasi dapat dioptimalkan dengan menggunakan *cross-validation gridsearch* [25]. Metode ini adalah pemilihan kombinasi model dan *hyperparameter* dengan menguji setiap kombinasi dan memvalidasi setiap kombinasi. Tujuan dari *grid search* adalah untuk mengidentifikasi kombinasi yang menghasilkan performansi model terbaik yang dapat dipilih untuk digunakan sebagai model prediktif [26]. Kontinuitas *grid search* sering dikombinasikan dengan validasi silang *k-fold*, yang menghasilkan indeks untuk model klasifikasi [31]. *K-fold cross-validation* dapat mengulang data latih dan uji hingga k iterasi dan $1/k$ divisi dari *dataset* yang digunakan sebagai data uji [27]. Keakuratan model k dapat diperoleh, dan kinerja klasifikasi validasi silang *k-fold* dievaluasi berdasarkan akurasi rata-rata model k . Kemudian, parameter *classifier* diubah berdasarkan pencarian *grid* dan akurasi *classifier* dihitung ulang. Proses optimasi berlanjut, seperti yang ditunjukkan pada Gambar 3. Akurasi model klasifikasi dengan semua kombinasi parameter dibandingkan untuk menghasilkan nilai akurasi yang maksimal.



Gambar 3. K-Fold Cross-Validation dan Grid search

2.5. Evaluasi Model

Evaluasi model *random forest* yang digunakan adalah *confusion matrix*. Kriteria klasifikasi dengan menggunakan *confusion matrix* dapat dilihat pada Tabel 2.

Tabel 2. *Confusion Matrix*

Kelas	Diklasifikasikan sebagai Positif	Diklasifikasikan sebagai Negatif
+	True Positive (TP)	False Positive (FP)
-	True Negative (TN)	False Negative (FN)

Tabel 2 adalah klasifikasi dari *confusion matrix*. *True positive* (TP) menunjukkan bahwa model klasifikasi dengan benar memberi label jumlah tupel positif. *True Negative* (TN) menunjukkan bahwa model klasifikasi dengan benar memberi label jumlah tupel negatif. *False positive* (FP) menunjukkan bahwa model klasifikasi memberi label yang salah untuk jumlah tupel negatif. *False Negative* (FN) menunjukkan bahwa model klasifikasi memberi label yang salah untuk jumlah tupel positif. Akurasi adalah ukuran kinerja model klasifikasi dan merupakan persentase jumlah data yang diprediksi dengan benar dari total data [28]. Persamaan untuk menghitung ketelitian dapat dilihat pada Persamaan 1.

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

Pengukuran *performance metrics* terdiri dari presisi, *recall* dan *f1-score* [27]. Presisi adalah rasio positif atau derajat keandalan, yaitu proporsi prediksi berlabel positif yang benar terhadap prediksi positif keseluruhan [29]. Persamaan untuk menghitung presisi dapat dilihat pada Persamaan 2.

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

Recall juga dikenal sebagai *true positive rate* atau sensitivitas. *Recall* juga disebut sebagai derajat keandalan model dalam mendeteksi data berlabel positif dengan benar [29]. Persamaan untuk menghitung *recall* dapat dilihat pada persamaan 3.

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

F1-score merangkum semua hasil perhitungan *precision* dan *recall* dengan membuat rata-rata harmonik [27]. Persamaan untuk menghitung *F1-score* dapat dilihat pada persamaan 4.

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

3. HASIL DAN PEMBAHASAN

3.1. Transformasi Data

Dataset telah mengalami proses transformasi data hingga layak untuk diproses dengan model klasifikasi. Dataset diubah menjadi versi numerik sehingga model klasifikasi *random forest* dapat memprosesnya. Pada field kategori kelas *prehypertension* diubah menjadi angka 0, *stage 1 hypertension* menjadi 1, *stage 2 hypertension* menjadi 2 dan normal menjadi 3. Pada field *sex* juga diubah nilainya menjadi format *number* sehingga jika *male* menjadi 0 dan *female* menjadi 1. Pada field *bmi* angka desimal juga diubah menjadi bilangan bulat sebagai contoh angka 32.66 diformat menjadi 136 dengan *label encoder*. Adapun beberapa perubahan transformasi data dapat dilihat pada Tabel 3.

Tabel 3. Hasil Transformasi Dataset

index	subject	sex	age	height	weight	sbp	dbp	hr	bmi
0	1	2	0	45	152	63	161	89	97
1	2	3	0	50	157	50	160	93	76
2	3	6	0	47	150	47	101	71	79
3	4	8	1	45	172	65	136	93	87
4	5	9	0	46	155	65	123	73	73
...	...	...	...	...	...	...	...	...	...

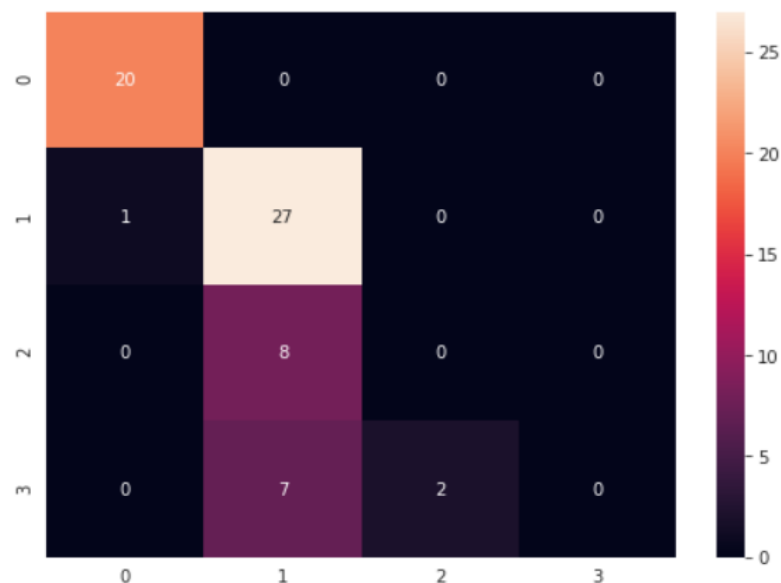
3.2. Performa Model Klasifikasi Random Forest

Modul seleksi memisahkan kelompok data menjadi data pelatihan dan pengujian. Komposisi perbandingan *dataset* data latih dan data uji yaitu 70% berbanding 30%. Penerapan algoritma klasifikasi *random forest* kemudian diuji pada *dataset* yang telah diolah dan dipisahkan sebelumnya. Tabel 4 merupakan hasil performa pengujian (*testing*) dari 219 data uji klasifikasi algoritma *random forest* dengan jumlah *split k-fold* yaitu $n=10$.

Tabel 4. Hasil Klasifikasi *Random Forest*

Kelas	Precision	Recall	F1-Score	Support
Normal	0.95	1.00	0.98	20
Prehypertension	0.64	0.96	0.77	28
Stage 1 hypertension	0.00	0.00	0.00	8
Stage 2 hypertension	0.00	0.00	0.00	9
accuracy	72,3%			65

Performa model klasifikasi ini juga menghasilkan *confusion matrix*. Blok pada posisi diagonal merupakan nilai TP dan TN. Matrik ini dapat dilihat pada gambar 4. Pada posisi diagonal kita bisa melihat angka 20,27,0 dan 0. Matrik ini bisa kita gunakan sebagai perhitungan nilai akurasi sesuai dengan persamaan 1. Untuk mempermudah perhitungan ini, hasil *confusion matrik* disajikan pada Tabel 5.



Gambar 4. Hasil *Confusion Matrix*Tabel 5. Hasil *Confusion Matrix Testing*

0	1	2	3	JUMLAH DATA
20	0	0	0	20
1	27	0	0	28
0	8	0	0	8
0	7	2	0	9
Akurasi : 72,3%				65

Berdasarkan Tabel 5, kita dapat menghitung nilai akurasi pengujian dengan persamaan 1 dimana nilai secara diagonal yang dijumlahkan dibagi total data uji sehingga $(20 + 27 + 0 + 0) / 65 = 72,3\%$. Tabel *confusion matrix* model SVM ini dapat dijelaskan secara detail sebagai berikut:

1. 20 data *test* untuk kelas *Normal*, sistem tersebut memprediksi 20 *Normal*, 0 *Prehypertension*, 0 *Stage 1 hypertension*, dan 0 *Stage 2 hypertension*.
2. 28 data *test* untuk kelas *Prehypertension*, sistem tersebut memprediksi 1 *Normal*, 27 *Prehypertension*, 0 *Stage 1 hypertension* dan 0 *Stage 2 hypertension*.
3. 8 data *test* untuk kelas *Stage 1 hypertension*, sistem tersebut memprediksi 0 *Normal*, 7 *Prehypertension*, 0 *Stage 1 hypertension* dan 0 *Stage 2 hypertension*.
4. 9 data *test* untuk kelas *Stage 2 hypertension*, sistem tersebut memprediksi 0 *Normal*, 7 *Prehypertension*, 2 *Stage 1 hypertension* dan 0 *Stage 2 hypertension*.

3.3. Optimasi *GridSearch CV*

Berdasarkan Tabel 4 dan Tabel 5, akurasi yang dihasilkan oleh model klasifikasi *random forest* adalah 72,3%. Nilai akurasi ini masih bisa dioptimalkan untuk menghasilkan nilai yang lebih baik. Kami memanfaatkan optimasi *gridsearch cv* untuk mendapatkan parameter terbaik. Hasil dari optimasi ini kami menemukan parameter sebagai berikut.

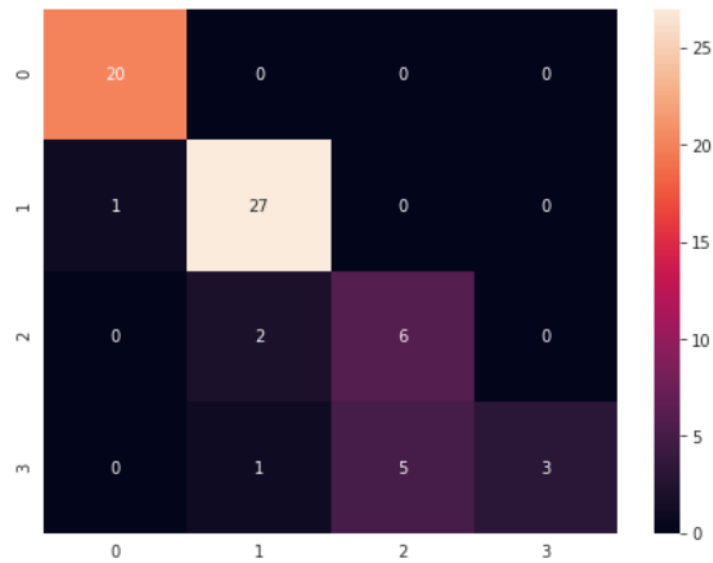
```
RandomForestClassifier(max_depth=80, max_features=3,
min_samples_leaf=3,min_samples_split=8, n_estimators=1000)
```

Selanjutnya kami mengaplikasikan kembali parameter tersebut untuk diujikan pada dataset sebelumnya. Tabel 6 merupakan hasil performa pengujian (*testing*) dari 219 data uji klasifikasi algoritma *random forest* setelah menggunakan parameter terbaik dari *optimasi gridsearch*.

Tabel 6. Hasil Klasifikasi *Random Forest* Setelah Optimasi

Kelas	Precision	Recall	F1-Score	Support
Normal	0.95	1.00	0.98	20
Prehypertension	0.90	0.96	0.93	28
Stage 1 hypertension	0.55	0.75	0.63	8
Stage 2 hypertension	1.00	0.33	0.50	9
accuracy	86%			65

Berdasarkan perbandingan antara Tabel 4 dan Tabel 6, terdapat pengaruh yang signifikan ketika menggunakan metode optimasi *grid search cv*. Hasil akurasi meningkat 13,7%, dengan akurasi baru 86%. Sedangkan untuk hasil *confusion matrix* dapat dilihat pada Gambar 4 dan Tabel 7.

Gambar 4. Hasil *Confusion Matrix* Setelah OptimasiTabel 7. Hasil *Confusion Matrix* Setelah Optimasi

0	1	2	3	JUMLAH DATA
20	0	0	0	20
1	27	0	0	28
0	2	6	0	8
0	1	5	3	9
Akurasi : 86,1%				65

Berdasarkan Tabel 5, kita dapat menghitung nilai akurasi pengujian dengan persamaan 1 dimana nilai secara diagonal yang dijumlahkan dibagi total data uji sehingga $(20 + 27 + 6 + 3) / 65 = 86,1\%$. Tabel *confusion matrix* model *random forest* ini dapat dijelaskan secara detail sebagai berikut:

1. 20 data *test* untuk kelas *Normal*, sistem tersebut memprediksi 20 *Normal*, 0 *Prehypertension*, 0 *Stage 1 hypertension*, dan 0 *Stage 2 hypertension*.
2. 28 data *test* untuk kelas *Prehypertension*, sistem tersebut memprediksi 1 *Normal*, 27 *Prehypertension*, 0 *Stage 1 hypertension* dan 0 *Stage 2 hypertension*.
3. 8 data *test* untuk kelas *Stage 1 hypertension*, sistem tersebut memprediksi 0 *Normal*, 2 *Prehypertension*, 6 *Stage 1 hypertension* dan 0 *Stage 2 hypertension*.
4. 9 data *test* untuk kelas *Stage 2 hypertension*, sistem tersebut memprediksi 0 *Normal*, 1 *Prehypertension*, 5 *Stage 1 hypertension* dan 3 *Stage 2 hypertension*.

4. KESIMPULAN

Berdasarkan hasil penelitian, ditemukan bahwa model klasifikasi *random forest* dapat digunakan sebagai alat klasifikasi untuk prediksi jenis penyakit hipertensi. Pengolahan data diperlukan dalam penelitian ini karena model *random forest* membutuhkan transformasi *dataset* dalam format numerik. Hasil akurasi sebelum optimasi adalah 72,3%. Optimasi dilakukan dengan melakukan iterasi dengan validasi silang terhadap model yang digunakan. Implementasi *gridsearch cv* menghasilkan parameter terbaik untuk model *random forest* yaitu *max_depth=80*, *max_features=3*, *min_samples_leaf=3*, *min_samples_split=8*, *n_estimators=1000*. Parameter ini

kemudian diterapkan kembali pada tahap klasifikasi jenis penyakit hipertensi. Hasilnya adalah peningkatan nilai akurasi hingga 86,1%. Optimasi *gridsearch cv* ternyata memberikan dampak yang signifikan terhadap klasifikasi *randomforest*, dimana terjadi peningkatan akurasi sebesar 13,7% antara sebelum dan sesudah menggunakan teknik ini. Penelitian selanjutnya adalah mencoba mengimplementasikannya dalam algoritma klasifikasi lain dan menganalisis hasilnya.

DAFTAR PUSTAKA

- [1] D. Masruroh, E. M. M.Has, and R. Fauziningtyas, "Pengaruh Terapi Humor dengan Media Film Komedi terhadap Penurunan Tekanan Darah Pada Lansia Dengan Hipertensi," *Indones. J. Community Heal. Nurs.*, vol. 4, no. 1, p. 29, 2019, doi: 10.20473/ijchn.v4i1.12496.
- [2] M. Nour and K. Polat, "Automatic Classification of Hypertension Types Based on Personal Features by Machine Learning Algorithms," *Math. Probl. Eng.*, vol. 2020, pp. 1–14, 2020, doi: 10.1155/2020/2742781.
- [3] A. C. Telaumbanua and Y. Rahayu, "Penyuluhan Dan Edukasi Tentang Penyakit Hipertensi," *J. Abdimas Saintika*, vol. 3, no. 1, p. 119, 2017, doi: 10.30633/jas.v3i1.1069.
- [4] Y. Nursakinah and A. Handayani, "Faktor-Faktor Risiko Hipertensi Diastolik Pada Usia Dewasa Muda," *J. Pandu Husada*, vol. 2, no. 1, p. 21, 2021, doi: 10.30596/jph.v2i1.5426.
- [5] G. Mancia *et al.*, "2018 ESC/ESH Guidelines for the management of arterial hypertension," *Eur. Heart J.*, vol. 39, pp. 3021–3104, 2018, doi: doi:10.1093/eurheartj/ehy339.
- [6] D. Tryastuti, "Determinan Pre-Hipertensi Di Kelurahan Curug Kecamatan Cimanggis Kota Depok," *Indones. J. Heal. Sci.*, vol. 11, no. 1, p. 71, 2019, doi: 10.32528/ijhs.v11i1.2240.
- [7] Y. T. G. Arum, "Hipertensi pada Penduduk Usia Produktif (15-64 Tahun)," *Higeia J. Public Heal. Res. Dev.*, vol. 3, no. 3, pp. 84–94, 2019, doi: <https://doi.org/10.15294/higeia.v3i3.30235>.
- [8] A. Syntya, "Hypertension and heart disease: literature review," *J. Ilm. Permas J. Ilm. STIKES Kendal*, vol. 11, no. 4, pp. 541–550, 2021, doi: doi.org/10.32583/pskm.v11i4.1621.
- [9] F. D. Telaumbanua, P. Hulu, T. Z. Nadeak, R. R. Lumbantong, and A. Dharma, "Penggunaan Machine Learning Di Bidang Kesehatan," *J. Teknol. dan Ilmu Komput.*, vol. 3, no. 1, pp. 57–64, 2019, doi: <https://doi.org/10.34012/jutikomp.v2i2.657>.
- [10] A. Mustafa and M. Rahimi Azghadi, "Automated machine learning for healthcare and clinical notes analysis," *Computers*, vol. 10, no. 2, pp. 1–31, 2021, doi: 10.3390/computers10020024.
- [11] A. Roihan, P. Abas Sunarya, and A. S. Rafika, "Pemanfaatan Machine Learning dalam Berbagai Bidang: Review paper," *IJCIT (Indonesian J. Comput. Inf. Technol.)*, vol. 5, no. 1, pp. 75–82, 2020, doi: <https://doi.org/10.31294/ijcit.v5i1.7951>.
- [12] M. M. Santoni, N. Chamidah, and N. Matondang, "Prediksi Hipertensi menggunakan Decision Tree, Naïve Bayes dan Artificial Neural Network pada software KNIME," *Techno.Com*, vol. 19, no. 4, pp. 353–363, 2020, doi: 10.33633/tc.v19i4.3872.
- [13] W. Apriliah, I. Kurniawan, M. Baydhowi, and T. Haryati, "Prediksi Kemungkinan Diabetes pada Tahap Awal Menggunakan Algoritma Klasifikasi Random Forest," *Sistemasi*, vol. 10, no. 1, p. 163, 2021, doi: 10.32520/stmsi.v10i1.1129.
- [14] L. B. Moreira and A. A. Namen, "A hybrid data mining model for diagnosis of patients with clinical suspicion of dementia," *Comput. Methods Programs Biomed.*, vol. 165, pp. 139–149, 2018, doi: 10.1016/j.cmpb.2018.08.016.
- [15] N. Chamidah, M. Mega Santoni, and N. Matondang, "Pengaruh Oversampling pada Klasifikasi Hipertensi dengan Algoritma Naïve Bayes, Decision Tree, dan Artificial Neural Network (ANN)," *J. RESTI (Rekayasa Sist. dan Teknol. Inf.)*, vol. 4, no. 4, pp. 635–641, 2020, doi: <https://doi.org/10.29207/resti.v4i4.2015>.
- [16] A. Akbar Ritonga, Ibnu Rasyid Munthe, Masrizal, "LVQ Algorithm for The Classification

- of Hypertension Based on ESH Guideline,” *J. Mantik*, vol. 4, no. 3, pp. 1772–1778, 2020, doi: <https://doi.org/10.35335/mantik.Vol4.2020.1006.pp1772-1778>.
- [17] E. Martinez-Ríos, L. Montesinos, and M. Alfaro-Ponce, “A machine learning approach for hypertension detection based on photoplethysmography and clinical data,” *Comput. Biol. Med.*, vol. 145, no. March, p. 105479, 2022, doi: 10.1016/j.compbimed.2022.105479.
- [18] “PPG Blood Pressure,” *IEEE Dataport*, 2019. <https://ieee-dataport.org/open-access/ppg-blood-pressure>.
- [19] A. Yakimovich, A. Beaugnon, Y. Huang, and E. Ozkirimli, “Labels in a haystack: Approaches beyond supervised learning in biomedical applications,” *Patterns*, vol. 2, no. 12, pp. 1–11, 2021, doi: 10.1016/j.patter.2021.100383.
- [20] J. Zheng, Y. Liu, and Z. Ge, “Dynamic ensemble selection based improved random forests for fault classification in industrial processes,” *IFAC J. Syst. Control*, vol. 20, p. 100189, 2022, doi: 10.1016/j.ifacsc.2022.100189.
- [21] A. Toha, P. Purwono, W. Gata, and A. Toha, “Model Prediksi Kualitas Udara dengan Support Vector Machines dengan Optimasi Hyperparameter GridSearch CV,” *Bul. Ilm. Sarj. Tek. Elektro*, vol. 4, no. 1, pp. 12–21, 2022, doi: 10.12928/biste.v4i1.6079.
- [22] P. Purwono, A. Wirasto, and K. Nisa, “Comparison of Machine Learning Algorithms for Classification of Drug Groups,” *Sisfotenika*, vol. 11, no. 2, p. 196, 2021, doi: 10.30700/jst.v11i2.1134.
- [23] T. Emmanuel, T. Maupong, D. Mpoeleng, T. Semong, B. Mphago, and O. Tabona, “A survey on missing data in machine learning,” *J. Big Data*, vol. 8, no. 1, 2021, doi: 10.1186/s40537-021-00516-9.
- [24] P. Purwono, A. Ma’arif, I. S. Mangku Negara, W. Rahmawati, and J. Rahmawan, “Linkage Detection of Features that Cause Stroke using Fuzzy Machine Learning Model,” *J. Ilm. Tek. Elektro Komput. dan Inform.*, vol. 7, no. 3, p. 423, 2021, doi: 10.26555/jiteki.v7i3.22237.
- [25] T. Yan, S. L. Shen, A. Zhou, and X.-S. Chen, “Prediction of geological characteristics from shield operational parameters using integrating grid search and K-fold cross validation into stacking classification algorithm,” *J. Rock Mech. Geotech. Eng.*, p. 100310, 2022, doi: <https://doi.org/10.1016/j.jrmge.2022.03.002>.
- [26] G. S. K. Ranjan, A. Kumar Verma, and S. Radhika, “K-Nearest Neighbors and Grid Search CV Based Real Time Fault Monitoring System for Industries,” in *2019 IEEE 5th International Conference for Convergence in Technology, I2CT 2019*, 2019, no. March, doi: 10.1109/I2CT45611.2019.9033691.
- [27] X. Xiong, S. Hu, D. Sun, S. Hao, H. Li, and G. Lin, “Detection of false data injection attack in power information physical system based on SVM–GAB algorithm,” *Energy Reports*, vol. 8, pp. 1156–1164, 2022, doi: 10.1016/j.egy.2022.02.290.
- [28] S. Katoch, V. Singh, and U. S. Tiwary, “Indian Sign Language Recognition System using SURF with SVM and CNN,” *Array*, p. 100141, 2022, doi: <https://doi.org/10.1016/j.array.2022.100141>.
- [29] A. Luque, A. Carrasco, A. Martín, and A. de las Heras, “The impact of class imbalance in classification performance metrics based on the binary confusion matrix,” *Pattern Recognit.*, vol. 91, pp. 216–231, 2019, doi: 10.1016/j.patcog.2019.02.023.