

Optimasi Prediksi Gagal Jantung dengan Teknik *Ensemble Bagging* Pada Neural Network

Angga Ariawan

Ilmu Komputer, Universitas Media Nusantara Citra, Jakarta Barat, Indonesia

E-mail: angga.ariawan@mncu.ac.id

Abstract

Prediction of heart failure is an important step in the early management of serious cardiovascular disease. This research uses the Ensemble bagging algorithm with Neural Network. The dataset is taken from Heart Failure Clinical Records available in the UCI Machine Learning Repository. The use of training data in this research was 80% of the total data set, and 20% of the test data. The dataset is divided into two feature models, namely features with categorical data and continuous data. At the data transformation stage, continuous data is subjected to value scaling. several single classifier machine learning algorithms have been tested in this research such as Logistic Regression, Artificial Neural Networks (ANN), Naïve Bayes, SVM. Single classifier Artificial Neural Networks (ANN) produces an accuracy value of 82%. Ensemble learning using the bagging method on Artificial Neural Networks (ANN) was carried out to get a higher accuracy value. Ensemble learning using the bagging method on Artificial Neural Networks (ANN) obtained an accuracy value of 98%. This method is proven to have increased the accuracy value by 16% better than just using a Single Classifier Artificial Neural Networks (ANN) in the case of the Heart Failure Clinical Records dataset.

Keywords: Ensemble, Bagging, Neural Network, Heart Failure, Machine Learning

Abstrak

Prediksi kegagalan jantung merupakan langkah penting dalam penanganan dini penyakit kardiovaskular yang serius. Penelitian ini menggunakan Algoritme Ensemble bagging dengan Neural Network. Dataset diambil dari Heart Failure Clinical Records yang tersedia di UCI Machine Learning Repository. Penggunaan data training pada penelitian ini adalah sebanyak 80% pada jumlah data set, dan 20 % pada data uji. Pada dataset tersebut dibagi dua model fitur yaitu fitur dengan data katagori dan data kontinu. Pada tahap transformasi data, data kontinu dilakukan penskalaan nilai. beberapa algoritme pembelajaran mesin single classifier telah diujikan pada penelitian ini seperti Logistic Regression, Artificial Neural Networks (ANN), Naïve Bayes, SVM. Single clasifier Artificial Neural Networks (ANN) menghasilkan nilai akurasi 82%. Ensemble learning dengan metode bagging pada Artificial Neural Networks (ANN) dilakukan untuk mendapatkan nilai akurasi yang lebih tinggi, pada Ensemble learning dengan metode bagging Artificial Neural Networks (ANN) didapatkan nilai Akurasi sebesar 98%. Methode ini terbukti telah meningkatkan nilai akurasi sebanyak 16% lebih baik dibandingkan hanya menggunakan Single Classifier Artificial Neural Networks (ANN) pada kasus dataset Heart Failure Clinical Records.

Kata Kunci: Ensemble, Bagging, Neural Network, Gagal Jantung, Machine Learning

1. Pendahuluan

Jantung merupakan organ tubuh yang sangat vital pada manusia. Menurut statistik dunia, terdapat 9,4 juta kematian setiap tahun yang disebabkan oleh penyakit kardiovaskular, dan 45% kematian tersebut disebabkan oleh penyakit jantung koroner, dengan angka yang terus meningkat setiap tahunnya[1]. Penyakit jantung koroner adalah gangguan fungsi jantung akibat penyempitan pada pembuluh darah koroner atau arteri koroner, yang dapat dicegah dengan pola hidup sehat[2]. Selain itu, beberapa pemicu lainnya seperti tekanan darah tinggi yang terus menerus, diabetes,

usia, kadar natrium dalam darah, dan kadar kolesterol juga berkontribusi terhadap penyakit jantung koroner[2], [3]. Pendeteksian dini penyakit jantung koroner sangat diperlukan guna mengurangi risiko kematian mendadak dan untuk pengelolaan penyakit yang lebih efektif. Dengan mengidentifikasi masalah jantung pada tahap awal, pasien dapat mendapatkan terapi yang tepat dan menerapkan perubahan pola hidup serta gaya hidup sehat seperti diet yang seimbang, olahraga teratur, dan menghindari kebiasaan buruk seperti merokok. Selain itu, deteksi dini penyakit jantung koroner dapat mencegah kerusakan jantung lebih lanjut akibat penundaan diagnosis[3].

Beberapa penelitian sebelumnya menggunakan data medis pernah diusulkan. Penelitian yang dilakukan oleh Fakhurrazi, Studi terkait menggunakan teknik *ensemble* seperti *Bagging*, menunjukkan peningkatan kinerja dalam memprediksi hasil klinis. Eksperimen yang dilakukan pada tiga dataset klinis menunjukkan bahwa model yang menggunakan *Bagging neural networks* secara konsisten mencapai tingkat akurasi yang lebih tinggi daripada jaringan saraf tradisional dengan nilai akurasi sebesar 96%.

Selain itu azizul hakim,dkk [4]melakukan penelitian dengan menggunakan data medis untuk memprediksi *stroke*, teknik *Bagging* dimodifikasi dan diterapkan pada prediksi *stroke*. Metode ini, yang menggabungkan pemangkasan berbasis akurasi, menunjukkan peningkatan signifikan dalam akurasi prediksi, mengungguli teknik *bagging* tradisional, memiliki nilai akurasi sebesar 96%[4]. Studi terbaru juga mengevaluasi model prediksi penyakit jantung dengan menggunakan algoritme *Machine Learning*, seperti *Logistic Regression*, *Random Forest Classifier*, dan *k-NN*. Penelitian ini menunjukkan bahwa model yang menggunakan teknik *ensemble*, khususnya *Bagging*, memiliki kinerja yang lebih baik dalam memprediksi penyakit jantung daripada model individual [5], [6]. Penelitian dengan model data non medis juga pernah diusulkan dengan menggunakan metode *Bagging* dalam pengembangan model hibrida untuk indurasi pelet bijih besi. Dengan memadukan teknik *ensemble* dengan pemodelan kinetik, penelitian ini menunjukkan peningkatan akurasi dan kinerja dibandingkan dengan model kinetik tradisional[7]. Penelitian yang dilakukan oleh pebrianti,dkk tentang penggunaan *ensamble learning* dapat meningkatkan nilai akurasi pernah diusulkan dengan model data non medis hasilnya bisa meningkatkan nilai akurasi sebesar 13,3% dibanding menggunakan *single classifier*[8]. Berdasarkan analisis dari penelitian sebelumnya, ditemukan bahwa berbagai peneliti telah mengusulkan teknik yang berbeda untuk memprediksi data-data klinis. Dapat disimpulkan bahwa *algoritme ensemble* dengan teknik *Bagging neural network* merupakan suatu metode yang memiliki akurasi cukup tinggi. Oleh karena itu, penelitian ini akan mengembangkan model yang sesuai untuk prediksi gagal jantung dengan akurasi yang lebih tinggi menggunakan *ensemble learning*.

Penelitian ini bertujuan untuk menjawab pertanyaan-pertanyaan berikut: Bagaimana teknik *Bagging* dengan *neural network* dapat meningkatkan akurasi prediksi penyakit jantung koroner? Apakah metode ini lebih efektif dibandingkan dengan metode prediksi lainnya? Manfaat dari penelitian ini diharapkan dapat memberikan kontribusi signifikan dalam bidang teknologi informasi, sistem informasi, rekayasa perangkat lunak, multimedia, dan ilmu komputer, khususnya dalam pengembangan model prediksi penyakit yang lebih akurat dan efisien.

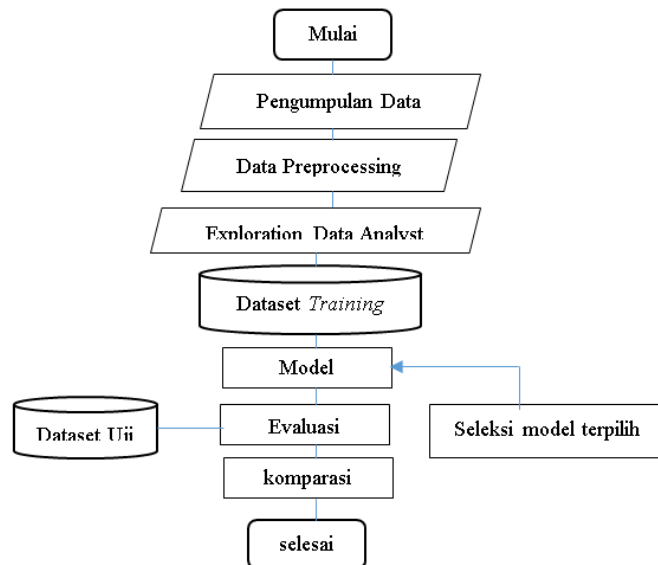
2. Metodologi Penelitian

Berikut adalah deskripsi singkat dari setiap tahapan dalam diagram alur proses:

- 1) Pengumpulan Data: Tahap ini melibatkan pengumpulan dataset yang akan digunakan dalam analisis dan pemodelan. Data dapat dikumpulkan dari *Heart Failure Clinical Records* yang tersedia di *UCI Machine Learning Repository*.
- 2) Data Preprocessing: Pada tahap ini, data yang telah dikumpulkan dibersihkan dan dipersiapkan untuk analisis lebih lanjut. Ini melibatkan penanganan nilai yang hilang, normalisasi data, penghapusan *outlier*, dan transformasi fitur.
- 3) Exploration Data: Tahap eksplorasi data melibatkan analisis data secara mendalam untuk memahami distribusi, hubungan antar fitur, dan pola yang mungkin ada dalam data. Ini sering dilakukan dengan visualisasi data dan teknik statistik deskriptif.

- 4) **Dataset Training:** Data yang telah diproses dan dieksplorasi kemudian dibagi menjadi dataset pelatihan. Dataset ini digunakan untuk melatih model *machine learning*.
- 5) **Model:** Pada tahap ini, berbagai algoritme *machine learning* diterapkan pada dataset pelatihan untuk membuat model prediktif.
- 6) **Evaluation:** Model yang telah dibuat dievaluasi menggunakan dataset uji. Kinerja model diukur berdasarkan akurasinya.
- 7) **Komparasi:** Hasil evaluasi dari berbagai model dibandingkan untuk menentukan model yang memiliki kinerja terbaik.
- 8) **Seleksi model terpilih dengan nilai akurasi tinggi :** Model dengan kinerja terbaik dipilih berdasarkan hasil evaluasi dan perbandingan.

Berikut Gambar 1 Merupakan tahapan yang dilakukan dalam penelitian ini:



Gambar 1. Tahapan penelitian

3. Hasil dan Pembahasan

Berikut adalah penjelasan setiap tahapan yang digunakan dalam penelitian ini

3.1. Pengumpulan Data

Dataset ini diambil dari *Heart Failure Clinical Records* yang tersedia di *UCI Machine Learning Repository*. Dataset ini mengandung 299 data pasien dengan 13 fitur yang relevan untuk menganalisis kondisi kesehatan pasien yang menderita gagal jantung. Berikut adalah penjelasan tentang setiap fitur yang ada di dalam dataset:

1. *Age*: Usia pasien (dalam tahun).
2. *Anaemia*: Mengindikasikan apakah pasien mengalami anemia (0 = Tidak, 1 = Ya).
3. *High Blood Pressure*: Mengindikasikan apakah pasien memiliki tekanan darah tinggi (0 = Tidak, 1 = Ya).
4. *Creatinine Phosphokinase (CPK)*: Tingkat enzim CPK dalam darah (mcg/L).
5. *Diabetes*: Mengindikasikan apakah pasien menderita diabetes (0 = Tidak, 1 = Ya).
6. *Ejection Fraction*: Persentase darah yang keluar dari jantung pada setiap kontraksi (%).
7. *Platelets*: Jumlah platelet dalam darah (kiloplatelet/mL).
8. *Serum Creatinine*: Tingkat kreatinin serum dalam darah (mg/dL).
9. *Serum Sodium*: Tingkat natrium serum dalam darah (mEq/L).
10. *Sex*: Jenis kelamin pasien (0 = Wanita, 1 = Pria).
11. *Smoking*: Mengindikasikan apakah pasien merokok (0 = Tidak, 1 = Ya).
12. *Time*: Waktu (dalam hari) setelah follow-up ketika kejadian terjadi.

13. **DEATH_EVENT**: Mengindikasikan apakah pasien meninggal selama periode follow-up (0 = Tidak, 1 = Ya).

Dataset ini umumnya digunakan untuk memprediksi potensi kematian pasien berdasarkan fitur-fitur yang ada, mengidentifikasi faktor risiko utama yang berkontribusi terhadap gagal jantung, serta membangun model prediksi yang dapat membantu dokter dalam membuat keputusan klinis. Contoh penggunaannya meliputi analisis statistik deskriptif yang menyediakan gambaran umum tentang distribusi fitur-fitur dalam dataset, seperti *mean*, *median*, dan distribusi frekuensi. Selain itu, visualisasi data melalui grafik dan plot membantu memahami hubungan antara berbagai fitur dan outcome (**DEATH_EVENT**). Pemodelan prediktif menggunakan teknik pembelajaran mesin seperti regresi logistik, svm, *neural networks*, *naïve bayes* dan *ensemble learning* yang nanti akan digunakan untuk memprediksi kematian pasien. Analisis faktor risiko juga dilakukan untuk mengidentifikasi fitur mana yang memiliki pengaruh paling signifikan terhadap kematian akibat gagal jantung melalui analisis regresi atau metode statistik lainnya. Dataset ini sangat berguna dalam bidang kesehatan untuk meningkatkan penanganan pasien dengan gagal jantung dan menyediakan wawasan yang lebih baik bagi dokter dan peneliti.

3.2. Data Preprocessing

Dari dataset yang memiliki 13 fitur tersebut, fitur yang akan dijadikan kelas prediksi adalah **DEATH_EVENT**, sementara yang lainnya akan digunakan sebagai fitur input. Untuk mempersiapkan dataset ini sebelum digunakan dalam analisis atau pemodelan, beberapa langkah preprocessing perlu dilakukan.

Langkah pertama adalah melakukan eksplorasi data. Ini meliputi:

- Shape* data: Mengecek dimensi dataset untuk memahami jumlah baris dan kolom yang ada.
- Cek apakah data tersebut mengandung unsur NAN: Memastikan tidak ada nilai yang hilang atau null dalam dataset yang dapat mengganggu analisis.
- Cek mana data yang termasuk katagorikal dan mana data yang termasuk *numerikal*: Mengidentifikasi jenis data untuk menentukan metode preprocessing yang tepat.

Jika data numerikal memiliki rentang nilai yang terlalu besar, langkah selanjutnya adalah melakukan *scaling* atau transformasi data agar nilainya berada dalam rentang 0-1 atau dalam bilangan yang lebih mudah diproses. Hal ini bertujuan untuk memudahkan komputasi.

Beberapa teknik *preprocessing* yang dapat diterapkan meliputi:

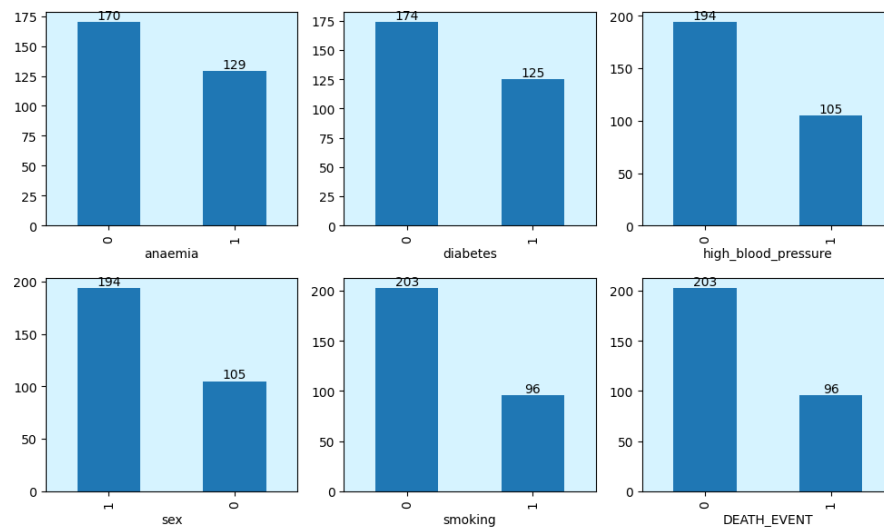
- Data Cleaning*: Tahap pertama adalah membersihkan data mentah dengan menghapus data yang tidak lengkap, tidak relevan, dan tidak akurat. Ini penting untuk menghindari kesalahan saat menganalisis data, dalam case ini dilakukan pencarian outlier.
- Transformasi/ *Data: Scaling* data untuk fitur yang bersifat kontinu dilakukan pada tahap transformasi data dalam proses preprocessing. Setelah membersihkan data dari nilai yang tidak lengkap, tidak relevan, atau tidak akurat pada tahap data *cleaning*, serta memastikan format yang konsisten pada tahap data integration.

Langkah berikutnya adalah transformasi data, Pada tahap ini, kita mengidentifikasi fitur-fitur kontinu dan kategorikal. Fitur kontinu, seperti usia, kadar enzim *creatinine phosphokinase*, *ejection fraction*, jumlah platelet, kadar serum *creatinine*, kadar *serum* sodium, dan waktu, perlu *discalling* agar nilai-nilainya berada dalam rentang yang lebih kecil, seperti 0-1, untuk memudahkan komputasi. Teknik yang umum digunakan untuk scaling termasuk *Min-Max Scaling*, yang mengubah nilai data kontinu ke dalam rentang [0, 1], dan *standardization*, yang mengubah nilai data kontinu sehingga memiliki mean 0 dan standar deviasi 1. Selain itu, fitur kategorikal juga dapat diubah menjadi bentuk numerikal jika diperlukan, misalnya dengan teknik *one-hot encoding*. Setelah transformasi data, langkah terakhir adalah mengurangi jumlah data dengan teknik seperti dimensionality reduction, numerosity reduction, dan data compression untuk memastikan dataset siap digunakan dalam analisis atau pemodelan prediksi. Dengan mengikuti langkah-langkah ini, dataset akan lebih mudah diproses dan dianalisis, memberikan hasil yang lebih akurat dan efisien.

3.3. Exploratory Data Analyst (EDA)

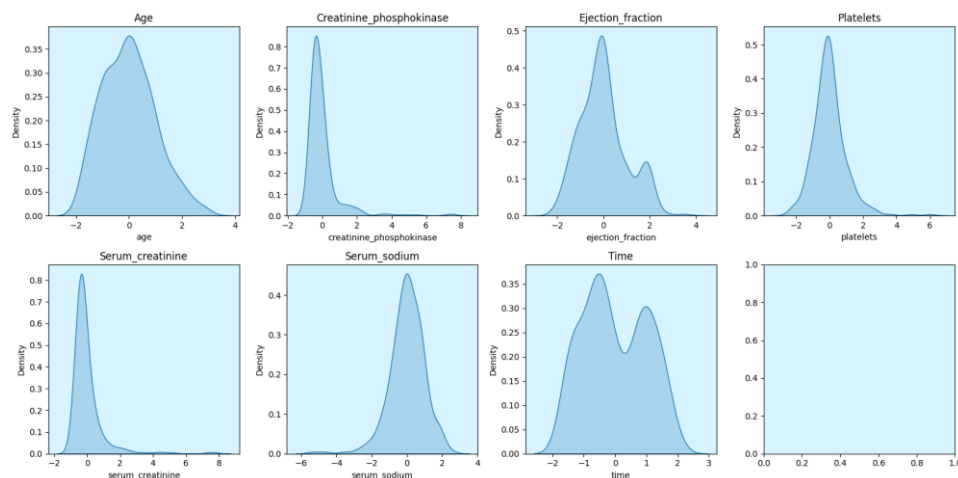
Setelah melakukan proses preprocessing pada dataset Heart Failure Clinical Records, langkah selanjutnya adalah melakukan *Exploratory Data Analysis (EDA)* untuk memahami karakteristik data yang telah diproses.

- a. **Distribusi data Kategorikal:** Langkah pertama adalah memeriksa distribusi fitur kategorikal dalam dataset. Fitur kategorikal dalam dataset ini termasuk anaemia, diabetes, high_blood_pressure, sex, dan smoking. Distribusi data kategorikal bisa diperiksa dengan metode tabulasi frekuensi atau visualisasi seperti bar plot. Ini membantu dalam memahami proporsi masing-masing kategori dalam fitur tersebut, misalnya, berapa banyak pasien yang menderita anemia atau berapa banyak pasien yang merokok. Distribusi data katagorikal ditunjukkan pada Gambar 2.



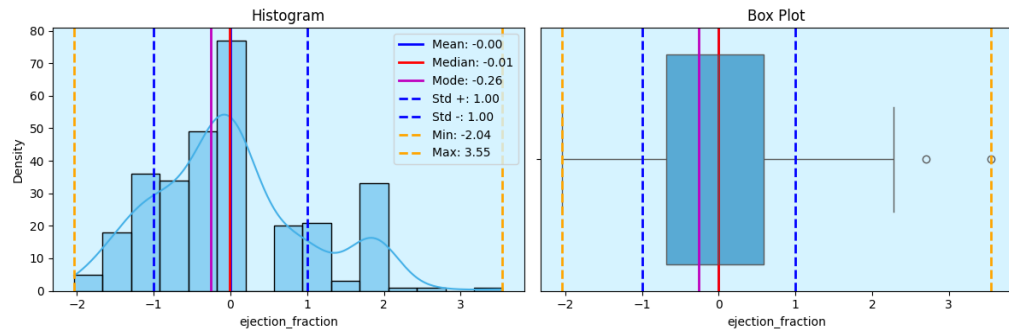
Gambar 2. Distribusi data katagorikal

- b. **Density Dataset:** Memeriksa kepadatan data (*density*) dari fitur-fitur numerikal untuk memahami distribusi nilai mereka. Density plot bisa digunakan untuk visualisasi ini. Density plot akan menunjukkan bagaimana nilai-nilai fitur seperti usia, kadar enzim creatinine phosphokinase, ejection fraction, jumlah platelet, kadar serum creatinine, dan waktu terdistribusi dalam dataset. Ini membantu dalam mengidentifikasi pola umum, seperti apakah data tersebut berdistribusi normal atau memiliki skewness tertentu. Density dataset skewness ditunjukkan oleh Gambar 3.



Gambar 3. Grafik density data kontinu

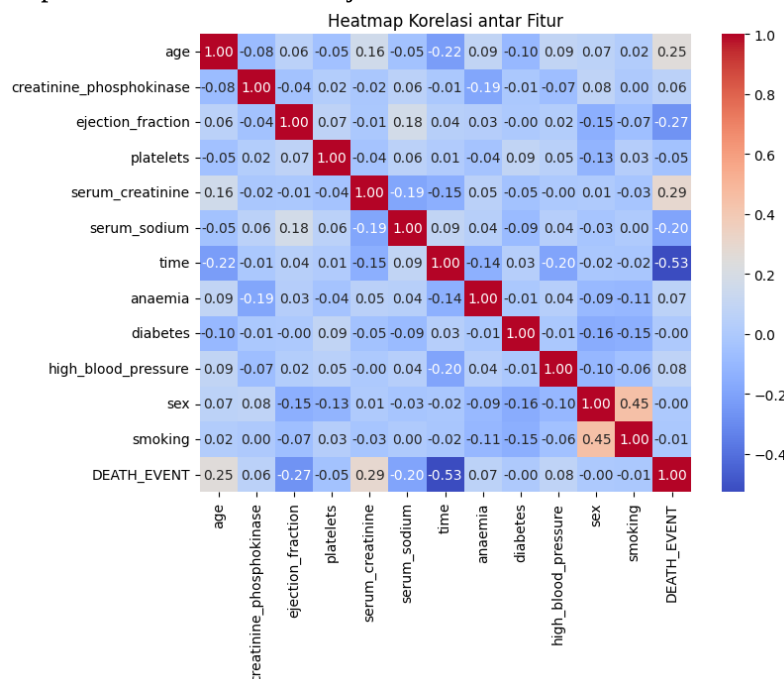
- c. **Handling Outlier.** Outlier bisa mempengaruhi hasil analisis dan pemodelan, sehingga perlu memeriksa dan menentukan apakah outlier perlu ditangani. Outlier dapat diidentifikasi menggunakan *box plot*, yang akan menyoroti nilai-nilai ekstrim yang jauh dari kuartil bawah dan atas. Setelah mengidentifikasi outlier, perlu diputuskan apakah outlier tersebut harus dihapus, ditransformasi, atau dipertahankan. Keputusan ini didasarkan pada pengaruh outlier terhadap analisis dan hasil prediksi. Gambar sampling Outlier fitur *ejection fraction* ditunjukkan oleh Gambar 4



Gambar 4. Outlier fitur *ejection fraction*

- d. **Korelasi Data:** Memeriksa korelasi antara fitur-fitur dalam dataset untuk memahami hubungan antara fitur. Korelasi membantu dalam memahami sejauh mana perubahan dalam satu fitur mungkin mempengaruhi fitur lain. Matriks korelasi dapat digunakan untuk menunjukkan nilai korelasi antara setiap pasangan fitur, yang kemudian divisualisasikan menggunakan *heatmap*. Ini membantu mengidentifikasi fitur yang memiliki hubungan kuat satu sama lain.

Gambar Heatmap korelasi antar fitur ditunjukkan oleh Gambar 5.



Gambar 5. Grafik *heatmap* korelasi antar fitur

Dengan melakukan langkah-langkah EDA ini setelah *preprocessing*, diperoleh wawasan yang lebih mendalam tentang data yang telah diproses, identifikasi hubungan penting antara fitur, dan persiapan data untuk tahap pemodelan lebih lanjut dengan pemahaman yang lebih baik tentang karakteristik *dataset*.

3.4. Dataset Training dan dataset Uji

Setelah melakukan langkah-langkah EDA ini dan memperoleh wawasan yang lebih mendalam tentang data yang telah diproses serta mengidentifikasi hubungan penting antara fitur, dataset siap untuk tahap pemodelan lebih lanjut. Dataset tersebut kemudian dibagi menjadi dua bagian 80% dari dataset digunakan sebagai data *training*, yang akan digunakan untuk melatih model, dan 20% sisanya digunakan sebagai data uji, yang akan digunakan untuk mengevaluasi kinerja model. Data uji diambil dari dataset yang terpisah dari data *training* untuk memastikan bahwa model yang dikembangkan dapat diuji dengan data yang belum pernah dilihat selama proses pelatihan. Ini penting untuk memberikan evaluasi yang lebih objektif dan realistis terhadap performa model, karena menguji model dengan data yang sama yang digunakan untuk melatihnya dapat menyebabkan *overfitting*, di mana model tampil sangat baik pada data *training* tetapi kurang mampu menggeneralisasi pada data baru atau tak terlihat. Pembagian ini membantu memastikan bahwa model tidak hanya menghafal data *training*, tetapi juga belajar pola yang dapat diaplikasikan pada data lain.

3.5. Model, Evaluasi dan Komparasi

Berikut adalah beberapa algoritme yang akan dilakukan pengujian

a. Super Vector Machine

Dalam proses pengujian model, *Support Vector Machine* (SVM) digunakan dengan berbagai jenis kernel untuk menentukan kernel mana yang memberikan akurasi terbaik. Kernel adalah fungsi yang digunakan oleh SVM untuk memproyeksikan data ke dalam dimensi yang lebih tinggi, sehingga memudahkan dalam pemisahan kelas. Pada percobaan ini, empat jenis kernel digunakan: linear, radial basis function (rbf), polynomial (poly), dan sigmoid. Hasil pengujian menunjukkan bahwa kernel linear memberikan akurasi tertinggi sebesar 80%, sedangkan kernel rbf dan poly masing-masing memberikan akurasi sebesar 73%, dan kernel sigmoid memberikan akurasi sebesar 77%. Hasil komparasi akurasi dari berbagai kernel dengan parameter $c = 100$ yang digunakan ditunjukkan oleh Tabel 1.

Tabel 1. Hasil Evaluasi Akurasi SVM dengan Berbagai Kernel

No	Kernel	Nilai Akurasi
1	Linear	80%
2	RBF	73%
3	Polynomial	73%
4	Sigmoid	77%

b. Artificial Neural Network (Multilayer Perceptron)

Dalam pengujian model, digunakan *Multi-Layer Perceptron* (MLP) classifier dengan berbagai konfigurasi *hidden layer* untuk menentukan arsitektur mana yang memberikan akurasi terbaik. MLP adalah jenis jaringan saraf tiruan yang terdiri dari beberapa lapisan (*layer*) yang masing-masing berisi sejumlah *node* (*neuron*). Setiap lapisan *neuron* memproses *input* dari lapisan sebelumnya dan mengirimkan outputnya ke lapisan berikutnya. Pada percobaan ini, lima konfigurasi *hidden layer* diuji:

1. (4,5,6): Artinya, model memiliki tiga lapisan tersembunyi (*hidden layers*). Lapisan pertama memiliki 4 *neuron*, lapisan kedua memiliki 5 *neuron*, dan lapisan ketiga memiliki 6 *neuron*.
2. (4,7,9): Model memiliki tiga lapisan tersembunyi dengan 4 *neuron* di lapisan pertama, 7 *neuron* di lapisan kedua, dan 9 *neuron* di lapisan ketiga.
3. (3,7,9,10): Model memiliki empat lapisan tersembunyi dengan 3 *neuron* di lapisan pertama, 7 *neuron* di lapisan kedua, 9 *neuron* di lapisan ketiga, dan 10 *neuron* di lapisan keempat.
4. (2,4,6): Model memiliki tiga lapisan tersembunyi dengan 2 *neuron* di lapisan pertama, 4 *neuron* di lapisan kedua, dan 6 *neuron* di lapisan ketiga.
5. (8,8,8): Model memiliki tiga lapisan tersembunyi, masing-masing dengan 8 *neuron*.

Hasil pengujian menunjukkan bahwa konfigurasi (4,7,9) memberikan akurasi tertinggi sebesar 82%.

Berikut adalah tabel yang merangkum hasil komparasi akurasi dari berbagai konfigurasi *hidden layer* yang digunakan ditunjukkan oleh Tabel 2.

Tabel 2. Hasil Evaluasi Akurasi *Neural Network* dengan *Multilayer Perceptron*

No.	Jumlah Layer dan Node	Nilai Akurasi
1	(4,5,6)	78%
2	(4,7,9)	82%
3	(3,7,9,10)	80%
4	(2,4,6)	76%
5	(8,8,8)	69%

c. *Logistic Regression*

Dalam pengujian model selanjutnya, kita melakukan evaluasi menggunakan *logistic regression* dengan iterasi sebanyak 1000 kali. Penggunaan parameter “*max_iter=1000*” menunjukkan bahwa algoritma optimasi akan melakukan penyesuaian parameter model hingga 1000 iterasi, dengan tujuan meminimalkan fungsi kerugian. Hasil dari pengujian ini menunjukkan bahwa model *logistic regression* mencapai akurasi sebesar 80%. Ini menunjukkan bahwa model mampu mengklasifikasikan data dengan baik. Hasil ini menunjukkan efektivitas dari peningkatan jumlah iterasi dalam mencapai konvergensi yang lebih baik, sehingga model dapat mempelajari pola dalam data secara lebih baik. Namun masih ada beberapa algoritme yang harus di Uji coba untuk mencari nilai akurasi yang lebih tinggi.

d. *Naïve bayes*

Selanjutnya, dilakukan evaluasi dengan *Naive Bayes Gaussian* yang menghasilkan akurasi sebesar 70%. Meskipun akurasi ini lebih rendah daripada *logistic regression*, namun penilaian terhadap berbagai algoritme tersebut membuka kesempatan untuk mencari solusi yang lebih baik dan memastikan bahwa model yang dipilih memiliki kinerja optimal untuk tugas klasifikasi pada *dataset* yang digunakan.

e. *Ensamble Learning*

Dalam langkah berikutnya, akan disajikan hasil evaluasi dari beberapa algoritme klasifikasi yang telah diuji sebelumnya. Evaluasi ini mencakup akurasi masing-masing model dalam melakukan prediksi pada *dataset* yang digunakan. Dengan menyajikan hasil evaluasi ini, maka dibuat keputusan yang lebih baik dalam pemilihan algoritma untuk pengujian selanjutnya. Tabel 3 menunjukkan komparasi hasil *Single Classifier* yang telah diujikan.

Tabel 3 komparasi hasil pengujian *single Classifier*

No	Algoritme	Konfigurasi	Nilai Akurasi
1	SVM	Linear	80%
2	SVM	RBF	73%
3	SVM	Polynomial	73%
4	SVM	Sigmoid	77%
5	MLP	(4,5,6)	78%
6	MLP	(4,7,9)	82%
7	MLP	(3,7,9,10)	80%
8	MLP	(2,4,6)	76%
9	MLP	(8,8,8)	69%
10	Logistic Regression	1000 iterasi	80%
11	Naïve Bayes	Gaussian NB	70%

Setelah mengevaluasi beberapa algoritma klasifikasi, diputuskan untuk menggunakan teknik *ensemble Neural Network Bagging*. Keputusan ini didasarkan pada observasi bahwa dalam pengujian sebelumnya, *Neural Network* menunjukkan kinerja yang unggul dibandingkan dengan algoritma lainnya. Dengan demikian, dengan menerapkan *ensemble Neural Network Bagging*,

didapatkan keunggulan *Neural Network* untuk meningkatkan kinerja prediksi secara signifikan. Bagging adalah teknik *ensemble* yang menggabungkan hasil prediksi dari beberapa model *Neural Network* yang berbeda yang dilatih pada subset acak dari data *training*. Melalui proses ini, variasi antar model dapat dikurangi, meningkatkan stabilitas dan akurasi keseluruhan dari model *ensemble*. Dengan menggunakan teknik *ensemble Neural Network Bagging*, tujuannya adalah menghasilkan model klasifikasi yang lebih kuat dan dapat diandalkan untuk tugas prediksi pada dataset yang digunakan. Dalam konteks dataset kegagalan jantung, teknik *Bagging* diterapkan untuk melatih model *neural network*. Setelah evaluasi dilakukan, hasil yang diperoleh menunjukkan tingkat akurasi sebesar 98%. Berikut adalah langkah-langkah yang diambil dalam penerapan teknik Bagging dengan *Neural Network*.

1. Pemuatan dan Persiapan *Dataset*:

Misalkan dataset asli adalah Q , dengan jumlah sampel N .

2. Pembuatan *Dataset Bootstrap*:

- Buat *BS subset bootstrap* Q_n dari dataset Q .
- Setiap *subset bootstrap* Q_n dibuat dengan pengambilan sampel acak dengan penggantian dari Q , sehingga setiap subset memiliki jumlah sampel yang sama dengan Q .
- Setiap *subset bootstrap* Q_n dibuat dengan pengambilan sampel acak dengan penggantian dari Q

$$Q_n = \{x_i | x_i \in Q, i = 1, 2, \dots, N\} \quad (1)$$

3. Inisialisasi Model Dasar

- Siapkan *BS model neural network* M_n . Untuk setiap *subset bootstrap* Q_n .
- Setiap model M_n direpresentasikan sebagai model *Neural Network* yang akan dilatih dengan *dataset bootstrap* Q_n .

4. Pelatihan Model Dasar :

- Latih setiap model M_n dengan *dataset bootstrap* Q_n .
 $M_n \leftarrow \text{Train}(Q_n) \text{ untuk } n = 1, 2, \dots, BS$ (2)
- Proses pelatihan dilakukan dengan memasukkan *dataset bootstrap* Q_n ke dalam model *Neural Network* M_n .
- Hasil dari proses pelatihan ini adalah model *Neural Network* yang terlatih M_n

5. Prediksi menggunakan *bagging*

a. Prediksi Model Dasar

- Setelah semua model dasar M_n dilatih, gunakan model-model ini untuk melakukan prediksi pada dataset pengujian X_{test} .
- Output dari setiap model dasar M_n adalah prediksi kelas $\hat{y}_{n,j}$ untuk setiap sampel data uji x_j . di mana $j = 1, 2, \dots, N_{test}$ (jumlah sampel data uji). $\hat{y}_{n,j} = M_n(x_j)$

b. Penggabungan Prediksi:

Hitung prediksi akhir untuk setiap sampel data uji x_j

4. Kesimpulan

Dalam konteks penggunaan teknik *Bagging* untuk melatih model *neural network* pada dataset kegagalan jantung, evaluasi menunjukkan tingkat akurasi yang tinggi, mencapai 98%. Langkah-langkah yang diambil dalam penerapan teknik ini mencakup pembuatan dataset *bootstrap*, inisialisasi model dasar, pelatihan model dasar, penggabungan output model dasar, dan penentuan hasil akhir melalui metode *voting*.

- Hasil akhir dari teknik *Bagging* ini menunjukkan bahwa kelas yang paling sering diprediksi oleh sebagian besar model dasar menjadi hasil akhir dari proses klasifikasi. Dengan demikian, dapat disimpulkan bahwa teknik Bagging Tradisional efektif digunakan untuk meningkatkan akurasi prediksi pada dataset kegagalan jantung.
- Pengujian Model pada Data yang Berbeda: Meskipun tingkat akurasi yang tinggi diperoleh pada dataset yang digunakan, penting untuk menguji model pada dataset yang berbeda untuk menilai kestabilan dan generalisasi dari model yang dilatih.

3. Penggunaan Metrik Evaluasi yang Lebih Komprehensif: Selain akurasi, metrik evaluasi lain seperti presisi, recall, F1-score, dan area di bawah kurva ROC (AUC-ROC) dapat memberikan pemahaman yang lebih holistik tentang kinerja model.

Daftar Pustaka

- [1] L. Ghani *Et Al.*, 'Faktor Risiko Dominan Penyakit Jantung Koroner Di Indonesia Dominant Risk Factors Of Coronary Heart Disease In Indonesia'.
- [2] A. Dwi Erawati, 'Peningkatan Pengetahuan Tentang Penyakit Jantung Koroner', *Jurnal Abdimas-Hip*, Vol. 2, 2021.
- [3] Mamat Supriyono, 'Tesis Faktor-Faktor Risiko Yang Berpengaruh Terhadap Kejadian Penyakit Jantung Koroner Pada Kelompok Usia < 45 Tahun Program Pasca Sarjana-Magister Epidemiologi Universitas Diponegoro Semarang Tahun 2008'.
- [4] D. Das, Md. I. Abu Bakar, Rajshahi University. Faculty Of Engineering, Institute Of Electrical And Electronics Engineers. Bangladesh Section., And Institute Of Electrical And Electronics Engineers, *5th International Conference On Computer, Communication, Chemical, Materials And Electronic Engineering : Ic⁴me² : 11-12 July, 2019*.
- [5] A. Mehmood *Et Al.*, 'Prediction Of Heart Disease Using Deep Convolutional Neural Networks', *Arab J Sci Eng*, Vol. 46, No. 4, Pp. 3409–3422, Apr. 2021, Doi: 10.1007/S13369-020-05105-1.
- [6] M. Ward, K. Malmsten, H. Salamy, And C. H. Min, 'Data Balanced Bagging Ensemble Of Convolutional-Lstm Neural Networks For Time Series Data Classification With An Imbalanced Dataset', In *Proceedings - Ieee International Symposium On Circuits And Systems*, Institute Of Electrical And Electronics Engineers Inc., 2021. Doi: 10.1109/Iscas51556.2021.9401389.
- [7] Dongbei Gong Xue Yuan, Ieee Control Systems Society, And Ieee Singapore Section. Industrial Electronics Chapter, *The 26th Chinese Control And Decision Conference (2014 Ccdc) : 31 May - 2 June 2014, Changsha, China*.
- [8] D. Pebrianti, 'Technical Job Distribution At Bsd Sharp Service Center Using Combination Of Naïve Bayes And K-Nearest Neighbour'.