

# MODEL DETEKSI GANGGUAN TIDUR BERDASARKAN KUALITAS TIDUR DAN GAYA HIDUP MENGGUNAKAN MACHINE LEARNING

Sevira Destiany Putri  
Informatics engineering  
Faculty of Computer  
Sriwijaya University  
South Sumatera, Indonesia  
[Seviraaputri@gmail.com](mailto:Seviraaputri@gmail.com)

Annisa Darmawahyuni  
Intelligent System Research Grup  
Faculty of Computr  
Sriwijaya University  
South Sumatera  
[riset.annisadarmawahyuni@gmail.com](mailto:riset.annisadarmawahyuni@gmail.com)

**Abstract**—Penelitian ini berfokus pada pengembangan model klasifikasi untuk gangguan tidur dengan menggunakan Machine Learning melalui analisis data kesehatan serta pola hidup individu. Dataset yang digunakan berasal dari Kaggle dan berisi informasi dari berbagai responden dengan fitur seperti lama tidur, tingkat stres, tekanan darah, detak jantung, aktivitas fisik, serta indeks massa tubuh. Proses pra-pemrosesan meliputi pembersihan data, normalisasi dengan Min-Max Scaling, dan pemisahan data dengan proporsi 80% untuk pelatihan, 10% untuk validasi, dan 10% untuk pengujian. Dalam studi ini, lima algoritma Machine Learning diterapkan, yaitu Random Forest, Decision Tree, K-Nearest Neighbors, Support Vector Machine, dan Naïve Bayes.. Evaluasi hasil menunjukkan bahwa algoritma Random Forest dan Decision Tree memberikan kinerja terbaik, dengan akurasi masing-masing 95%, diikuti K-Nearest Neighbors yang mencapai akurasi 94%, sementara Support Vector Machine dan Naïve Bayes memiliki akurasi yang lebih rendah, yakni 86% dan 85%. Temuan ini mengindikasikan bahwa algoritma berbasis Decision Tree lebih efektif dalam mengklasifikasikan gangguan tidur. Dengan demikian, model ini memiliki potensi untuk dikembangkan lebih lanjut sebagai sistem yang mendukung diagnosis awal gangguan tidur menggunakan Machine Learning yang efisien dan akurat.

**Keywords**—Gangguan Tidur, Kualitas Tidur, Prediksi, Machine Learning, Kesehatan

## I. PENDAHULUAN

Tidur merupakan kebutuhan fundamental yang sangat signifikan baik dari segi fisiologis maupun psikologis. Proses tidur memiliki peran yang sangat penting dalam pemulihan energi, peningkatan fungsi kognitif, serta pemeliharaan kesehatan fisik dan mental secara keseluruhan. Manusia menghabiskan sekitar sepertiga dari hidupnya untuk tidur, menjadikannya kebutuhan yang tidak kalah penting dengan makanan dan minuman [1]. Kualitas tidur yang baik mendukung kesiapan individu dalam menjalani aktivitas sehari-hari serta mempertahankan keseimbangan emosi.

Namun, banyak orang di seluruh dunia menghadapi berbagai masalah tidur, mulai dari insomnia, sleep apnea, hingga gangguan tidur lainnya yang berakibat pada penurunan kualitas tidur secara drastis. Gangguan tidur tidak hanya menimbulkan rasa lelah dan kurang konsentrasi, tetapi juga meningkatkan kemungkinan terserang penyakit kronis seperti hipertensi, diabetes, dan gangguan mental [2], [4]. Setiap tahunnya, sekitar 20% hingga 50% laporan menunjukkan

bahwa 17% orang dewasa mengalami masalah gangguan tidur[3].

Faktor-faktor yang memengaruhi kualitas tidur sangat beragam, meliputi kondisi medis, psikologis, dan lingkungan. Pola hidup modern seperti ketergantungan pada gadget, paparan cahaya biru di malam hari, pola makan yang tidak sehat, serta stres yang berlebihan juga berkontribusi pada semakin meluasnya gangguan tidur [5]. Ini menunjukkan perlunya metode diagnosis dan deteksi gangguan tidur yang lebih cepat, efektif, dan terjangkau dibandingkan dengan metode tradisional yang sering kali memakan waktu dan biaya tinggi [5], [6], [9].

Seiring dengan perkembangan teknologi, penggunaan metode Machine Learning memberikan solusi yang menarik untuk memprediksi kualitas tidur dan mendeteksi gangguan yang mungkin terjadi secara otomatis. Algoritma Machine Learning memiliki kemampuan untuk menganalisis data besar dari berbagai parameter fisiologis dan kebiasaan sehari-hari, yang memungkinkan pengambilan prediksi yang akurat dan mendukung diagnosa yang lebih efisien. Metode ini dapat mempercepat proses identifikasi bagi pasien yang memerlukan perawatan medis lebih lanjut.

Tujuan penelitian ini adalah untuk menciptakan model prediksi gangguan tidur berdasarkan Machine Learning dengan menggunakan data kesehatan pribadi, seperti durasi tidur, aktivitas fisik, tingkat stres, tekanan darah, dan indeks massa tubuh. Diharapkan bahwa temuan penelitian dapat berkontribusi pada pengembangan teknologi kesehatan digital yang dapat melakukan deteksi awal gangguan tidur secara tepat dan praktis, serta memberikan informasi yang berguna bagi tenaga medis dan masyarakat dalam upaya meningkatkan kualitas tidur dan kesehatan secara keseluruhan.

## II. KAJIAN PUSTAKA

Tidur adalah kondisi yang tidak sadar dimana seseorang dapat dibangun oleh stimulus atau sensoris yang sesuai atau juga dapat dikatakan sebagai keadaan tidak sadarkan diri yang relatif. Tidur yang tidak berkualitas dapat menyebabkan gangguan kesehatan seperti kelelahan kronis, penurunan fungsi kognitif, bahkan meningkatkan risiko penyakit kronis seperti hipertensi dan depresi [4], [7]. Tidur sendiri adalah kebutuhan biologis dasar yang memiliki fungsi restoratif, membantu proses metabolisme, serta mengatur sistem kekebalan tubuh [5].

Gangguan tidur adalah kondisi di mana seseorang mengalami kesulitan untuk memulai, mempertahankan, atau mendapatkan tidur yang berkualitas. Masalah ini mencakup insomnia, sleep apnea, narcolepsy, dan restless leg syndrome. Di Indonesia, jumlah kasus gangguan tidur terus meningkat, seringkali tanpa disadari oleh mereka yang mengalaminya [2], [3]. Selain mempengaruhi orang dewasa, gangguan tidur juga dapat terjadi pada anak-anak dan berdampak pada perkembangan mereka [4].

Berbagai faktor dapat memengaruhi kualitas tidur, termasuk gaya hidup, keadaan mental, stres, penggunaan gadget, konsumsi kafein, serta kondisi lingkungan seperti pencahayaan dan suhu ruangan [5]. Pola tidur yang buruk bisa membuat seseorang merasa lelah secara terus-menerus, sulit berkonsentrasi, bahkan mengalami depresi [8].

Salah satu solusi yang banyak dikembangkan untuk mengatasi gangguan tidur adalah penerapan deteksi dini berbasis Machine Learning. Pendekatan ini memungkinkan identifikasi gangguan tidur dilakukan secara cepat dan akurat melalui analisis data kesehatan dan gaya hidup individu. Beberapa penelitian menunjukkan bahwa metode berbasis Machine Learning mampu membantu proses diagnosis awal gangguan tidur secara efisien serta mendukung pengambilan keputusan medis berbasis data [9], [11], [12].

Perkembangan teknologi, khususnya dalam bidang Machine Learning (ML), telah membuka peluang untuk mengembangkan sistem deteksi gangguan tidur yang dapat membantu dalam proses diagnosis gangguan tidur (*sleep disorder*). Sistem ini memanfaatkan berbagai data seperti kebiasaan tidur, detak jantung, aktivitas fisik, serta faktor lingkungan untuk menghasilkan prediksi yang lebih akurat. Dalam penelitian ini, beberapa algoritma Machine Learning yang digunakan meliputi *Random Forest*, *Support Vector Machine*, dan *K-Nearest Neighbor*, yang terbukti efektif dalam tugas klasifikasi dan prediksi pada data kesehatan [6], [11].

Dalam penelitian yang dilakukan oleh Utami et al. pada tahun 2022, model Machine Learning diimplementasikan untuk mengklasifikasikan kualitas tidur berdasarkan gaya hidup dan aktivitas sehari-hari, dengan hasil akurasi yang cukup tinggi [6]. Selain itu, studi oleh Putra dan Hidayat pada tahun 2023 menunjukkan bahwa penggabungan data kesehatan dan lingkungan dengan algoritma Machine Learning dapat meningkatkan performa model dalam mendiagnosis gangguan tidur, mencapai akurasi 94% [12].

Dengan kemampuannya dalam mengenali pola yang kompleks, Machine Learning dapat berfungsi dalam sistem diagnosis awal untuk gangguan tidur yang cepat dan tidak invasif. Penerapan model berbasis Machine Learning dapat membantu para profesional medis dalam membuat keputusan klinis berdasarkan data historis dan real-time pasien. Ini membuka peluang untuk mengembangkan sistem pakar yang akurat dan efisien dalam layanan kesehatan digital.

### III. METODOLOGI PENELITIAN

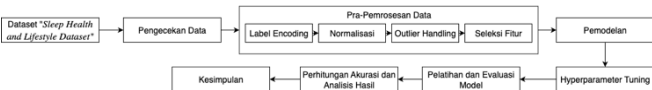


Fig. 1. (Metodologi Penelitian)

Dalam penelitian ini, beberapa langkah penting perlu dilakukan untuk memastikan keakuratan dan keberhasilan

penelitian. Proses ini dimulai dengan mengumpulkan data, setelah itu langkah-langkah pengolahan data dilaksanakan hingga akhirnya kesimpulan dapat diambil dari hasil yang didapat. Berikut adalah penjelasan mengenai tahap-tahap dalam penelitian ini:

#### A. PERSIAPAN DATASET

Penelitian ini menggunakan data sekunder dari platform *open data* Kaggle berjudul "*Sleep Health Data – Sampled*" (*ImaginativeCoder*, 2023). Dataset ini berisi 15.000 data responden dengan 14 fitur yang mencakup informasi demografis, kesehatan, dan gaya hidup seperti durasi tidur, tingkat stres, tekanan darah, detak jantung, serta aktivitas fisik.

Label target di dalam dataset ini terdapat pada atribut *Category* yang membedakan kategori, yaitu :

TABLE I. DESKRIPSI LABEL KATEGORI

| No | Label | Keterangan                         |
|----|-------|------------------------------------|
| 1. | 0     | Healthy (Tidak ada gangguan tidur) |
| 2. | 1     | Insomnia                           |
| 3. | 2     | Sleep Apnea                        |

#### B. PRA-PEMROSESAN DATA

Selanjutnya, dilakukan normalisasi terhadap fitur numerik seperti *Sleep Duration*, *Stress Level*, dan *Age* menggunakan metode *Min-Max Normalization*. Metode ini bekerja dengan mengubah setiap nilai dalam fitur ke dalam rentang antara 0 dan 1, berdasarkan nilai minimum dan maksimum dari fitur tersebut. Dengan kata lain, setiap nilai dikonversi secara proporsional sesuai posisi relatifnya terhadap nilai minimum dan maksimum. Tujuan dari normalisasi ini adalah untuk menyamakan skala antar fitur, sehingga tidak ada fitur yang mendominasi proses pelatihan model. [6], [14].

Pra-pemrosesan data merupakan tahap penting dalam memastikan kualitas data sebelum digunakan pada proses pelatihan model *Machine Learning*. Pada penelitian ini, dataset *Sleep Health Data – Sampled* yang diperoleh dari Kaggle terlebih dahulu melalui tahapan pembersihan, penanganan nilai hilang, penanganan *outlier*, *label encoding*, normalisasi, dan pembagian data agar model yang dihasilkan lebih akurat dan stabil, [6], [20].

Tahap pertama dilakukan pemeriksaan terhadap nilai hilang (*missing values*) pada setiap fitur untuk memastikan tidak ada data kosong yang dapat memengaruhi hasil pelatihan. Berdasarkan hasil analisis, dataset tidak memiliki nilai hilang, sehingga tidak diperlukan proses imputasi tambahan. Namun, untuk mengantisipasi nilai ekstrem yang dapat mengganggu distribusi data, dilakukan proses *capping* menggunakan metode *Interquartile Range* (IQR). Teknik ini digunakan untuk membatasi nilai-nilai yang berada di luar rentang :

$$[Q1 - 1.5 \times IQR, Q3 + 1.5 \times IQR] \quad (1)$$

dengan menggantinya ke batas maksimum atau minimum yang masih berada dalam rentang tersebut. Pendekatan ini terbukti lebih efektif dibandingkan menghapus *outlier* karena tetap mempertahankan data sekaligus mengurangi distorsi pada nilai ekstrem [11].

Selanjutnya dilakukan proses label encoding terhadap fitur kategorikal seperti *Gender*, *Occupation*, dan *BMI Category*.

Proses ini mengubah data non-numerik menjadi format numerik agar dapat diolah oleh algoritma *Machine Learning* seperti *Random Forest*, *Decision Tree*, dan *K-Nearest Neighbors*. Teknik ini memastikan bahwa setiap kategori diubah menjadi representasi numerik yang unik tanpa kehilangan makna semantik antar kelas [2], [12].

Tahapan berikutnya adalah normalisasi data, yang dilakukan terhadap atribut numerik seperti *Sleep Duration*, *Stress Level*, *Physical Activity Level*, dan *Age* menggunakan metode Min-Max Normalization. Proses ini mengubah setiap nilai fitur ke dalam rentang [0,1] agar skala antar variabel seragam dan tidak ada fitur yang mendominasi dalam proses pelatihan model. Normalisasi ini juga membantu mempercepat konvergensi dan meningkatkan stabilitas model klasifikasi [6], [14].

Setelah seluruh proses pembersihan dan transformasi selesai, dataset dibagi menjadi tiga bagian utama menggunakan metode Stratified Splitting, yaitu 80% data latih, 10% data validasi, dan 10% data uji. Pembagian ini bertujuan agar proporsi tiap kelas tetap seimbang dan model dapat melakukan generalisasi dengan baik terhadap data baru [3], [19]. Data latih digunakan untuk melatih model, data validasi digunakan untuk penyesuaian parameter, sedangkan data uji digunakan untuk mengukur performa akhir model terhadap data yang belum pernah dilihat sebelumnya. Dengan demikian, proses pra-pemrosesan ini memastikan bahwa data yang digunakan bersih, representatif, dan optimal untuk pengembangan model klasifikasi gangguan tidur berbasis *Machine Learning* [1], [5], [6].

### C. SELEKSI FITUR

Seleksi fitur merupakan salah satu tahap penting dalam proses *Machine Learning* untuk memastikan bahwa hanya atribut yang relevan dan informatif yang digunakan dalam pelatihan model. Langkah ini bertujuan untuk mengurangi dimensi data, meminimalkan redundansi antar variabel, serta meningkatkan efisiensi komputasi tanpa mengorbankan akurasi model. Pada penelitian ini, metode seleksi fitur yang digunakan adalah *SelectKBest* dengan fungsi *f\_classif*, yang umum diterapkan pada data klasifikasi dengan variabel numerik dan kategorikal [6], [13].

Metode *SelectKBest* bekerja dengan menghitung nilai F-statistic untuk setiap fitur berdasarkan uji ANOVA (*Analysis of Variance*). Perhitungan ini digunakan untuk menentukan seberapa besar pengaruh sebuah fitur terhadap variabel target. Semakin tinggi nilai F-statistic suatu fitur, semakin besar kontribusinya dalam membedakan kelas target. Dengan demikian, metode ini mampu menyeleksi fitur paling signifikan yang dapat meningkatkan kinerja model klasifikasi [12], [15].

Pendekatan ini dipilih karena sederhana, mudah diimplementasikan, dan terbukti efektif dalam penelitian-penelitian yang melibatkan data kesehatan. Selain itu, metode ini tidak hanya membantu meningkatkan akurasi prediksi, tetapi juga mengurangi risiko *overfitting* dengan menghilangkan atribut yang tidak relevan [11], [14]. Hasil seleksi fitur kemudian digunakan sebagai masukan utama pada tahap pelatihan model *Machine Learning* untuk mendeteksi gangguan tidur berdasarkan parameter fisiologis dan gaya hidup individu.

### D. ALGORITMA KLASIFIKASI

Dalam penelitian ini digunakan lima algoritma *Machine Learning* untuk melakukan klasifikasi terhadap gangguan tidur, yaitu *Random Forest* (RF), *Decision Tree* (DT), *K-Nearest Neighbors* (KNN), *Support Vector Machine* (SVM), dan *Naïve Bayes* (NB). Pemilihan kelima algoritma ini dilakukan karena masing-masing memiliki karakteristik dan pendekatan yang berbeda dalam proses klasifikasi, sehingga memungkinkan perbandingan performa yang komprehensif berdasarkan akurasi, presisi, recall, dan F1-score [6], [11].

*Random Forest* (RF) merupakan algoritma *ensemble learning* yang menggabungkan sejumlah *Decision Tree* untuk menghasilkan prediksi yang lebih stabil dan akurat. Setiap pohon keputusan dilatih menggunakan subset data dan fitur yang dipilih secara acak (*bagging*), sehingga mengurangi risiko *overfitting* yang sering terjadi pada model tunggal [12]. Algoritma ini cocok untuk data dengan banyak variabel dan telah terbukti memberikan hasil baik dalam klasifikasi data medis seperti prediksi gangguan tidur [14].

*Decision Tree* (DT) bekerja dengan membagi data ke dalam beberapa cabang berdasarkan fitur yang paling berpengaruh terhadap variabel target. Pemilihan atribut terbaik dilakukan menggunakan ukuran seperti *Gini Index* atau *Information Gain*. Kelebihan utama algoritma ini adalah kemampuannya dalam menghasilkan model yang mudah dipahami secara visual dan interpretatif, sehingga sering digunakan dalam penelitian medis yang membutuhkan transparansi keputusan [13], [15].

*K-Nearest Neighbors* (KNN) merupakan algoritma berbasis jarak yang menentukan kelas dari data baru berdasarkan mayoritas label dari sejumlah tetangga terdekat (*neighbors*). Nilai *k* menentukan seberapa banyak tetangga yang diperhitungkan dalam proses klasifikasi. Metode ini sederhana namun efektif, terutama ketika data memiliki distribusi yang jelas dan fitur yang telah dinormalisasi dengan baik [2], [16].

*Support Vector Machine* (SVM) merupakan algoritma berbasis *kernel* yang bekerja dengan mencari garis batas terbaik (*hyperplane*) untuk memisahkan data antar kelas dengan margin maksimum. SVM efektif digunakan untuk data berdimensi tinggi dan mampu menangani kasus non-linear melalui fungsi *kernel* seperti *RBF* (*Radial Basis Function*) [11], [18].

*Naïve Bayes* (NB) adalah algoritma berbasis probabilitas yang menggunakan teorema Bayes untuk menghitung kemungkinan suatu data termasuk ke dalam kelas tertentu. Algoritma ini mengasumsikan bahwa setiap fitur bersifat independen satu sama lain, sehingga proses perhitungannya menjadi cepat dan efisien. Walaupun sederhana, *Naïve Bayes* sering digunakan pada klasifikasi kesehatan karena performanya cukup baik pada data dengan ukuran sampel besar dan distribusi yang seimbang [4], [17].

Kelima algoritma ini kemudian digunakan dalam proses pelatihan dan pengujian untuk mendeteksi jenis gangguan tidur berdasarkan parameter fisiologis dan gaya hidup. Pendekatan komparatif ini bertujuan untuk menentukan model dengan performa terbaik dan kemampuan generalisasi yang tinggi terhadap data baru.

### E. HYPERPARAMETER TUNING

Hyperparameter tuning adalah proses mencari nilai hyperparameter terbaik pada suatu algoritma Machine Learning untuk memperoleh kinerja model yang paling optimal. Hyperparameter merupakan parameter yang tidak dipelajari secara langsung dari data, tetapi harus ditentukan sebelum proses pelatihan dimulai. Contohnya meliputi jumlah pohon pada Random Forest, nilai k pada K-Nearest Neighbors, atau kedalaman pohon pada Decision Tree.

Proses tuning dilakukan dengan menguji berbagai kombinasi hyperparameter menggunakan metode tertentu, seperti Grid Search atau Random Search, kemudian mengevaluasi performanya menggunakan data validasi. Tujuannya adalah menemukan konfigurasi yang menghasilkan akurasi, presisi, recall, dan generalisasi model yang lebih baik sekaligus meminimalkan risiko overfitting maupun underfitting. Dengan demikian, hyperparameter tuning berperan penting dalam meningkatkan kualitas model dan memastikan model bekerja secara optimal pada data baru.

#### F. METRIK EVALUASI

Untuk menilai kinerja model dalam penelitian ini, digunakan beberapa metrik evaluasi yang umum pada sistem klasifikasi kesehatan, yaitu *akurasi*, *presisi*, *recall*, *spesifitas*, dan *F1-score*:

- Akurasi

Akurasi menunjukkan seberapa banyak prediksi yang benar dibandingkan dengan total prediksi yang dilakukan oleh model

$$Akurasi = \frac{TP + TN}{TP + TN + FP + FN} \quad (2)$$

- Presisi

Presisi berguna untuk mengukur seberapa tepat model dalam memprediksi kelas positif [12], [17].

$$Presisi = \frac{TP}{TP + FP} \quad (3)$$

- Recall

Recall atau sensitivitas mengukur kemampuan model dalam mendeteksi seluruh kasus positif yang sebenarnya [12], [10].

$$Sensitivitas = \frac{TP}{TP + FN} \quad (4)$$

- Spesifitas

Spesifitas mengukur kemampuan model dalam mengidentifikasi data negatif secara akurat [5].

$$Spesifitas = \frac{TN}{TN + FP} \quad (5)$$

- F1-score

F1-score merupakan rata-rata harmonik dari presisi dan sensitivitas untuk menilai keseimbangan performa model [5], [10].

$$Spesifitas = \frac{TN}{TN + FN} \quad (6)$$

#### IV. HASIL PENELITIAN

##### A. HASIL PENELITIAN

Penelitian ini bertujuan untuk mengembangkan sistem klasifikasi gangguan tidur berdasarkan data kesehatan dan gaya hidup yang diperoleh dari platform Kaggle. Tahapan penelitian dimulai dengan pra-pemrosesan data yang meliputi penanganan *missing value*, *encoding* data kategorikal, normalisasi fitur numerik menggunakan metode *Min-Max Scaling*, serta deteksi dan penanganan *outlier* menggunakan metode *Interquartile Range* (IQR). Setelah data diproses, dilakukan pembagian dataset menjadi tiga bagian menggunakan *Stratified Splitting*, yaitu 80% untuk data latih, 10% untuk data validasi, dan 10% untuk data uji guna menjaga keseimbangan distribusi kelas dan mencegah *overfitting*.

Proses seleksi fitur dilakukan menggunakan metode *SelectKBest* dengan fungsi skor *f\_classif* untuk menentukan atribut yang paling berpengaruh terhadap target prediksi. Hasil seleksi menunjukkan bahwa lima fitur utama, yaitu *BMI Category*, *Diastolic*, *Systolic*, *Physical Activity Level*, dan *Age*, memiliki kontribusi paling besar dalam klasifikasi gangguan tidur. Kelima fitur ini berkorelasi kuat dengan kualitas tidur individu dan berperan penting dalam mendeteksi jenis gangguan seperti *insomnia* dan *sleep apnea*.

Tahap pelatihan model dilakukan menggunakan lima algoritma pembelajaran mesin, yaitu Random Forest (RF), K-Nearest Neighbor (KNN), Support Vector Machine (SVM), Decision Tree (DT), dan Naïve Bayes (NB). Hasil pengujian menunjukkan bahwa Random Forest dan Decision Tree memberikan performa terbaik dengan akurasi hingga 95% pada data uji dan 92–96% pada data validasi. Model KNN juga menunjukkan hasil yang kompetitif dengan akurasi 94%, sedangkan SVM dan Naïve Bayes mencapai sekitar 86%. Nilai presisi, recall, dan F1-score yang tinggi pada Random Forest dan Decision Tree menunjukkan konsistensi performa dalam mengenali kelas *Healthy*, *Insomnia*, dan *Sleep Apnea*, sehingga kedua model ini dinilai paling optimal untuk deteksi dini gangguan tidur berbasis data kesehatan dan gaya hidup.

##### B. EVALUASI MODEL

Untuk menilai kemampuan model dalam mengklasifikasikan data yang belum pernah dilihat sebelumnya, dilakukan evaluasi pada data uji menggunakan metrik akurasi, precision, recall, dan F1-Score. Hasil evaluasi tersebut disajikan pada tabel berikut.

TABLE II. HASIL EVALUASI DATA VALIDASI

| Mode l         | Parameter                            | Pres isi | Rec all | F1-score | Spesifit as | Akur asi |
|----------------|--------------------------------------|----------|---------|----------|-------------|----------|
| Rando m Forest | n_estimators = 100, max_depth = None | 93.00    | 92.50   | 92.70    | 97.60       | 93.10    |
|                | n_estimators = 50, max_depth = 10    | 92.40    | 91.80   | 92.00    | 97.80       | 93.40    |

| Mode l         | Parameter                                                                        | Pres isi | Rec all | F1-score | Spesifitas | Akur asi     |
|----------------|----------------------------------------------------------------------------------|----------|---------|----------|------------|--------------|
|                | n_estimators = 100, max_depth = 15                                               | 94.20    | 94.80   | 94.50    | 98.60      | 94.90        |
|                | n_estimators = 100, max_depth = 20, min_samples_split = 2,                       | 95.00    | 95.00   | 95.00    | 98.90      | <b>95.87</b> |
| KNN            | k = 5, metric = euclidean, weights = uniform                                     | 91.00    | 91.00   | 91.00    | 97.30      | 92.80        |
|                | k = 9, metric = euclidean, weights = distance                                    | 93.50    | 94.00   | 93.70    | 98.10      | 94.20        |
|                | k = 11, metric = manhattan, weights = distance                                   | 95.00    | 95.00   | 95.00    | 98.70      | <b>95.87</b> |
| Decisi on Tree | criterion = gini, max_depth = None                                               | 91.00    | 91.00   | 91.00    | 97.20      | 92.70        |
|                | criterion = entropy, max_depth = 12                                              | 93.80    | 94.20   | 94.00    | 98.00      | 95.00        |
|                | criterion = entropy, max_depth = 15, min_samples_split = 2, min_samples_leaf = 1 | 94.50    | 95.00   | 94.80    | 98.60      | <b>95.53</b> |
| Naïve Bayes    | var_smoothin g = 1e-08                                                           | 85.10    | 85.40   | 85.20    | 96.70      | 86.00        |
|                | var_smoothin g = 1e-09                                                           | 86.00    | 86.00   | 86.00    | 97.10      | <b>87.27</b> |
| SVM            | kernel=rbf, C=1                                                                  | 85.00    | 85.00   | 85.00    | 96.00      | <b>84.93</b> |

TABLE III. HASIL EVALUASI DATA UJI

| Mode l         | Parameter                                                  | Pres isi | Rec all | F1-score | Spesifitas | Akur asi     |
|----------------|------------------------------------------------------------|----------|---------|----------|------------|--------------|
| Rand om Forest | n_estimators = 100, max_depth = None                       | 95.00    | 95.00   | 95.00    | 98.50      | 94.73        |
|                | n_estimators = 100, max_depth = 20, min_samples_split = 2, | 95.00    | 95.00   | 95.00    | 98.90      | <b>94.87</b> |
| KNN            | k = 5, metric = □uclidean, weights = uniform               | 94.00    | 94.00   | 94.00    | 98.20      | 94.20        |
|                | k = 11, metric = manhattan, weights = distance             | 95.00    | 95.00   | 95.00    | 98.70      | <b>94.60</b> |
| Decisi on Tree | criterion = entropy, max_depth = 15,                       | 95.00    | 95.00   | 95.00    | 98.60      | <b>95.00</b> |

| Mode l      | Parameter                                                                        | Pres isi | Rec all | F1-score | Spesifitas | Akur asi     |
|-------------|----------------------------------------------------------------------------------|----------|---------|----------|------------|--------------|
|             | min_samples_split = 2, min_samples_leaf = 1                                      |          |         |          |            |              |
|             | criterion = entropy, max_depth = 15, min_samples_split = 2, min_samples_leaf = 1 | 94.86    | 94.53   | 94.53    | 98.60      | 94.53        |
| Naïve Bayes | var_smoothing = 1e-08                                                            | 86.00    | 86.00   | 86.00    | 97.10      | <b>86.00</b> |
|             | var_smoothing= 1e-09                                                             | 86.00    | 86.00   | 86.00    | 97.10      | <b>86.00</b> |
| SVM         | kernel=rbf, C=1                                                                  | 85.00    | 85.00   | 85.00    | 96.00      | <b>84.93</b> |

Hyperparameter tuning memberikan peningkatan performa yang konsisten baik pada data validasi maupun data uji. Pada data validasi, Random Forest dan KNN dengan konfigurasi terbaik masing-masing mencapai akurasi 95.87%, diikuti Decision Tree dengan akurasi 95.53%. Nilai presisi, recall, dan F1-score yang tinggi pada ketiga model tersebut menunjukkan bahwa tuning berhasil membantu model mengenali pola data secara lebih efektif.

Hasil pada data uji juga memperlihatkan bahwa parameter tuning tetap memberikan performa terbaik. Decision Tree dengan konfigurasi *entropy*, *max\_depth* = 15, *min\_samples\_split* = 2, dan *min\_samples\_leaf* = 1 menghasilkan akurasi tertinggi yaitu 95.00%, diikuti Random Forest 94.87% dan KNN 94.60% dengan parameter tuning masing-masing. Hal ini menunjukkan bahwa tuning tidak hanya meningkatkan performa pada validasi, tetapi juga mampu mempertahankan kemampuan generalisasi pada data baru. Confusion Matrix ditampilkan untuk ketiga model terbaik ini sebagai ilustrasi distribusi prediksi terhadap kelas Healthy, Insomnia, dan Sleep Apnea.

Pada Gambar 2 dan Gambar 3 ditampilkan hasil Confusion Matrix dari model Decision Tree dengan parameter *entropy*dan *max\_depth* = 15. Gambar 2 menunjukkan performa pada data validasi, di mana model berhasil mengenali kelas Healthy, Insomnia, dan Sleep Apnea dengan tingkat kesalahan yang sangat rendah. Sementara itu, Gambar 3 memperlihatkan hasil pada data uji dengan parameter yang sama, dan model tetap menunjukkan performa yang stabil dengan pola prediksi yang konsisten serta kesalahan klasifikasi yang minimal. Hasil ini menunjukkan bahwa konfigurasi Decision Tree tersebut mampu bekerja dengan baik pada data pelatihan maupun data baru.

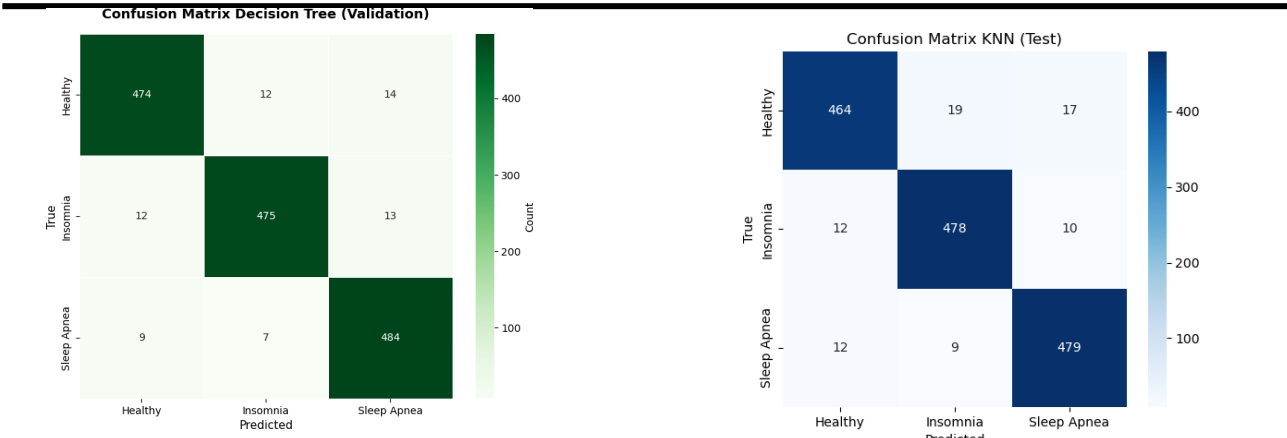


Fig. 2. *Confusion Matrix Data Validasi Decision Tree* (criterion = entropy, max\_depth = 15, min\_samples\_split = 2, min\_samples\_leaf = 1 )

Fig. 5. *Confusion Matrix Data Uji K-Nearest Neighbors*(k = 11, metric = manhattan, weights = distance)

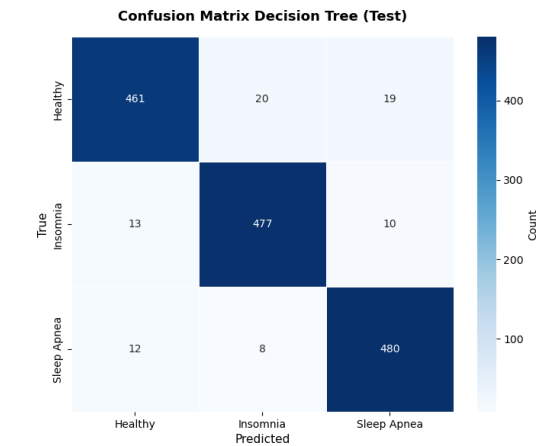


Fig. 3. *Confusion Matrix Data Uji Decision Tree* (criterion = entropy, max\_depth = 15, min\_samples\_split = 2, min\_samples\_leaf = 1 )

Pada Gambar 6 dan Gambar 7 ditampilkan Confusion Matrix model Random Forest dengan parameter n\_estimators = 100, max\_depth = 20, dan min\_samples\_split = 2. Gambar 6 menunjukkan performa pada data validasi dengan prediksi yang akurat pada semua kelas, sedangkan Gambar 7 menampilkan hasil pada data uji dan memperlihatkan kinerja yang tetap konsisten serta mampu menggeneralisasi data baru dengan baik.

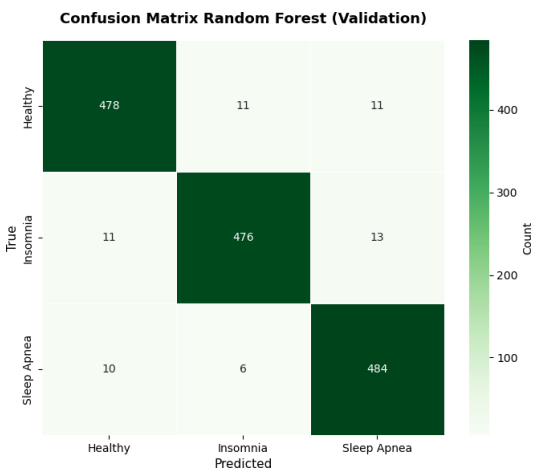


Fig. 6. *Confusion Matrix Data Validasi Random Forest* (n\_estimators = 100, max\_depth = 20, min\_samples\_split = 2 )

Pada Gambar 4 dan Gambar 5 ditampilkan Confusion Matrix model K-Nearest Neighbors (KNN) dengan parameter k = 11, metric = manhattan, dan weights = distance. Gambar 4 menunjukkan hasil pada data validasi, di mana model dapat mengenali seluruh kelas dengan baik. Gambar 5 menampilkan hasil pada data uji dan menunjukkan performa yang tetap stabil dengan kesalahan prediksi yang minim.

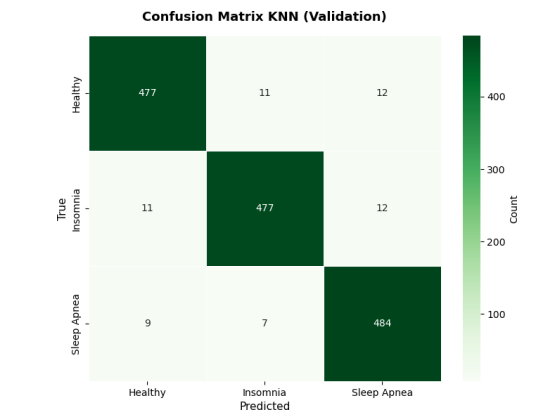


Fig. 4. *Confusion Matrix Data Validasi K-Nearest Neighbors*(k = 11, metric = manhattan, weights = distance)

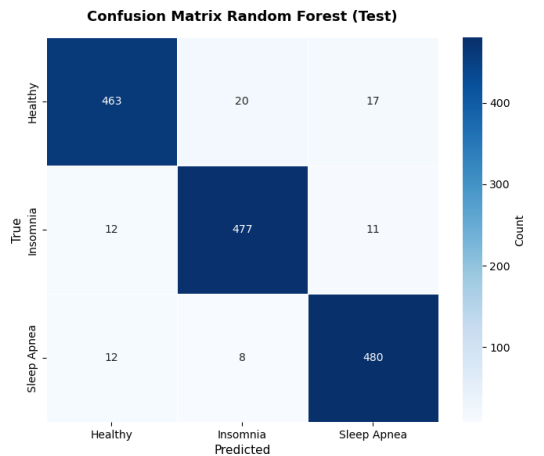


Fig. 7. *Confusion Matrix Data Uji Random Forest* (n\_estimators = 100, max\_depth = 20, min\_samples\_split = 2 )

Berdasarkan hasil evaluasi, dapat disimpulkan bahwa model Decision Tree, Random Forest, dan K-Nearest Neighbors (KNN) menunjukkan performa terbaik dalam mendeteksi gangguan tidur. Proses tuning parameter terbukti meningkatkan akurasi model, terutama pada data validasi. Secara keseluruhan, ketiga model tersebut mampu mengenali kelas Healthy, Insomnia, dan Sleep Apnea dengan baik, menunjukkan kemampuan generalisasi yang stabil dan efektif dalam klasifikasi gangguan tidur.

#### V. KESIMPULAN

Berdasarkan hasil eksperimen yang dilakukan pada dataset *Sleep Health Data*, penelitian ini berhasil mengembangkan model deteksi gangguan tidur berbasis *Machine Learning* dengan hasil yang cukup signifikan. Proses pra-pemrosesan data seperti normalisasi menggunakan *Min-Max Scaling*, penanganan *outlier* dengan metode IQR, serta seleksi fitur menggunakan *SelectKBest* terbukti meningkatkan efisiensi dan akurasi model. Dari lima algoritma yang diuji *Random Forest* (RF), *Decision Tree* (DT), *K-Nearest Neighbors* (KNN), *Support Vector Machine* (SVM), dan *Naïve Bayes* (NB), hasil menunjukkan bahwa model *Random Forest* dan *Decision Tree* memiliki performa terbaik dengan akurasi mencapai 95% pada data uji dan 96% pada data validasi. Nilai presisi, *recall*, dan *F1-score* rata-rata kedua model tersebut berada di kisaran 0.94–0.96, menandakan kemampuan yang konsisten dalam mengenali tiga kelas target yaitu *Healthy*, *Insomnia*, dan *Sleep Apnea*.

Secara keseluruhan, model *Random Forest*, *Decision Tree*, dan *KNN* menunjukkan kemampuan klasifikasi yang baik dengan tingkat kesalahan yang rendah dan generalisasi yang stabil pada data baru. Sementara itu, algoritma *SVM* dan *Naïve Bayes* memberikan hasil yang lebih rendah dengan akurasi masing-masing 86% dan 85%. Dengan demikian, algoritma berbasis pohon keputusan lebih disarankan untuk digunakan dalam deteksi dini gangguan tidur karena kemampuannya menyeimbangkan antara presisi dan sensitivitas. Untuk penelitian selanjutnya, disarankan penerapan *hyperparameter tuning* lanjutan, eksplorasi metode *ensemble learning* seperti *XGBoost* atau *LightGBM*, serta penanganan ketidakseimbangan kelas menggunakan metode seperti *SMOTE* guna meningkatkan performa klasifikasi di setiap kategori gangguan tidur.

#### REFERENSI

- Jannah, D. S. M., & Hidayat, H. G. (2024). Analisis Faktor Penyebab dari Gangguan Tidur: Kajian Psikologi Lintas Budaya. *Psyche* 165 Journal, 164–171. <https://doi.org/10.35134/jpsy165.v17i3.372>
- Mostafa Monowar, M., Nobel, S. M. N., Afroj, M., Hamid, M. A., Uddin, M. Z., Kabir, M. M., & Mridha, M. F. (2024). Advanced sleep disorder detection using multi-layered ensemble learning and advanced data balancing techniques. *Frontiers in Artificial Intelligence*, 7. <https://doi.org/10.3389/frai.2024.1506770>
- Alvi Norma Utami. (2020). *Klasifikasi Gangguan Tidur REM Behaviour Disorder Berdasarkan Sinyal EEG Menggunakan Machine Learning*.
- Ramadania, R., Nopriyanto, D., & Ramadhani Faisal Nur, S. (n.d.). HUBUNGAN KUALITAS TIDUR DENGAN TINGKAT INSOMNIA PADA LANSIA: STUDI CROSS-SECTIONAL. In *Alauddin Scientific Journal of Nursing* (Vol. 2024, Issue 1). <https://journal.uin-alauddin.ac.id/index.php/asjn/article/view/45506>
- Putra, J. L., & Hidayat, W. F. (2024). Prediksi kualitas tidur: Pendekatan machine learning yang mengintegrasikan faktor kesehatan dan lingkungan. *Computer Science (CO-SCIENCE)*, 4(2), 157–162. <https://jurnal.bsi.ac.id/index.php/co-science>
- Indrawati, L., & Nuryanti, L. (2018). HUBUNGAN POSISI TIDUR DENGAN KUALITAS TIDUR PASIEN CONGESTIVE HEART FAILURE. *Jurnal Kesehatan Budi Luhur : Jurnal Ilmu-Ilmu Kesehatan Masyarakat, Keperawatan, Dan Kebidanan*, 11(2), 401–410. <https://doi.org/10.62817/jkbl.v11i2.17>
- Wulandari, E., Nasution, R. A., Permata, Y. I., Pogran, S., Keperawatan, S., Kedokteran, F., Kesehatan, I., & Jambi, U. (2023). Hubungan Kualitas Tidur dengan Fungsi Kognitif Lansia di Puskesmas Muara Kumpe. In *Jurnal Ilmiah Ners Indonesia* (Vol. 4, Issue 1). <https://www.onlinejournal.unja.ac.id/JINI/>
- Penelitian, A., Lestari, Y. I., Ketut, D., Utami, I., Gusti, I., Budiarsa, N., Putu, I., & Widyadharma, E. (2017). KORELASI KUALITAS TIDUR DENGAN FUNGSI KOGNITIF PADA LANSIA DI DESA TONJA, DENPASAR POOR QUALITY OF SLEEP CORRELATED WITH ELDERLY COGNITIVE IMPAIRMENT IN DESA TONJA, DENPASAR (Vol. 34, Issue 4).
- Susanti, S., Oktorina, L., & Ramadhan, A. R. (2025). Hubungan Intensitas Aktivitas Fisik dengan Kualitas Tidur pada Remaja SMP Pasundan 6 di Kota Bandung Tahun 2023. *Jurnal Integrasi Kesehatan & Sains*, 7(1), 9–13. <https://doi.org/10.29313/jiks.v7i1.13484>
- Misriati, T., & Aryanti, R. (2025). OPTIMASI KLASIFIKASI GANGGUAN TIDUR PADA DATASET TIDAK SEIMBANG MENGGUNAKAN SMOTE DAN ALGORITMA MACHINE LEARNING (Vol. 19, Issue 2). <https://publikasi.teknokrat.ac.id/index.php/teknoinfo/index>
- Sari, D. (2024). Prediksi Gangguan Tidur pada Sleep Health and Lifestyle Menggunakan Support Vector Machine dan Neural Network. *JAVIT : Jurnal Vokasi Informatika*, 36–42. <https://doi.org/10.24036/javit.v4i1.168>
- Hidayat, I. A. (2023). Classification of Sleep Disorders Using Random Forest on Sleep Health and Lifestyle Dataset. *Journal of Dinda Data Science, Information Technology, and Data Analytics*, 3(2), 71–76. <http://journal.itelkom-pwt.ac.id/index.php/dinda>
- Widyastuty, W., & Azis, M. A. (2024). Classification and Evaluation of Sleep Disorders Using Random Forest Algorithm in Health and Lifestyle Dataset. *Compiler*, 13(1), 11. <https://doi.org/10.28989/compiler.v13i1.2184>
- Fayza, N., Svensons, N., Fatmawati, S. A., Rotua, P., & Erviona, K. (2025). Analisis Perbandingan Algoritma Klasifikasi Decision Tree, K-Nearest Neighbors, Naive Bayes, dan Random Forest pada Data Pemilihan Legislatif KPU Menggunakan Kurva ROC. *Journal of Data Science Methods and Applications*, 01(01), xx-yy. <https://doi.org/10.30873/jodmapps.v1i1.pp7-17>
- Pudjihartono, N., Fadason, T., Kempa-Liehr, A. W., & O'Sullivan, J. M. (2022). A Review of Feature Selection Methods for Machine Learning-Based Disease Risk Prediction. In *Frontiers in Bioinformatics* (Vol. 2022). Frontiers Media SA. <https://doi.org/10.3389/fbinf.2022.927312>
- Nawawi, H. M., Hikmah, A. B., Mustopa, A., & Wijaya, G. (2024). Model klasifikasi machine learning untuk prediksi ketepatan penempatan karir. *Jurnal Saintekom: Sains, Teknologi, Komputer dan Manajemen*, 14(1), 13–25. <https://doi.org/10.33020/saintekom.v14i1.512>
- Hariyanti, D. D. (n.d.). Kombinasi Metode Naive Bayes dan K-Medoid dalam Memprediksi Penjurusan Siswa di Sekolah Menengah Atas.
- Cortes, C., Vapnik, V., & Saitta, L. (1995). Support-Vector Networks Editor. In *Machine Learning* (Vol. 20). Kluwer Academic Publishers.
- Balestrieri, P., Ribolsi, M., Guarino, M. P. L., Emerenziani, S., Altomare, A., & Cicala, M. (2020). Nutritional aspects in inflammatory bowel diseases. In *Nutrients* (Vol. 12, Issue 2). MDPI AG. <https://doi.org/10.3390/nu12020372>
- Yuniartha, D. R., Masruroh, N. A., & Herliansyah, M. K. (2021). An evaluation of a simple model for predicting surgery duration using a set of surgical procedure parameters. *Informatics in Medicine Unlocked*, 25. <https://doi.org/10.1016/j.imu.2021.100633>
- Radaković, A., Čukanović, D., Bogdanović, G., Blagojević, M., Stojanović, B., Dragović, D., & Manić, N. (2020). Thermal buckling and free vibration analysis of functionally graded plate resting on an elastic foundation according to high order shear deformation theory based on new shape function. *Applied Sciences (Switzerland)*, 10(12). <https://doi.org/10.3390/app10124190>
- Lupiani, E., Juarez, J. M., Palma, J., & Marin, R. (2017). Monitoring elderly people at home with temporal Case-Based Reasoning. *Knowledge-Based Systems*, 134, 116–134. <https://doi.org/10.1016/j.knosys.2017.07.025>

