

**KLASIFIKASI KEBUTUHAN KONSULTASI KESEHATAN MENTAL
PEKERJA *REMOTE* MENGGUNAKAN METODE
*RANDOM FOREST***

SKRIPSI

Oleh:
FAUZIAH PUTRI UTAMI
NIM. 220605110172



**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2026**

**KLASIFIKASI KEBUTUHAN KONSULTASI KESEHATAN MENTAL
PEKERJA *REMOTE* MENGGUNAKAN METODE
*RANDOM FOREST***

SKRIPSI

Diajukan kepada:
Universitas Islam Negeri Maulana Malik Ibrahim Malang
Untuk memenuhi Salah Satu Persyaratan dalam
Memperoleh Gelar Sarjana Komputer (S.Kom)

Oleh:
FAUZIAH PUTRI UTAMI
NIM. 220605110172

**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2026**

HALAMAN PERSETUJUAN

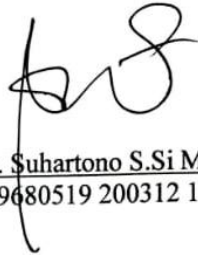
**KLASIFIKASI KEBUTUHAN KONSULTASI KESEHATAN MENTAL
PEKERJA *REMOTE* MENGGUNAKAN METODE
*RANDOM FOREST***

SKRIPSI

Oleh:
FAUZIAH PUTRI UTAMI
NIM. 220605110172

Telah Diperiksa dan Disetujui untuk Diuji:
Tanggal: 17 Juni 2026

Pembimbing I,



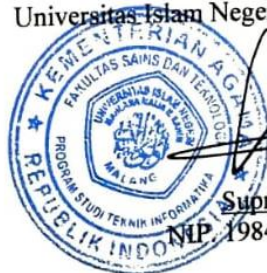
Prof. Dr. Suhartono S.Si M.Kom
NIP. 19680519 200312 1 001

Pembimbing II,



Nurizal Dwi Priandani, M.Kom
NIP. 19920830 202203 1 001

Mengetahui,
Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi
Universitas Islam Negeri Maulana Malik Ibrahim Malang



Supriyono, M.Kom
NIP. 19841010 201903 1 012

HALAMAN PENGESAHAN

KLASIFIKASI KEBUTUHAN KONSULTASI KESEHATAN MENTAL PEKERJA *REMOTE* MENGGUNAKAN METODE *RANDOM FOREST*

SKRIPSI

Oleh:

FAUZIAH PUTRI UTAMI

NIM. 220605110172

Telah Dipertahankan di Depan Dewan Penguji Skripsi
dan Dinyatakan Diterima Sebagai Salah Satu Persyaratan
Untuk Memperoleh Gelar Sarjana Komputer (S.Kom)
Tanggal: 22 Juni 2026

Susunan Dewan Penguji

Ketua Penguji : Dr. Totok Chamidy, M.Kom
NIP. 19691222 200604 1 001

Anggota Penguji I : Nur Fitriyah Tunjung Sari, M.Cs
NIP. 19911226 202012 2 001

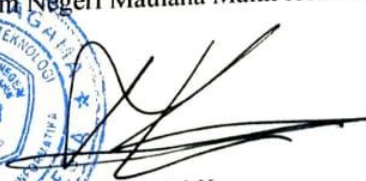
Anggota Penguji II : Prof. Dr. Suhartono S.Si M.Kom
NIP. 19680519 200312 1 001

Anggota Penguji III : Nurizal Dwi Priandani, M.Kom
NIP. 19920830 202203 1 001

()
()
()
()

Mengetahui dan Mengesahkan,
Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi
Universitas Islam Negeri Maulana Malik Ibrahim Malang




Supriyono, M.Kom

NIP. 19841010 201903 1 012

PERNYATAAN KEASLIAN TULISAN

Saya yang bertanda tangan di bawah ini:

Nama : Fauziah Putri Utami
NIM : 220605110172
Fakultas / Program Studi : Sains dan Teknologi / Teknik Informatika
Judul Skripsi : Klasifikasi Kebutuhan Konsultasi Kesehatan
Mental Pekerja *Remote* Menggunakan Metode
Random Forest

Menyatakan dengan sebenarnya bahwa Skripsi yang saya tulis ini benar-benar merupakan hasil karya saya sendiri, bukan merupakan pengambil alihan data, tulisan, atau pikiran orang lain yang saya akui sebagai hasil tulisan atau pikiran saya sendiri, kecuali dengan mencantumkan sumber cuplikan pada daftar pustaka.

Apabila dikemudian hari terbukti atau dapat dibuktikan skripsi ini merupakan hasil jiplakan, maka saya bersedia menerima sanksi atas perbuatan tersebut.

Malang, 26 Juni 2026
Yang membuat pernyataan,



Fauziah Putri Utami
NIM. 220605110172

MOTTO

... *“Lelah boleh singgah, tapi menyerah bukan pilihan bagi jiwa yang telah berjuang sejauh ini” ...*

... *“Dan bersabarlah kamu, sesungguhnya janji Allah Adalah benar” ...*
(Qs. Ar-Ruum:60)

HALAMAN PERSEMBAHAN

Bismillahirrahmanirrahim

Dengan mengucapkan syukur kepada Allah Subhanahu Wa Ta'ala atas segala rahmat dan karunia-Nya, karya sederhana ini penulis persembahkan kepada:

Kedua Orang Tua Tercinta, yang selalu memberikan doa, kasih sayang, dukungan, dan pengorbanan tanpa henti.

Seluruh Keluarga, yang senantiasa memberikan semangat, doa, dan motivasi dalam setiap langkah penulis.

Dosen Pembimbing dan Dosen Penguji, atas ilmu, bimbingan, serta arahan yang diberikan selama proses penyusunan skripsi.

Teman-Teman Seperjuangan, atas kebersamaan, dukungan, dan semangat selama menempuh pendidikan.

KATA PENGANTAR

Assalamu'alaikum Wr. Wb.

Alhamdulillah, segala puji dan rasa syukur penulis panjatkan ke hadirat Allah Subhanahu Wa Ta'ala atas rahmat, karunia, dan hidayah-Nya, sehingga penulis dapat menyelesaikan skripsi yang berjudul “Klasifikasi Kebutuhan Konsultasi Kesehatan Mental Pekerja *Remote* Menggunakan Metode *Random Forest*”. Sholawat serta salam tetap tercurahkan kepada junjungan kita Nabi Muhammad Shallallahu Alaihi Wasallam, yang menjadi suri tauladan dan semoga kita mendapat syafaat di akhir kelak, Aamiin.

Dalam penyusunan skripsi ini, penulis menyadari bahwa banyak pihak telah memberikan doa dan dukungan. Oleh karena itu, penulis mengucapkan terima kasih yang sebesar-besarnya kepada:

1. Prof. Dr. Hj. Ilfi Nur Diana, M.Si., CAHRM., CRMP., selaku Rektor Universitas Islam Negeri Maulana Malik Ibrahim Malang.
2. Dr. Agus Mulyono, M.Kes., selaku Dekan Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang.
3. Supriyono, M.Kom., selaku Ketua Program Studi Teknik Informatika Universitas Islam Negeri Maulana Malik Ibrahim Malang.
4. Prof. Dr. Suhartono, S.Si., M.Kom., selaku Dosen Pembimbing I yang telah meluangkan waktu, tenaga, dan pikiran untuk memberikan bimbingan, arahan, masukan, serta motivasi kepada penulis selama proses penyusunan skripsi.

5. Nurizal Dwi Priandani, M.Kom., selaku Dosen Pembimbing II yang telah memberikan saran, masukan, serta bimbingan dengan penuh kesabaran sehingga penelitian ini dapat diselesaikan dengan baik.
6. Dr. Totok Chamidy, M.Kom., selaku Ketua Penguji yang telah memberikan kritik, saran, serta masukan yang sangat bermanfaat untuk penyempurnaan skripsi ini.
7. Nur Fitriyah Tunjung Sari, M.Cs., selaku Dosen Penguji yang telah memberikan evaluasi, saran, dan masukan yang membangun sehingga skripsi ini menjadi lebih baik.
8. Kedua orang tua tercinta, yang selalu memberikan kasih sayang, doa, dukungan, pengorbanan, serta semangat yang tidak pernah putus kepada penulis. Terima kasih atas segala perjuangan dan kepercayaan yang diberikan sehingga penulis dapat menyelesaikan pendidikan hingga tahap ini.
9. Seluruh keluarga, yang senantiasa memberikan doa, perhatian, motivasi, dan dukungan sehingga penulis tetap semangat dalam menyelesaikan skripsi ini.
10. Nopal, yang selalu hadir memberikan bantuan, dukungan, semangat, serta menjadi tempat berbagi cerita dan berdiskusi selama proses penyusunan skripsi maupun dalam berbagai hal lainnya.
11. Mila, yang telah memberikan bantuan, motivasi, serta dukungan kepada penulis selama proses penelitian dan penyusunan skripsi.
12. Hilya, yang telah banyak membantu dalam proses penulisan skripsi, memberikan masukan, serta selalu memberikan semangat kepada penulis.

13. Silvi, yang telah memberikan bantuan, motivasi, dan dukungan selama proses penyusunan skripsi hingga dapat terselesaikan dengan baik.
14. Teman-teman Angkatan 2022 "INFINITY", yang telah menjadi bagian dari perjalanan selama masa perkuliahan. Terima kasih atas kebersamaan, pengalaman, kerja sama, dan dukungan yang telah diberikan selama menempuh pendidikan.
15. Diri sendiri, yang telah mampu bertahan, berjuang, bangkit, dan tidak menyerah dalam menghadapi setiap tantangan selama proses perkuliahan hingga penyusunan skripsi ini. Terima kasih karena telah terus berusaha dan percaya bahwa setiap proses akan menghasilkan sesuatu yang bermakna.

Akhir kata, semoga Allah Subhanahu Wa Ta'ala senantiasa membalas segala kebaikan, bantuan, dan doa yang telah diberikan oleh semua pihak dengan balasan yang berlipat ganda. Semoga skripsi ini dapat memberikan manfaat bagi perkembangan ilmu pengetahuan dan menjadi amal baik bagi semua pihak yang telah berkontribusi dalam penyusunannya.

Wassalamu'alaikum Wr, Wb.

Malang, 26 Juni 2026

Penulis

DAFTAR ISI

HALAMAN PERSETUJUAN.....	iii
HALAMAN PENGESAHAN.....	iv
PERNYATAAN KEASLIAN TULISAN	v
MOTTO	vi
HALAMAN PERSEMBAHAN.....	vii
KATA PENGANTAR.....	viii
DAFTAR ISI.....	xi
DAFTAR GAMBAR.....	xiii
DAFTAR TABEL.....	xiv
ABSTRAK	xv
ABSTRACT	xvi
مستخلص البحث.....	xvii
BAB I PENDAHULUAN.....	1
1.1 Latar Belakang	1
1.2 Rumusan Masalah	7
1.3 Batasan Masalah	7
1.4 Tujuan Penelitian	7
1.5 Manfaat Penelitian	7
BAB II STUDI PUSTAKA	9
2.1 Penelitian Terdahulu	9
2.2 Kesehatan Mental pada Pekerja <i>Remote</i>	18
2.3 Decision Tree	20
2.4 <i>Random Forest</i>	22
BAB III DESAIN DAN IMPLEMENTASI	28
3.1 Desain Penelitian	28
3.2 Pengumpulan Data	29
3.3 Implementasi Model	34
3.4 Input Data.....	34
3.5 <i>Preprocessing</i>	35
3.5.1 Data Mentah	36
3.5.2 <i>Cleaning Data</i>	36
3.5.3 <i>Label Encoding</i>	37
3.5.4 Normalisasi Data	38
3.6 <i>Feature Selection</i>	40
3.7 Model <i>Random Forest</i>	42
3.7.1 Perhitungan Manual	47

3.8	<i>Classification</i>	50
3.9	Skenario Pengujian	51
3.10	Evaluasi Model	53
3.10.1	Strategi Validasi (<i>K-Fold Cross Validation</i> dan <i>Out-of-Bag Score</i>)...	53
3.10.2	<i>Confusion Matrix</i>	54
BAB IV HASIL DAN PEMBAHASAN.....		58
4.1	Data Penelitian	58
4.2	Hasil <i>Preprocessing</i> Data	59
4.2.1	<i>Cleaning</i> Data.....	60
4.2.2	<i>Label Encoding</i>	61
4.2.3	Normalisasi Data	63
4.3	Hasil <i>Feature Selection</i>	64
4.4	Pembentukan Model <i>Random Forest</i>	66
4.5	Evaluasi Model	68
4.5.1	Evaluasi Model Dengan <i>Stratified K-Fold Cross Validation</i>	69
4.5.2	Evaluasi Model Dengan <i>Out-of-Bag (OOB) Score</i>	70
4.6	Hasil Skenario Uji Coba	71
4.6.1	Alur Pengujian Model	72
4.6.2	Skenario 1.....	73
4.6.3	Skenario 2.....	77
4.6.4	Skenario 3.....	82
4.7	Pembahasan.....	86
4.7.1	Analisis Performa Model	87
4.7.2	Analisis <i>Confusion Matrix</i>	88
4.7.3	Analisis ROC-AUC.....	89
4.7.4	Analisis Pengaruh Pembagian Data dan Jumlah <i>Fold</i>	90
4.8	Perbandingan Kinerja Model	91
BAB V KESIMPULAN DAN SARAN		94
5.1	Kesimpulan	94
5.2	Saran	95
DAFTAR PUSTAKA		

DAFTAR GAMBAR

Gambar 2. 1 Arsitektur <i>Random Forest</i>	24
Gambar 3. 1 Desain Penelitian Alur Proses	29
Gambar 3. 2 Implementasi Model.....	34
Gambar 3. 3 <i>Preprocessing</i>	36
Gambar 3. 4 Flowchart <i>Decision Tree</i>	43
Gambar 3. 5 Flowchart <i>Random Forest</i>	44
Gambar 4. 1 Hasil <i>Feature Importance</i>	65
Gambar 4. 2 Stabilisasi Akurasi Jumlah Pohon Dari Proses Pelatihan Model	67
Gambar 4. 3 <i>Confusion Matrix</i> (60:40 K 5).....	74
Gambar 4. 4 <i>Confusion Matrix</i> (60:40 K 10).....	75
Gambar 4. 5 Kurva ROC(60:40 K 5).....	76
Gambar 4. 6 Kurva ROC(60:40 K 10).....	76
Gambar 4. 7 <i>Confusion Matrix</i> (70:30 K 5).....	79
Gambar 4. 8 <i>Confusion Matrix</i> (70:30 K 10).....	79
Gambar 4. 9 Kurva ROC(70:30 K 5).....	80
Gambar 4. 10 Kurva ROC(70:30 K 10).....	81
Gambar 4. 11 <i>Confusion Matrix</i> (80:20 K 5).....	83
Gambar 4. 12 <i>Confusion Matrix</i> (80:20 K 10).....	84
Gambar 4. 13 Kurva ROC (80:20 K 5).....	85
Gambar 4. 14 Kurva ROC (80:20 K 10).....	85

DAFTAR TABEL

Tabel 2. 1 Penelitian Terdahulu	15
Tabel 3. 1 Detail Atribut Dataset	30
Tabel 3. 2 Sample Dataset.....	33
Tabel 3. 3 Data Sebelum <i>Encoding</i>	37
Tabel 3. 4 Data Setelah <i>Encoding</i>	38
Tabel 3. 5 Hasil Akhir Seleksi Fitur yang Digunakan dalam Pemodelan.....	41
Tabel 3. 6 <i>Tuning Hyperparameter</i>	42
Tabel 3. 7 <i>Confusion Matrix</i>	55
Tabel 4. 1 Hasil <i>Encoding</i>	62
Tabel 4. 2 Hasil Normalisasi Atribut <i>Age</i>	64
Tabel 4. 3 Hasil Evaluasi <i>Stratified K-Fold Cross Validation</i>	69
Tabel 4. 4 Hasil Evaluasi <i>Out-of-Bag (OOB) Score</i>	70
Tabel 4. 5 Hasil Evaluasi Model Skenario 1	73
Tabel 4. 6 Hasil Evaluasi Model Skenario 2.....	77
Tabel 4. 7 Hasil Evaluasi Model Skenario 3.....	82
Tabel 4. 8 Perbandingan Kinerja Model	91

ABSTRAK

Utami, Fauziah Putri. 2026. **Klasifikasi Kebutuhan Konsultasi Kesehatan Mental Pekerja Remote Menggunakan Metode Random Forest**. Skripsi. Program Studi Teknik Informatika Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang. Pembimbing: (I) Prof. Dr. Suhartono S.Si M.Kom. (II) Nurizal Dwi Priandani, M.Kom.

Kata kunci: *Random Forest, Feature Importance, Stratified K-Fold Cross Validation, Klasifikasi, Kesehatan Mental, Pekerja Remote.*

Kesehatan mental pekerja *remote* menjadi salah satu isu yang semakin mendapat perhatian karena sistem kerja jarak jauh berpotensi meningkatkan stres, kecemasan, dan isolasi sosial yang berdampak pada produktivitas kerja. Penelitian ini bertujuan untuk mengklasifikasikan kebutuhan konsultasi kesehatan mental pekerja *remote* menggunakan metode *Random Forest* berdasarkan riwayat *treatment*. Dataset yang digunakan merupakan data survei kesehatan mental pekerja di sektor teknologi yang telah melalui tahapan *data cleaning*, *label encoding*, normalisasi data, dan *feature selection* menggunakan *Feature Importance* sehingga diperoleh atribut yang paling relevan dalam proses klasifikasi. Evaluasi model dilakukan menggunakan *Stratified K-Fold Cross Validation* dengan nilai *k* sebanyak 5 dan 10 pada tiga skenario pembagian data, yaitu 60:40, 70:30, dan 80:20. Hasil penelitian menunjukkan bahwa metode *Random Forest* mampu menghasilkan performa klasifikasi yang baik berdasarkan nilai *accuracy*, *precision*, *recall*, *F1-score*, dan *Area Under Curve* (AUC) pada seluruh skenario pengujian. Skenario pembagian data 60:40 dengan *Stratified K-Fold* 10 dan 5 memberikan hasil evaluasi terbaik dibandingkan skenario lainnya, sehingga menunjukkan bahwa model memiliki kemampuan yang baik dalam membedakan pekerja yang membutuhkan konsultasi kesehatan mental dengan yang tidak membutuhkan konsultasi kesehatan mental. Dengan demikian, metode *Random Forest* dapat digunakan sebagai pendekatan yang efektif untuk mendukung proses klasifikasi kebutuhan konsultasi kesehatan mental pekerja *remote* sehingga dapat membantu pengambilan keputusan secara lebih cepat, objektif, dan akurat.

ABSTRACT

Utami, Fauziah Putri. 2026. **Classification of Remote Workers' Mental Health Consultation Needs Using the Random Forest Method**. Thesis. Informatics Engineering Study Program Faculty of Science and Technology Universitas Islam Negeri Maulana Malik Ibrahim Malang. Advisor: (I) Prof. Dr. Suhartono S.Si M.Kom. (II) Nurizal Dwi Priandani, M.Kom.

Keywords: Random Forest, Feature Importance, Stratified K-Fold Cross Validation, Classification, Mental Health, Remote Workers.

Mental health among remote workers has become an increasingly important issue because remote working arrangements may increase stress, anxiety, and social isolation, which can negatively affect work productivity. This research aims to classify the mental health consultation needs of remote workers using the Random Forest method based on treatment history. The dataset used in this study was obtained from a mental health survey of technology workers and underwent several preprocessing stages, including data cleaning, label encoding, data normalization, and feature selection using Feature Importance to obtain the most relevant attributes for the classification process. Model evaluation was conducted using Stratified K-Fold Cross Validation with k values of 5 and 10 under three data split scenarios, namely 60:40, 70:30, and 80:20. The experimental results indicate that the Random Forest method achieves good classification performance based on accuracy, precision, recall, F1-score, and Area Under the Curve (AUC) across all testing scenarios. The 60:40 data split scenario with 10 and 5-fold Stratified K-Fold Cross Validation produced the best performance among all scenarios, demonstrating that the proposed model can effectively distinguish between remote workers who require mental health consultation and those who do not. Therefore, the Random Forest method can serve as an effective approach for supporting the classification of remote workers' mental health consultation needs, enabling faster, more objective, and accurate decision-making.

مستخلص البحث

أوتامي، فوزية فوتري. ٢٠٢٦. تصنيف الحاجة إلى الاستشارة في الصحة النفسية لدى العاملين عن بُعد باستخدام طريقة الغابة العشوائية. بحث جامعي. قسم الهندسة المعلوماتية، كلية العلوم والتكنولوجيا بجامعة مولانا مالك إبراهيم الإسلامية الحكومية مالانج. المشرف الأول: الأستاذ الدكتور سوهارتونو، بكالوريوس العلوم، ماجستير في علوم الحاسوب. المشرف الثاني: نوريزال دوي بريانداني، ماجستير في علوم الحاسوب.

الكلمات الرئيسية: الغابة العشوائية، أهمية الخصائص، التحقق المتقاطع الطبقي بطريقة ك-الطيات، التصنيف، الصحة النفسية

تُعدّ الصحة النفسية لدى العاملين عن بُعد من القضايا التي تحظى باهتمام متزايد، لأن نظام العمل عن بُعد قد يزيد من مستويات التوتر والقلق والعزلة الاجتماعية، مما يؤثر في إنتاجية العاملين. ويهدف هذا البحث إلى تصنيف الحاجة إلى الاستشارة في واعتمدت (Treatment) الصحة النفسية لدى العاملين عن بُعد باستخدام طريقة الغابة العشوائية اعتماداً على سجل العلاج الدراسة على بيانات مستمدة من استبيان الصحة النفسية للعاملين في قطاع التكنولوجيا، وقد خضعت البيانات لعدة مراحل من المعالجة، شملت تنظيف البيانات، وترميز البيانات، وتطبيع البيانات، واختيار الخصائص باستخدام أهمية الخصائص، وذلك للحصول على أكثر الخصائص ارتباطاً بعملية التصنيف. كما جرى تقييم أداء النموذج باستخدام التحقق المتقاطع الطبقي بطريقة ك-الطيات بقيمتي ٥ و ١٠ ضمن ثلاثة سيناريوهات لتقسيم البيانات، وهي ٦٠:٤٠ و ٧٠:٣٠ و ٨٠:٢٠. وأظهرت نتائج الدراسة أن طريقة والإحكام (Accuracy)، الغابة العشوائية حققت أداءً جيداً في جميع سيناريوهات الاختبار استناداً إلى مقياس الدقة كما أظهر سيناريو تقسيم (AUC) ومساحة ما تحت المنحنى، F1-Score ودرجة (Recall) والاستدعاء (Precision) البيانات ٦٠:٤٠ باستخدام التحقق المتقاطع الطبقي بطريقة ك-الطيات بقيمتي ٥ و ١٠ أفضل نتائج التقييم مقارنةً بالسيناريوهات الأخرى، مما يدل على قدرة النموذج على التمييز بين العاملين الذين يحتاجون إلى الاستشارة في الصحة النفسية والعاملين الذين لا يحتاجون إليها. وبناءً على ذلك، يمكن اعتماد طريقة الغابة العشوائية بوصفها أسلوباً فعالاً لدعم عملية تصنيف الحاجة إلى الاستشارة في الصحة النفسية لدى العاملين عن بُعد، بما يساهم في اتخاذ قرارات أكثر سرعة وموضوعية ودقة.

BAB I

PENDAHULUAN

1.1 Latar Belakang

Kesehatan mental merupakan salah satu isu kesehatan global yang mendapat perhatian serius dari Organisasi Kesehatan Dunia. Menurut laporan *World Mental Health Report* (WHO) pada *Transforming Mental Health for All* tahun 2022, lebih dari 970 juta orang di seluruh dunia hidup dengan gangguan kesehatan mental, dengan depresi dan kecemasan sebagai penyebab utama disabilitas global. Selain itu juga tercatat, lebih dari 700.000 orang meninggal akibat bunuh diri setiap tahun. Hal ini menjadikan penyebab utama kematian di kalangan usia produktif. Kondisi ini menunjukkan bahwa kesehatan mental tidak hanya berdampak pada individu, tetapi juga menimbulkan beban sosial dan ekonomi yang besar bagi masyarakat, sehingga deteksi dini dan intervensi yang tepat sangat diperlukan untuk mencegah dampak yang lebih serius (WHO, 2022).

Secara global, WHO memperkirakan bahwa satu dari empat orang akan mengalami gangguan kesehatan mental sepanjang hidupnya. Kondisi ini menegaskan bahwa gangguan mental telah menjadi masalah kesehatan utama yang sejajar dengan penyakit fisik kronis seperti kanker dan penyakit jantung. Dampak gangguan kesehatan mental tidak hanya dirasakan oleh individu, melainkan juga berimplikasi pada kehidupan sosial dan ekonomi. Sementara itu, Nurdiansyah et al. (2025) menambahkan bahwa kondisi psikologis yang terganggu dapat memicu masalah sosial lain seperti penurunan kualitas hubungan interpersonal, meningkatnya risiko konflik, bahkan kriminalitas. Fakta ini

memperkuat urgensi kesehatan mental sebagai isu publik yang harus ditangani secara sistematis.

Dalam konteks dunia kerja modern, isu kesehatan mental menjadi semakin kompleks dengan munculnya tren kerja jarak jauh atau *remote working*. Perubahan pola kerja ini semakin populer setelah pandemi COVID-19 yang mendorong banyak perusahaan beradaptasi dengan sistem digital. Secara teori, *remote working* memberikan fleksibilitas waktu, efisiensi biaya, serta kesempatan untuk bekerja lintas wilayah. Namun, kenyataannya pola kerja ini juga memunculkan berbagai permasalahan baru. Pekerja *remote* kerap menghadapi isolasi sosial, kesulitan menjaga keseimbangan antara pekerjaan dan kehidupan pribadi, serta keterbatasan akses terhadap dukungan kesehatan mental. Ropponen (2025) menyebutkan bahwa pekerja jarak jauh lebih rentan terhadap stres, *burnout*, dan depresi dibanding pekerja kantoran. Wu et al. (2024) menemukan bahwa lebih dari 70% pekerja *remote* melaporkan mengalami gejala psikologis berupa kecemasan, depresi, dan kelelahan emosional. Hal ini juga ditegaskan oleh Khan et al. (2025) yang menunjukkan bahwa pekerja *remote* menghadapi risiko tinggi terhadap gangguan psikologis yang berdampak langsung pada kualitas hidup maupun produktivitas kerja.

Minimnya dukungan psikososial dari perusahaan memperburuk kondisi tersebut. Febrian et al. (2025) menegaskan bahwa lingkungan kerja yang tidak mendukung kesehatan mental dapat memperbesar risiko gangguan psikologis pada pekerja. Choirunnisa (2025) menambahkan bahwa stigma terhadap masalah psikologis membuat banyak pekerja enggan mencari bantuan profesional

meskipun sudah merasakan gejala depresi atau kecemasan. Situasi ini menunjukkan perlunya sistem deteksi dini yang dapat membantu mengidentifikasi pekerja *remote* yang membutuhkan layanan konsultasi kesehatan mental sebelum kondisi mereka semakin memburuk.

Dalam perspektif Islam, menjaga kesehatan mental sama pentingnya dengan menjaga kesehatan fisik. Al-Qur'an dalam Surah Asy-Syu'ara ayat 80.

وَإِذَا مَرِضْتُ فَهُوَ يَشْفِينِ ﴿٨٠﴾

“Dan apabila aku sakit, Dialah yang menyembuhkanku.” (QS. Asy-Syu'ara-80).

Ayat tersebut menegaskan bahwa Allah adalah penyembuh segala penyakit, baik fisik maupun psikis, namun manusia tetap diwajibkan untuk berusaha mencari jalan kesembuhan. Dalam konteks penelitian ini, ikhtiar tersebut diwujudkan melalui pemanfaatan teknologi berbasis *machine learning*, khususnya algoritma *Random Forest*, untuk memprediksi kebutuhan konsultasi kesehatan mental mahasiswa. Upaya ini merupakan bentuk usaha manusia dalam menjaga amanah kesehatan jiwa yang dianugerahkan Allah, sekaligus wujud rahmat yang dapat mencegah kerugian lebih besar bagi individu maupun lingkungan kerja. Dengan demikian, penerapan teknologi untuk prediksi kesehatan mental sejalan dengan prinsip Islam dalam memelihara kehidupan dan mencegah kemudarat.

Seiring dengan perkembangan teknologi, upaya deteksi dini kesehatan mental kini dapat dilakukan dengan memanfaatkan algoritma *machine learning*. Salah satu algoritma yang efektif dan banyak digunakan adalah *Random Forest*. Algoritma ini merupakan metode *ensemble learning* yang membangun banyak pohon keputusan (*decision tree*) dan menggabungkan hasilnya untuk memperoleh

prediksi yang lebih akurat. Faisti et al. (2025) menjelaskan bahwa *Random Forest* memiliki keunggulan dalam mengatasi masalah *overfitting*, bekerja dengan baik pada dataset berukuran besar, serta mampu memberikan interpretasi mengenai fitur yang paling berpengaruh dalam klasifikasi. Choirunnisa (2025) menambahkan bahwa evaluasi *Random Forest* pada klasifikasi kesehatan mental menunjukkan performa yang stabil bahkan pada data yang heterogen. Penelitian Sebayang et al. (2023) membuktikan bahwa penggunaan *Random Forest* pada data pekerja teknologi menghasilkan akurasi mencapai 84%. Sun (2024) juga melaporkan bahwa algoritma ini mampu memberikan rata-rata *akurasi*, *presisi*, *recall*, dan *f1-score* di atas 80%.

Penelitian ini menggunakan dataset OSMI *Mental Health in Tech Survey* 2016, yang dikembangkan oleh *Open Sourcing Mental Illness* (OSMI) (OSMI, 2016). Dataset ini berisi hasil survei ribuan pekerja teknologi, termasuk pekerja *remote*, dengan variabel yang relevan untuk analisis kesehatan mental. Beberapa atribut penting di dalamnya meliputi riwayat kesehatan mental keluarga (*family_history*), pengalaman perawatan (*treatment*), status kerja jarak jauh (*remote_work*), dukungan perusahaan seperti tunjangan kesehatan (*benefits*), opsi layanan psikologis (*care_options*), serta persepsi pekerja terhadap stigma (*mental_health_consequence*). Sebayang et al. (2023) menilai dataset OSMI sangat representatif untuk riset kesehatan mental di sektor teknologi. Priyono et al. (2024) juga menegaskan bahwa keragaman variabel dalam dataset ini memungkinkan penerapan metode *machine learning* untuk klasifikasi kebutuhan konsultasi kesehatan mental dengan tingkat akurasi yang tinggi.

Meskipun penelitian mengenai klasifikasi kesehatan mental dengan algoritma *machine learning* telah banyak dilakukan, masih terdapat celah penelitian yang belum banyak disentuh. Sebagian besar penelitian terdahulu berfokus pada populasi umum atau pekerja teknologi secara keseluruhan. Sebayang et al. (2023) mengembangkan klasifikasi kesehatan mental berbasis *Random Forest* pada pekerja industri teknologi, tetapi penelitian tersebut belum mengidentifikasi tantangan khusus pekerja jarak jauh. Priyono et al. (2024) juga mengusulkan penerapan *Random Forest* untuk diagnosis kesehatan mental, namun konteks yang diteliti lebih luas dan tidak spesifik pada model kerja *remote*. Penelitian lain seperti yang dilakukan oleh Wijaya (2024) dan Kunhayati (2025) lebih menekankan pada penggunaan teknik penyeimbangan data seperti SMOTE untuk mengatasi *imbalanced dataset*. Pendekatan tersebut memang dapat meningkatkan akurasi, tetapi berpotensi tidak sepenuhnya menggambarkan kondisi nyata distribusi data di lapangan. Celah inilah yang menjadi dasar penelitian ini, yaitu fokus pada objek penelitian yaitu pekerja *remote* dengan penggunaan algoritma *Random Forest* tanpa penyeimbangan data untuk memperoleh hasil yang lebih representatif terhadap kondisi sebenarnya.

Penelitian ini memiliki keunikan yang membedakannya dari penelitian terdahulu. Fokus penelitian diarahkan pada pekerja *remote* yang hingga saat ini masih jarang dijadikan subjek kajian, padahal jumlahnya semakin meningkat dan menghadapi tantangan psikologis berbeda dari pekerja kantoran. Dataset OSMI yang digunakan dalam penelitian ini kaya akan variabel penting sehingga hasil klasifikasi lebih relevan untuk menggambarkan kondisi kesehatan mental pekerja

teknologi. Pendekatan yang dipilih juga menekankan pada evaluasi algoritma *Random Forest* pada data yang tidak diseimbangkan, sehingga hasil analisis lebih realistis. Dengan demikian, penelitian ini berpotensi memberikan kontribusi nyata baik dalam ranah akademik maupun praktik.

Urgensi penelitian ini didasarkan pada meningkatnya prevalensi gangguan kesehatan mental di kalangan pekerja remote. Ropponen (2025) menjelaskan bahwa pekerja jarak jauh lebih rentan mengalami stres, kecemasan, dan isolasi sosial yang berdampak pada *burnout* serta niat untuk mengundurkan diri. Hendra et al. (2025) menambahkan bahwa kondisi psikologis yang terganggu dapat menurunkan kinerja, meningkatkan absensi, serta menambah beban biaya kesehatan bagi perusahaan. Wells et al. (2023) menemukan bahwa penerapan sistem kerja jarak jauh memiliki dampak yang beragam terhadap kesehatan mental pekerja, baik berupa peningkatan fleksibilitas maupun peningkatan stres dan isolasi sosial. Kristman et al. (2022) menyatakan bahwa hubungan antara kerja remote dan kesejahteraan mental pekerja masih menunjukkan hasil yang belum konsisten karena dipengaruhi oleh berbagai faktor individu dan lingkungan kerja.

Meskipun berbagai penelitian telah mengkaji hubungan antara sistem kerja dan kesehatan mental, sebagian besar penelitian masih berfokus pada pekerja secara umum tanpa membedakan karakteristik pekerja remote dan pekerja non-remote secara spesifik. Wells et al. (2023) lebih menitikberatkan pada dampak kerja jarak jauh terhadap kesehatan dan keselamatan kerja, sedangkan Kristman et al. (2022) berfokus pada kesejahteraan mental pekerja secara umum tanpa melakukan analisis klasifikasi berdasarkan karakteristik pekerja remote dan non-

remote. Selain itu, penelitian yang menerapkan algoritma *machine learning* untuk melakukan klasifikasi kondisi kesehatan mental pada pekerja remote masih relatif terbatas. Oleh karena itu, penelitian ini dilakukan untuk mengisi kesenjangan penelitian tersebut dengan menerapkan algoritma *Random Forest* dalam mengklasifikasikan kondisi kesehatan mental pada pekerja *remote* berdasarkan riwayat *treatment*.

1.2 Rumusan Masalah

1. Bagaimana akurasi tiap model *Random Forest* dalam mengklasifikasi kebutuhan konsultasi kesehatan mental pekerja *remote*?

1.3 Batasan Masalah

1. Penelitian ini berupa hasil survei kesehatan mental pekerja pada bidang teknologi tahun 2016.

1.4 Tujuan Penelitian

Untuk mengetahui performa *Random Forest* dalam mengklasifikasi kebutuhan konsultasi kesehatan mental pada mahasiswa, dengan menggunakan dataset *OSMI Tech in Survey 2016*.

1.5 Manfaat Penelitian

1. Bagi Perusahaan Teknologi, model *Random Forest* dapat digunakan sebagai alat bantu deteksi kebutuhan konsultasi kesehatan mental pekerja *remote*, sehingga mendukung intervensi dini untuk menurunkan risiko *burnout*.

2. Bagi Pekerja *Remote*, model ini dapat membantu mengenali kondisi psikologis secara lebih cepat sehingga mendorong kesadaran untuk mencari dukungan profesional bila diperlukan.
3. Bagi Akademisi dan Peneliti, penelitian ini dapat dijadikan referensi dalam pengembangan metode klasifikasi kesehatan mental berbasis *machine learning*.

BAB II

STUDI PUSTAKA

2.1 Penelitian Terdahulu

Penelitian mengenai klasifikasi kesehatan mental dengan pendekatan machine learning telah berkembang pesat dalam lima tahun terakhir. Berbagai algoritma digunakan untuk mengatasi kompleksitas data kesehatan mental yang melibatkan banyak variabel psikologis, sosial, dan lingkungan. Penelitian terdahulu menjadi landasan penting untuk memahami tren, kelebihan, serta keterbatasan metode yang digunakan. Dalam konteks penelitian ini, yang berfokus pada prediksi kebutuhan konsultasi kesehatan mental pekerja *remote* menggunakan *Random Forest* tanpa teknik penyeimbangan data pada dataset OSMI Tech in Survey 2016, penelitian-penelitian ini memberikan wawasan tentang performa algoritma pada data imbalanced yang merepresentasikan kondisi nyata.

Priyono et al. (2024) melakukan klasifikasi diagnosis penyakit mental menggunakan metode *Random Forest* dengan data survei kesehatan mental. Dalam penelitian mereka, model mencapai akurasi sebesar 90.83% tanpa menerapkan teknik penanganan ketidakseimbangan data (imbalanced data). Penelitian ini menekankan kemampuan *Random Forest* dalam menangani variabel survei psikologis seperti gejala depresi dan kecemasan, yang sering kali tidak seimbang secara alami. Hasilnya menunjukkan bahwa algoritma ini stabil meskipun tanpa modifikasi data, sehingga mengurangi risiko overfitting pada kelas minoritas. Hubungan dengan masalah penelitian ini terletak pada

pendekatan tanpa balancing, yang sejalan dengan tujuan penelitian untuk mengevaluasi performa *Random Forest* pada dataset OSMI yang imbalanced, di mana kelas pekerja yang membutuhkan konsultasi lebih sedikit. Pendekatan ini memungkinkan penelitian ini untuk merefleksikan kondisi real-world pekerja *remote* di sektor teknologi, di mana data survei sering kali tidak seimbang dan memerlukan model yang robust tanpa intervensi artifisial.

Fahmi et al. (2025) meningkatkan klasifikasi kesehatan mental pekerja IT *remote* dengan *Random Forest* dan teknik feature engineering. Algoritma mereka berhasil mencapai akurasi hingga 89% tanpa manipulasi distribusi data, dengan fokus pada variabel seperti isolasi sosial, beban kerja, dan dukungan perusahaan. Penelitian ini melibatkan dataset survei serupa dengan OSMI, di mana fitur demografis dan psikososial diekstrak untuk meningkatkan interpretabilitas model. Keunggulan utama adalah kemampuan *Random Forest* dalam menangani multicollinearity antar variabel tanpa perlu balancing, sehingga hasil prediksi lebih akurat pada populasi pekerja *remote*. Hubungan dengan penelitian ini adalah pada konteks pekerja IT *remote*, yang mirip dengan subjek dataset OSMI. Dengan tidak menggunakan balancing, penelitian Fahmi et al. mendukung hipotesis penelitian ini bahwa *Random Forest* dapat memberikan performa yang baik pada data imbalanced, sehingga membantu dalam memprediksi kebutuhan konsultasi tanpa mengubah distribusi data asli, yang krusial untuk aplikasi praktis di perusahaan teknologi.

Ndikumana et al. (2025) menggunakan *Random Forest* untuk mengklasifikasi risiko kesehatan mental pada remaja di Rwanda berdasarkan

survei lintas sektoral. Model mereka mencapai akurasi 88.8% tanpa teknik balancing khusus, dengan variabel utama meliputi faktor lingkungan, riwayat keluarga, dan dukungan sosial. Penelitian ini menyoroti bagaimana *Random Forest* efektif dalam dataset survei yang heterogen dan imbalanced, di mana kelas risiko tinggi sering kali minoritas. Hasil evaluasi menggunakan metrik seperti precision dan recall menunjukkan kestabilan model meskipun data tidak seimbang. Hubungan dengan masalah penelitian ini terletak pada penggunaan survei sebagai sumber data, mirip dengan OSMI Tech in Survey 2016. Penelitian Ndikumana et al. memperkuat pendekatan penelitian ini untuk tidak menggunakan balancing, karena hal itu memungkinkan evaluasi performa *Random Forest* yang lebih autentik pada data pekerja *remote*, di mana faktor seperti isolasi sosial (seperti pada remaja Rwanda) dapat analog dengan tantangan kerja jarak jauh, sehingga meningkatkan relevansi klasifikasi kebutuhan konsultasi.

Retnaswamy et al. (2025) membandingkan *Random Forest* dengan Decision Tree dan SVM pada prediksi kebutuhan perawatan kesehatan mental. *Random Forest* memberikan hasil akurasi 81.2% dan AUC 0.813 pada dataset pekerja teknologi tanpa modifikasi data imbalanced. Penelitian ini melibatkan variabel seperti treatment history, *remote* work status, dan mental health consequences, yang dievaluasi tanpa oversampling untuk menjaga integritas data. Hasil perbandingan menunjukkan superioritas *Random Forest* dalam menangani noise dan variabel kategorikal. Hubungan dengan penelitian ini adalah pada fokus prediksi kebutuhan perawatan, yang langsung berkaitan dengan target klasifikasi biner (membutuhkan/tidak membutuhkan konsultasi) pada pekerja *remote*.

Dengan tidak menggunakan balancing, penelitian Retnaswamy et al. mendukung evaluasi performa murni *Random Forest* pada dataset OSMI yang imbalanced, membantu penelitian ini dalam mengukur metrik seperti F1-score untuk kelas minoritas, yang sering kali mewakili pekerja dengan risiko tinggi di lingkungan *remote*.

Choirunnisa (2025) mengembangkan model *Random Forest* yang akurat untuk identifikasi gangguan mental pada populasi besar, dengan hasil mendekati state-of-the-art tanpa menggunakan data balancing. Penelitian ini menggunakan dataset survei luas, di mana akurasi dicapai melalui tuning hyperparameter seperti jumlah tree dan depth, tanpa mengubah distribusi kelas. Fokus pada interpretasi feature importance menunjukkan variabel seperti family history dan care options sebagai prediktor utama. Hubungan dengan masalah penelitian ini terletak pada identifikasi gangguan mental pada populasi besar, yang analog dengan pekerja teknologi di dataset OSMI. Pendekatan tanpa balancing memperkuat tujuan penelitian ini untuk mengeksplorasi performa *Random Forest* pada data nyata imbalanced, sehingga dapat diterapkan pada pekerja *remote* untuk mendeteksi kebutuhan konsultasi secara dini tanpa bias artifisial.

Gamage et al. (2025) mengadopsi machine learning untuk mengklasifikasi tingkat stres kerja pada pekerja *remote* dengan *Random Forest*, menunjukkan peningkatan prediktabilitas dengan variabel psikososial dan demografis tanpa teknik balancing data. Penelitian ini mencapai akurasi tinggi pada dataset survei, dengan evaluasi metrik seperti recall untuk kelas stres tinggi yang minoritas. Hasilnya menekankan kemampuan *Random Forest* dalam menangani data tidak

seimbang melalui mekanisme ensemble. Hubungan dengan penelitian ini adalah pada konteks stres kerja pekerja *remote*, yang langsung relevan dengan isu kesehatan mental di sektor teknologi. Dengan tidak menggunakan balancing, penelitian Gamage et al. mendukung pendekatan penelitian ini untuk menilai performa *Random Forest* pada dataset OSMI, di mana variabel seperti *remote_work* dan *benefits* dapat diklasifikasi tanpa modifikasi, sehingga meningkatkan keakuratan deteksi kebutuhan konsultasi.

Kannan et al. (2024) menyajikan tinjauan terbaru mengenai kemajuan machine learning dan deep learning dalam diagnosis dan manajemen gangguan mental. Mereka menggarisbawahi pentingnya integrasi data biomarker, genomik, dan perilaku melalui model prediktif canggih seperti *Random Forest*, yang mampu mendeteksi gangguan mental lebih dini dibanding metode tradisional. Tinjauan ini mencakup kasus di mana *Random Forest* diterapkan pada data imbalanced tanpa balancing untuk menjaga validitas klinis. Hubungan dengan masalah penelitian ini terletak pada deteksi dini gangguan mental, yang sejalan dengan tujuan prediksi kebutuhan konsultasi pada pekerja *remote*. Tinjauan Kannan et al. memperkuat penggunaan *Random Forest* tanpa balancing pada dataset OSMI, karena memungkinkan integrasi variabel perilaku seperti *mental_health_consequence* untuk prediksi yang lebih akurat dan praktis di lingkungan kerja jarak jauh.

Rosenberg et al. (2021) mengevaluasi algoritma *Random Forest* dan decision tree sebagai teknik terbaik untuk memprediksi kesejahteraan mental berdasarkan data pribadi dan perilaku. Hasilnya menguatkan efektivitas *Random*

Forest dalam menangani dataset besar dan beragam untuk prediksi personalisasi kesehatan mental, dengan akurasi stabil meskipun data imbalanced. Penelitian ini menggunakan metrik seperti precision untuk menilai kelas minoritas. Hubungan dengan penelitian ini adalah pada prediksi personalisasi berdasarkan data pribadi, yang mirip dengan variabel demografis di OSMI seperti age dan gender. Pendekatan tanpa balancing mendukung evaluasi performa murni di penelitian ini, sehingga dapat diterapkan pada pekerja *remote* untuk mengidentifikasi individu yang membutuhkan konsultasi tanpa mengubah data asli.

Febrian et al. (2025) melakukan klasifikasi kesehatan mental remaja berdasarkan faktor lingkungan sekolah dengan menggunakan algoritma *Random Forest*. Penelitian ini mengumpulkan data dari 229 remaja dan mengkategorikan kondisi mental dalam 4 kelas, dengan akurasi model mencapai 80.43% dan F1-score tertinggi pada kategori tanpa gangguan mental. Faktor-faktor seperti rasa kesepian, tekanan akademik, dan stres keluarga menjadi prediktor utama, dievaluasi tanpa balancing untuk merefleksikan distribusi nyata. Hubungan dengan masalah penelitian ini terletak pada penggunaan faktor lingkungan sebagai prediktor, yang analog dengan variabel seperti coworkers dan supervisor di dataset OSMI untuk pekerja *remote*. Penelitian Febrian et al. memperkuat pendekatan tanpa balancing di penelitian ini, karena menunjukkan kemampuan *Random Forest* dalam mendeteksi kondisi mental berdasarkan faktor eksternal, yang dapat disesuaikan untuk konteks kerja jarak jauh guna mengklasifikasi kebutuhan konsultasi secara efektif.

Kunhayati (2025), dalam skripsinya yang berjudul "Klasifikasi Kesehatan Mental pada Imbalanced Data Menggunakan *Random Forest* dan SMOTE", mengeksplorasi performa *Random Forest* pada data kesehatan mental yang tidak seimbang. Penelitian ini menggunakan dataset survei serupa, dengan fokus pada kombinasi *Random Forest* dan SMOTE untuk meningkatkan akurasi pada kelas minoritas. Namun, bagian evaluasi tanpa SMOTE menunjukkan akurasi dasar yang stabil, meskipun recall untuk kelas gangguan mental lebih rendah. Skripsi ini mencakup tahap preprocessing, model building, dan evaluasi metrik komprehensif. Hubungan dengan penelitian ini adalah pada penanganan data imbalanced pada kesehatan mental, yang langsung relevan dengan dataset OSMI. Meskipun Kunhayati menggunakan SMOTE sebagai opsi, evaluasi tanpa balancing mendukung tujuan penelitian ini untuk menilai performa murni *Random Forest*, sehingga memberikan landasan untuk mengukur akurasi pada pekerja *remote* tanpa modifikasi data, yang lebih mencerminkan kondisi nyata di sektor teknologi.

Tabel 2.1 Penelitian Terdahulu

No	Nama Peneliti	Judul Penelitian	Metode	Hasil Penelitian
1	Priyono et al. (2024)	<i>Klasifikasi Diagnosis Penyakit Mental Menggunakan Metode Random Forest dengan Data Survei Kesehatan Mental</i>	<i>Random Forest</i>	Model mencapai akurasi 90,83% tanpa teknik penyeimbangan data. <i>Random Forest</i> terbukti stabil dalam menangani variabel psikologis yang tidak seimbang, seperti gejala depresi dan kecemasan, serta mampu mencerminkan kondisi nyata pada data survei psikologis.
2	Fahmi et al. (2025)	<i>Peningkatan Prediksi Kesehatan Mental Pekerja IT Remote Menggunakan Random Forest dan Feature Engineering</i>	<i>Random Forest, Feature Engineering</i>	Menghasilkan akurasi 89% tanpa manipulasi distribusi data. Penelitian ini menyoroti kemampuan <i>Random Forest</i> dalam menangani variabel psikososial seperti isolasi sosial dan dukungan perusahaan tanpa

No	Nama Peneliti	Judul Penelitian	Metode	Hasil Penelitian
				balancing data, relevan untuk konteks pekerja <i>remote</i> .
3	Ndikumana et al. (2025)	<i>Mental Health Prediction among Adolescents Using Random Forest</i>	<i>Random Forest</i>	Model mencapai akurasi 88,8% tanpa balancing dengan variabel utama seperti dukungan sosial dan faktor lingkungan. <i>Random Forest</i> terbukti efektif pada data survei yang heterogen dan imbalanced, mirip dengan karakteristik dataset OSMI.
4	Retnaswamy et al. (2025)	<i>Comparative Analysis of Random Forest, Decision Tree, and SVM for Mental Health Treatment Prediction</i>	<i>Random Forest, Decision Tree, SVM</i>	<i>Random Forest</i> memperoleh akurasi 81,2% dan AUC 0,813 tanpa oversampling. Model ini unggul dalam prediksi kebutuhan perawatan kesehatan mental dan tahan terhadap noise pada data pekerja teknologi.
5	Choirunnisa (2025)	<i>Random Forest for Large-Scale Mental Disorder Identification</i>	<i>Random Forest</i>	Model mencapai akurasi tinggi dengan tuning parameter tanpa balancing data. Fitur seperti <i>family history</i> dan <i>care options</i> terbukti menjadi prediktor utama gangguan mental, mendukung pendekatan penelitian ini.
6	Gamage et al. (2025)	<i>Predicting Work Stress among Remote Employees Using Random Forest</i>	<i>Random Forest</i>	<i>Random Forest</i> menunjukkan performa tinggi dalam memprediksi tingkat stres kerja pada pekerja <i>remote</i> , dengan evaluasi recall yang baik untuk kelas minoritas tanpa balancing.
7	Kannan et al. (2024)	<i>Advances in Machine Learning for Mental Disorder Diagnosis and Management</i>	<i>Random Forest, Deep Learning Review</i>	Menekankan efektivitas <i>Random Forest</i> pada data imbalanced tanpa balancing untuk menjaga validitas klinis. Studi ini mendukung integrasi variabel perilaku dalam model prediksi gangguan mental.
8	Rosenberg et al. (2021)	<i>Evaluation of Random Forest and Decision Tree in Predicting Mental Well-being</i>	<i>Random Forest, Decision Tree</i>	<i>Random Forest</i> memberikan akurasi stabil meskipun data besar dan imbalanced, membuktikan keunggulan metode ini dalam prediksi kesejahteraan mental berbasis data personal.
9	Febrian et al. (2025)	<i>Prediksi Kesehatan Mental Remaja Berdasarkan Faktor Lingkungan Sekolah Menggunakan Random Forest</i>	<i>Random Forest</i>	Model mencapai akurasi 80,43% dengan variabel lingkungan seperti tekanan akademik dan dukungan sosial sebagai prediktor utama. Pendekatan tanpa balancing mencerminkan distribusi nyata data psikologis.
10	Kunhayati (2025)	<i>Klasifikasi Kesehatan Mental pada Imbalanced</i>	<i>Random Forest, SMOTE</i>	Model <i>Random Forest</i> tetap stabil tanpa SMOTE dengan akurasi tinggi, namun recall untuk kelas

No	Nama Peneliti	Judul Penelitian	Metode	Hasil Penelitian
		<i>Data Menggunakan Random Forest dan SMOTE</i>		minoritas lebih rendah. Penelitian ini mendukung evaluasi murni <i>Random Forest</i> tanpa <i>balancing</i> .
11	Fauziah Putri Utami (skripsi 2025)	<i>Klasifikasi Kebutuhan Konsultasi Kesehatan Mental Pekerja Remote Menggunakan Metode Random Forest</i>	<i>Random Forest</i>	Penelitian ini dirancang untuk mengevaluasi performa algoritma <i>Random Forest</i> dalam mengklasifikasi kebutuhan konsultasi kesehatan mental pada pekerja <i>remote</i> dengan memanfaatkan dataset OSMI Tech in Survey 2016. Model diharapkan mampu memberikan hasil klasifikasi yang stabil dan akurat berdasarkan uji <i>k-fold cross validation</i> tanpa penerapan teknik <i>balancing</i> data, sehingga hasilnya dapat merepresentasikan kondisi nyata distribusi data kesehatan mental pekerja <i>remote</i> .

Tabel 2.1 menyajikan ringkasan penelitian terdahulu yang menggunakan berbagai metode dalam memprediksi kondisi kesehatan mental dengan pendekatan *machine learning*. Beberapa penelitian menunjukkan bahwa algoritma *Random Forest* memiliki performa yang unggul dalam mendeteksi serta mengklasifikasi kebutuhan konsultasi kesehatan mental pada berbagai kelompok, termasuk pekerja di sektor teknologi. Hasil dari penelitian-penelitian sebelumnya membuktikan bahwa *Random Forest* mampu mengolah data kompleks tanpa memerlukan teknik penyeimbangan data (*balancing*), sehingga memberikan hasil prediksi yang representatif terhadap kondisi nyata. Temuan tersebut menjadi landasan bagi penelitian ini yang menerapkan *Random Forest* pada dataset OSMI Tech in Survey 2016 untuk mengklasifikasi kebutuhan konsultasi kesehatan mental pekerja *remote* secara akurat dan realistis sesuai dengan distribusi data sebenarnya.

Pembaruan penelitian ini terletak pada objek yang digunakan, yaitu pekerja *remote*. Fokus utama dalam penelitian ini adalah mengenai klasifikasi konsultasi kesehatan mental. Sebagian besar penelitian terdahulu hanya membahas populasi umum atau pekerja teknologi secara keseluruhan tanpa menyoroti tantangan psikologis khas yang dihadapi oleh pekerja jarak jauh. Selain dari sisi objek, pembaruan juga tampak pada pendekatan analisis yang menekankan evaluasi performa algoritma *Random Forest* tanpa penerapan teknik *balancing* terhadap dataset OSMI Tech in Survey 2016. Pendekatan ini memberikan kontribusi baru karena mempertahankan distribusi alami data, sehingga hasil prediksi kebutuhan konsultasi kesehatan mental lebih realistis dan mencerminkan kondisi sesungguhnya di lingkungan kerja *remote* berbasis teknologi.

2.2 Kesehatan Mental pada Pekerja *Remote*

Kesehatan mental merupakan aspek fundamental dalam kesejahteraan individu yang memengaruhi cara berpikir, merasa, dan berperilaku dalam kehidupan sehari-hari. Menurut World Health Organization (2022), kesehatan mental adalah kondisi sejahtera ketika individu mampu menyadari potensinya, mengelola tekanan hidup secara wajar, bekerja produktif, serta berkontribusi bagi masyarakat. Dalam konteks pekerjaan modern, kondisi mental karyawan berpengaruh signifikan terhadap produktivitas dan kualitas kerja. Gangguan mental seperti stres, kecemasan, atau depresi dapat menurunkan performa dan menghambat pencapaian organisasi (Maramis, 2022).

Perubahan pola kerja akibat kemajuan teknologi dan pandemi global mendorong banyak perusahaan menerapkan sistem kerja jarak jauh (*remote working*). Meskipun sistem ini memberikan fleksibilitas waktu dan tempat, beberapa studi menunjukkan bahwa pekerja *remote* lebih berisiko mengalami gangguan psikologis. Ropponen (2025) melaporkan bahwa pekerja jarak jauh cenderung menghadapi kesulitan dalam memisahkan kehidupan pribadi dan pekerjaan, sehingga meningkatkan tekanan emosional dan risiko *burnout*. Temuan ini sejalan dengan penelitian Wu et al. (2024), yang menyebutkan bahwa 71% pekerja *remote* melaporkan peningkatan kecemasan dan penurunan motivasi akibat isolasi sosial dan kurangnya interaksi tatap muka.

Faktor sosial dan lingkungan kerja juga berperan penting dalam menjaga stabilitas mental pekerja *remote*. Menurut Khan et al. (2025), dukungan sosial dari rekan kerja serta komunikasi terbuka dengan atasan dapat menurunkan risiko stres dan meningkatkan kepuasan kerja. Tanpa adanya dukungan tersebut, pekerja berpotensi mengalami kelelahan emosional yang berkepanjangan. Oleh karena itu, perusahaan di era digital perlu menyesuaikan kebijakan mereka agar lebih ramah terhadap kesejahteraan psikologis, termasuk melalui penyediaan akses konsultasi mental daring atau pelatihan manajemen stres (Febrian et al., 2025).

Selain dampak individual, kesehatan mental pekerja *remote* juga berdampak langsung terhadap kinerja perusahaan. Studi oleh Febrian et al. (2025) menemukan bahwa organisasi yang menerapkan kebijakan dukungan psikologis mengalami peningkatan produktivitas karyawan hingga 23%. Hasil ini menunjukkan pentingnya deteksi dini terhadap kondisi mental karyawan sebagai

bagian dari strategi keberlanjutan organisasi. Oleh sebab itu, penelitian ini memanfaatkan pendekatan *machine learning* untuk mengklasifikasi kebutuhan konsultasi kesehatan mental pekerja *remote* agar perusahaan dapat melakukan intervensi secara tepat waktu.

Dengan demikian, pendekatan berbasis data menjadi solusi efektif dalam memahami pola risiko kesehatan mental di lingkungan kerja jarak jauh. Melalui penerapan algoritma *Random Forest*, penelitian ini bertujuan menghasilkan model prediksi yang akurat dalam mengidentifikasi individu yang berpotensi membutuhkan bantuan psikologis. Hasil klasifikasi ini diharapkan dapat menjadi dasar bagi perusahaan dalam merancang kebijakan kerja yang lebih adaptif, inklusif, dan mendukung kesejahteraan mental pekerja.

2.3 Decision Tree

Decision Tree merupakan algoritma pembelajaran terawasi (*supervised learning*) yang digunakan dalam analisis klasifikasi. Model ini berbentuk struktur pohon di mana setiap simpul (*node*) mewakili atribut atau fitur yang diuji, cabang menunjukkan hasil dari pengujian, dan daun (*leaf node*) menandakan hasil akhir keputusan (Han et al., 2012). Dengan representasi visual yang intuitif, *Decision Tree* memudahkan pengguna untuk memahami bagaimana keputusan diambil berdasarkan data.

Proses kerja *Decision Tree* melibatkan pembagian data (*splitting*) secara rekursif berdasarkan atribut yang memberikan *information gain* tertinggi. Kriteria seperti *Entropy*, *Gini Index*, atau *Gain Ratio* digunakan untuk menentukan pemisahan terbaik. Kusumarini et al. (2021) menjelaskan bahwa semakin besar

information gain suatu atribut, semakin baik atribut tersebut dalam memisahkan kelas data. Model ini kemudian menghasilkan aturan (*rules*) yang dapat digunakan untuk melakukan prediksi terhadap data baru dengan tingkat interpretabilitas yang tinggi.

Kelebihan utama *Decision Tree* adalah kemampuannya menangani data kategorikal maupun numerik tanpa memerlukan normalisasi. Namun, algoritma ini memiliki kecenderungan *overfitting* ketika pohon terlalu dalam atau jumlah fitur terlalu banyak (Marsuhandi et al., 2020). Oleh karena itu, diperlukan proses pemangkasan (*pruning*) agar model lebih general dan stabil terhadap data baru. Dalam penelitian ini, *Decision Tree* digunakan untuk mengidentifikasi faktor-faktor utama yang berkontribusi terhadap kebutuhan konsultasi kesehatan mental pekerja *remote*, seperti *family_history*, *benefits*, dan *work_interfere*.

Selain berdiri sendiri, *Decision Tree* juga menjadi komponen inti dalam metode *ensemble learning* seperti *Random Forest*. Retnaswamy et al. (2025) menjelaskan bahwa gabungan beberapa *Decision Tree* dalam satu model meningkatkan stabilitas dan akurasi prediksi hingga 15% dibandingkan model tunggal. Dengan memahami cara kerja *Decision Tree*, peneliti dapat mengetahui bagaimana atribut-atribut psikologis berinteraksi membentuk pola prediktif terhadap kebutuhan konsultasi kesehatan mental.

Dalam konteks penelitian ini, penerapan *Decision Tree* membantu menjelaskan bagaimana variabel-variabel seperti tingkat gangguan kerja akibat kondisi mental, manfaat perusahaan terhadap kesehatan mental, dan dukungan sosial berpengaruh terhadap hasil klasifikasi. Pemahaman mendalam terhadap

struktur pohon keputusan ini menjadi dasar penting untuk mengembangkan model *Random Forest* yang lebih kompleks namun tetap interpretatif dalam menganalisis kondisi kesehatan mental pekerja *remote*.

2.4 *Random Forest*

Random Forest merupakan algoritma *ensemble learning* yang dikembangkan oleh Breiman (2001) dengan tujuan meningkatkan akurasi dan mengurangi *overfitting* yang sering terjadi pada *Decision Tree* tunggal. Algoritma ini bekerja dengan membangun sejumlah *Decision Tree* secara acak pada subset data dan fitur, kemudian menggabungkan hasil klasifikasi dari setiap pohon melalui proses *majority voting* untuk klasifikasi (Friedman et al., 2009). Pendekatan ini menjadikan *Random Forest* lebih stabil dan akurat dibandingkan model prediksi tunggal.

Setiap pohon dalam *Random Forest* dilatih menggunakan teknik *bootstrap sampling*, di mana sebagian data dipilih secara acak dengan pengembalian (*sampling with replacement*). Proses ini menciptakan keragaman antar pohon sehingga mengurangi risiko kesalahan akibat ketergantungan data tertentu. Selain itu, algoritma ini juga mengukur *feature importance* untuk mengidentifikasi variabel paling berpengaruh terhadap hasil klasifikasi (Marsuhandi et al., 2020). Dalam penelitian ini, fitur seperti *work_interfere*, *mental_health_consequence*, dan *benefits* menjadi faktor utama yang diuji dalam model.

Dalam Al-Qur'an, pohon (syajarah) digunakan dalam berbagai konteks simbolik, yaitu sebagai gambaran tentang kehidupan, keputusan, dan kebaikan atau keburukan. Misalnya dalam Surat Ibrahim ayat 24-25:

أَلَمْ تَرَ كَيْفَ ضَرَبَ اللَّهُ مَثَلًا كَلِمَةً طَيِّبَةً كَشَجَرَةٍ طَيِّبَةٍ أَصْلُهَا ثَابِتٌ وَفَرْعُهَا فِي السَّمَاءِ
تُؤْتِي أُكْلَهَا كُلَّ حِينٍ ۖ بِإِذْنِ رَبِّهَا وَيَضْرِبُ اللَّهُ الْأَمْثَالَ لِلنَّاسِ لَعَلَّهُمْ يَتَذَكَّرُونَ

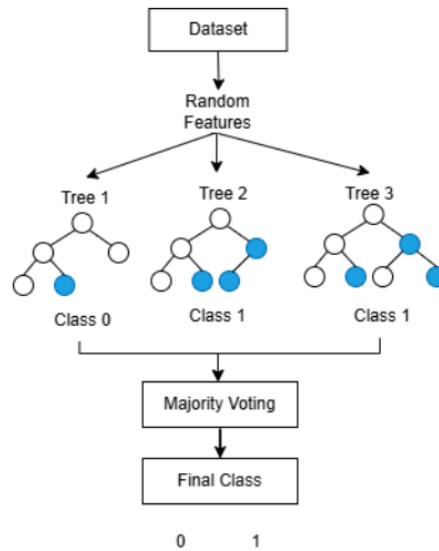
"Tidakkah kamu perhatikan bagaimana Allah telah membuat perumpamaan kalimat yang baik seperti pohon yang baik, akarnya teguh dan cabangnya (menjulang) ke langit, pohon itu memberikan buahnya pada setiap waktu dengan izin Tuhannya" (QS. Ibrahim: 24-25)

Ayat ini menggambarkan bahwa pohon yang baik (syajarah ṭayyibah) dengan akar kuat dan cabang yang menjulang melambangkan kebaikan, kemantapan, dan kebermanfaatan yang terus menerus (Sauda, 2020). Ini bisa diparalelkan dengan prinsip decision tree yang kuat, yang memiliki akar keputusan (fitur penting), cabang (split data), dan hasil (klasifikasi).

Keunggulan utama *Random Forest* terletak pada kemampuannya menangani dataset besar, kompleks, dan tidak seimbang (*imbalanced data*). Metode ini juga efektif ketika menghadapi data dengan *missing values* serta tidak memerlukan normalisasi. Menurut Sebayang et al. (2023), *Random Forest* menunjukkan performa yang unggul dalam memprediksi kondisi psikologis karena mampu mengakomodasi interaksi antarvariabel tanpa kehilangan akurasi. Oleh sebab itu, algoritma ini banyak digunakan dalam penelitian kesehatan mental berbasis data survei.

Dalam penelitian ini, *Random Forest* diterapkan pada dataset OSMI Tech in Survey 2016 untuk mengklasifikasi kebutuhan konsultasi kesehatan mental pekerja *remote*. Pendekatan tanpa teknik *balancing* dipilih agar hasil prediksi tetap mencerminkan kondisi nyata distribusi data, di mana proporsi individu yang membutuhkan konsultasi lebih sedikit dibandingkan yang tidak. Model dievaluasi

menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score* untuk menilai kinerja secara komprehensif.



Gambar 2.1 Arsitektur Random Forest

Pada Gambar 2.1 menunjukkan arsitektur dari algoritma *Random Forest* yang digunakan dalam penelitian ini. Dataset pekerja *remote* dibagi menjadi beberapa subset menggunakan teknik *bootstrap sampling*, di mana setiap subset digunakan untuk membangun satu *decision tree* (Tree 1, Tree 2, Tree 3). Setiap *tree* dilatih secara independen dengan subset fitur yang dipilih secara acak di setiap node, sehingga masing-masing *tree* dapat menghasilkan hasil prediksi yang berbeda.

Sebagai contoh, *Tree 1* mengklasifikasi bahwa seorang pekerja membutuhkan konsultasi kesehatan mental (*class 1*), sedangkan *Tree 2* dan *Tree 3* memprediksi bahwa pekerja tersebut tidak membutuhkannya (*class 0*). Hasil prediksi dari seluruh *tree* kemudian digabungkan menggunakan mekanisme *majority voting*, yaitu proses pengambilan keputusan berdasarkan suara terbanyak.

Pada contoh tersebut, karena dua dari tiga *tree* mengklasifikasi *class 0*, maka sistem menetapkan hasil akhir (*final class*) sebagai “tidak membutuhkan konsultasi”.

Algoritma *Random Forest* bekerja melalui beberapa tahapan berikut:

- a. Dataset berukuran n dibentuk menjadi beberapa subset menggunakan teknik *sampling with replacement*. Setiap subset memiliki ukuran yang sama dengan dataset asli, namun sebagian sampel dapat terambil lebih dari satu kali, sementara sebagian lainnya tidak terambil sama sekali. Berdasarkan teori *bootstrap* yang dijelaskan oleh Hartshorn (2023), sekitar 63,2% data digunakan sebagai *training data*, sedangkan 36,8% sisanya digunakan sebagai *Out-of-Bag (OOB)* data yang berfungsi sebagai *internal testing set* tanpa memerlukan dataset validasi tambahan. Pendekatan ini meningkatkan efisiensi proses pelatihan sekaligus mengurangi risiko *overfitting* dalam model *Random Forest*.
- b. Dalam proses pembentukan setiap *tree*, hanya sebagian fitur yang dipilih secara acak di tiap node. Strategi ini meningkatkan *diversity* antar *tree* dan menurunkan korelasi di antara mereka. Dari total 21 fitur prediktor pada dataset, hanya sekitar $\sqrt{21} \approx 4$ fitur acak yang digunakan di setiap node, misalnya *age*, *gender*, *remote_work*, dan *family_history*. Variasi ini membuat setiap *tree* memiliki karakteristik keputusan yang berbeda, yang pada akhirnya memperkuat kemampuan generalisasi model.
- c. Proses pemisahan data pada setiap node dilakukan menggunakan metrik Gini Impurity untuk mengukur tingkat ketidakmurnian data. Nilai Gini

menunjukkan seberapa baik suatu fitur dapat memisahkan kelas target.

Rumus *Gini Impurity* dapat dituliskan sebagai berikut:

$$Gini(s_i) = 1 - \sum_{j=1}^c p_{ij}^2 \quad 2.1$$

Keterangan:

s_i : node ke-i,

c : jumlah kelas target,

p_{ij} : proporsi kelas ke-j pada node ke-i.

Semakin kecil nilai Gini, semakin murni node tersebut, dan semakin baik fitur dalam membedakan kelas. Pada penelitian ini, kelas target adalah treatment dengan nilai 1 untuk pekerja yang membutuhkan konsultasi kesehatan mental, dan 0 untuk yang tidak membutuhkan konsultasi kesehatan mental.

- d. Setelah node terbaik ditentukan berdasarkan nilai Gini terkecil, proses pembentukan *tree* dilakukan secara rekursif hingga kondisi homogen (Gini = 0) atau kedalaman maksimum tercapai. Karena sebagian besar atribut pada dataset bersifat kategorikal seperti gender, benefits, dan work_interfere, maka pemisahan dilakukan dengan membandingkan kombinasi kategori antar atribut. Kombinasi dengan nilai *Gini Split* terkecil akan dipilih sebagai *split* terbaik.

Titik *split* terbaik ditentukan dengan rumus berikut:

$$Gini_{split} = \sum_{i=1}^k \frac{n_i}{n} \times Gini(s_i) \quad 2.2$$

Keterangan:

n_i : jumlah data subset ke-i,

n : total data dalam node,

$Gini(s_i)$: nilai Gini subset ke-i.

- e. Setelah semua *tree* terbentuk, hasil prediksi dari masing-masing *tree* digabungkan dengan teknik *majority voting*. Mekanisme ini menentukan hasil akhir berdasarkan kelas yang paling sering dipilih oleh seluruh *tree*. Sebagai contoh, jika tiga *tree* menghasilkan prediksi berbeda dan dua di antaranya menyatakan bahwa pekerja membutuhkan konsultasi kesehatan mental (*class I*), maka hasil akhir yang dipilih adalah “butuh konsultasi”.
- f. Kinerja model diukur menggunakan *Out-of-Bag Accuracy*, yaitu tingkat akurasi prediksi pada data yang tidak digunakan saat pelatihan. Rumus perhitungannya adalah:

$$OOB_{Accuracy} = \frac{1}{N} \sum_{i=1}^N I(\hat{y}_{OOB_i} = y_i) \quad 2.3$$

Keterangan:

N : jumlah total sampel.

i : Indeks untuk setiap sampel data (dari 1 sampai (N)).

$\hat{y}_{OOB_i}$: prediksi model untuk sampel ke- i .

y_i : label aktual dari sampel ke- i .

$I(.)$: fungsi indikator yang bernilai 1 jika prediksi benar dan 0 jika salah.

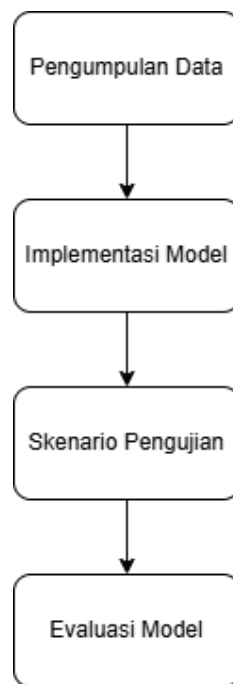
BAB III

DESAIN DAN IMPLEMENTASI

3.1 Desain Penelitian

Pada subbab ini menjelaskan tahapan yang dilaksanakan dalam proses penelitian. Desain penelitian disusun sebagai pedoman agar penelitian dapat berjalan secara sistematis, efisien, dan terarah sesuai dengan tujuan yang hendak dicapai. Penelitian ini menggunakan pendekatan kuantitatif dengan metode eksperimental yang bertujuan untuk menguji performa algoritma *machine learning* dalam mengklasifikasi kondisi kesehatan mental pekerja *remote*. Alur penelitian dirancang melalui beberapa tahapan utama, yaitu pengumpulan data, perancangan sistem, skenario pengujian, evaluasi, serta penarikan kesimpulan. Setiap tahapan memiliki fungsi yang saling berkaitan untuk menjamin keterpaduan proses penelitian. Dengan adanya desain ini, keseluruhan kegiatan penelitian diharapkan dapat tergambarkan secara jelas, terstruktur, dan sesuai dengan langkah-langkah ilmiah yang telah ditetapkan.

Tahapan penelitian yang disusun bertujuan untuk memastikan proses pengolahan data hingga evaluasi model dapat dilakukan secara sistematis dan konsisten. Desain penelitian ini juga memudahkan dalam melakukan analisis hasil pengujian serta penarikan kesimpulan berdasarkan data yang diperoleh. Selain itu, desain penelitian memberikan gambaran menyeluruh mengenai alur penelitian dalam membangun model klasifikasi kondisi kesehatan mental menggunakan algoritma *Random Forest*.



Gambar 3.1 Desain Penelitian Alur Proses

3.2 Pengumpulan Data

Data yang digunakan dalam penelitian ini bersifat sekunder, yang berarti peneliti tidak melakukan pengumpulan data secara langsung, melainkan menggunakan data yang telah tersedia secara publik. Dataset yang digunakan adalah *Open Sourcing Mental Illness (OSMI) Mental Health in Tech Survey 2016*, yaitu hasil survei terhadap para pekerja di industri teknologi yang memuat informasi mengenai kondisi kesehatan mental, lingkungan kerja, dan dukungan perusahaan terhadap kesehatan psikologis karyawan. Dataset ini dapat diakses secara terbuka melalui platform *Kaggle* pada tautan berikut: https://www.kaggle.com/datasets/osmi/mental-health-in-tech-2016?select=mental-health-in-tech-2016_20161114.csv. Dataset ini telah banyak digunakan sebagai acuan dalam penelitian-penelitian sejenis di bidang *mental health clasification*.

Dataset OSMI 2016 memiliki total 1.433 responden dengan 63 atribut. Untuk menyesuaikan dengan fokus penelitian, data disaring berdasarkan kolom “*Do you work remotely?*” dengan kategori “*Always*” dan “*Sometimes*”, yang masing-masing berjumlah 343 responden dan 757 responden. Kedua kategori tersebut mewakili pekerja *remote* penuh dan sebagian, sehingga total data yang digunakan dalam penelitian ini adalah 1.100 responden yang bekerja secara *remote*. Data tersebut dipilih karena relevan dengan konteks penelitian mengenai prediksi konsultasi kesehatan mental pada pekerja jarak jauh. Dari total 63 atribut, digunakan 21 atribut yang dianggap relevan dengan konteks penelitian, atribut-atribut tersebut ditampilkan pada Tabel 3.1 berikut.

Tabel 3.1 Detail Atribut Dataset

No	Atribut	Keterangan	Tipe Data	Jenis Data	Nilai / Rentang
1	<i>Age</i>	Usia responden dalam satuan tahun	Integer	Numerik	18 – 72 tahun
2	<i>Gender</i>	Jenis kelamin responden	String	Kategorikal	<i>Male / Female / Other</i>
3	<i>Country</i>	Negara tempat responden bekerja	String	Kategorikal	49 negara (US, UK, Canada, Indonesia, dll)
4	<i>State</i>	Negara bagian atau wilayah kerja (khusus Amerika)	String	Kategorikal	50 state (California, Texas, dll)
5	<i>Self_employed</i>	Status bekerja mandiri (freelance/self-employed)	String	Kategorikal	<i>Yes / No</i>
6	<i>Remote_work</i>	Frekuensi bekerja secara <i>remote</i>	String	Kategorikal	<i>Always / Sometimes / Never</i>
7	<i>Family_history</i>	Riwayat keluarga dengan gangguan kesehatan mental	String	Kategorikal	<i>Yes / No</i>
8	<i>Treatment</i>	Status pernah/tidaknya menjalani pengobatan kesehatan mental (label target)	Integer (0/1)	Kategorikal (Biner)	1 = Ya, 0 = Tidak
9	<i>Work_interfere</i>	Tingkat gangguan kesehatan mental	String	Kategorikal	<i>Never / Rarely /</i>

No	Atribut	Keterangan	Tipe Data	Jenis Data	Nilai / Rentang
		terhadap pekerjaan			<i>Sometimes / Often</i>
10	<i>No_employees</i>	Jumlah karyawan di perusahaan responden	String	Kategorikal	1–5 / 6–25 / 26–100 / 100–500 / 500+
11	<i>Tech_company</i>	Status bekerja di bidang teknologi	String	Kategorikal	<i>Yes / No</i>
12	<i>Benefits</i>	Fasilitas kesehatan mental di perusahaan	String	Kategorikal	<i>Yes / No / Don't know</i>
13	<i>Care_options</i>	Kemudahan akses ke layanan kesehatan mental	String	Kategorikal	<i>Easy / Difficult / Don't know</i>
14	<i>Wellness_program</i>	Adanya program kesejahteraan mental di perusahaan	String	Kategorikal	<i>Yes / No / Don't know</i>
15	<i>Anonymity</i>	Kebijakan anonimitas dalam mencari bantuan psikologis	String	Kategorikal	<i>Yes / No / Don't know</i>
16	<i>Leave</i>	Kemudahan mengambil cuti untuk alasan kesehatan mental	String	Kategorikal	<i>Very easy / Somewhat easy / Don't know / Very difficult</i>
17	<i>Mental_health_consequence</i>	Dampak terhadap karier jika membicarakan isu mental	String	Kategorikal	<i>Yes / No / Maybe</i>
18	<i>Coworkers</i>	Kenyamanan membicarakan isu mental dengan rekan kerja	String	Kategorikal	<i>Yes / No / Some of them</i>
19	<i>Supervisor</i>	Kenyamanan membicarakan isu mental dengan atasan	String	Kategorikal	<i>Yes / No / Some of them</i>
20	<i>Mental_health_interview</i>	Pengalaman membicarakan isu mental saat wawancara kerja	String	Kategorikal	<i>Yes / No / Maybe</i>
21	<i>Seek_help</i>	Kebijakan perusahaan dalam memfasilitasi karyawan mencari bantuan profesional	String	Kategorikal	<i>Yes / No / Don't know</i>

Pada Tabel 3.1 terdapat 21 atribut yang digunakan dalam penelitian ini, terdiri atas 20 atribut prediktor dan 1 atribut target, yaitu *treatment*. Atribut-atribut tersebut diperoleh dari dataset kesehatan mental pekerja teknologi dan digunakan sebagai variabel awal dalam proses pemodelan. Selanjutnya, dilakukan proses

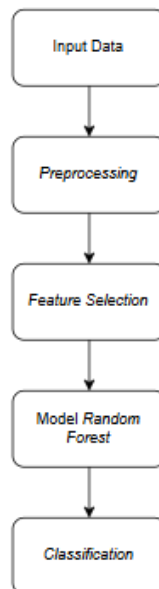
feature selection menggunakan metode *Feature Importance* pada algoritma *Random Forest* untuk menentukan atribut yang memiliki kontribusi terbesar terhadap klasifikasi variabel target treatment. Dengan demikian, atribut yang digunakan dalam penelitian ini merupakan hasil seleksi berdasarkan tingkat kepentingan fitur yang diperoleh dari model, sehingga fitur yang digunakan dinilai paling relevan dan representatif dalam proses prediksi kondisi kesehatan mental pekerja remote.

Tabel 3.2 Sample Dataset

<i>No</i>	<i>Age</i>	<i>Gender</i>	<i>Country</i>	<i>Self_employed</i>	<i>Remote_work</i>	<i>Family_history</i>	<i>Treatment</i>	<i>Work_interruption</i>	<i>No_employment</i>	<i>Tech_company</i>	<i>Benefits</i>	<i>Care_options</i>	<i>Wellness_program</i>	<i>Anonymity</i>	<i>Leave</i>	<i>Mental_health_consequence</i>	<i>Coworkers</i>	<i>Supervisor</i>	<i>Mental_health_interview</i>
1	28	Male	United States	No	Sometimes	Yes	1	Often	26–100	Yes	Yes	Don't know	No	Yes	Somewhat easy	No	Yes	No	Yes
2	35	Female	Canada	No	Always	No	0	Rarely	100–500	Yes	Yes	Easy	Yes	Yes	Very easy	Maybe	Yes	No	Yes
3	42	Male	United Kingdom	Yes	Never	Yes	1	Sometimes	6–25	Yes	No	Difficult	No	No	Very difficult	Yes	No	Maybe	Don't know
4	30	Female	Germany	No	Sometimes	No	0	Never	26–100	Yes	Don't know	Don't know	Yes	Yes	Somewhat easy	No	Yes	No	Yes
5	24	Male	India	No	Always	Yes	1	Often	100–500	Yes	Yes	Easy	No	Yes	Very easy	No	Some of them	No	Yes
6	29	Female	Australia	No	Sometimes	No	0	Rarely	6–25	Yes	Yes	Don't know	Yes	Yes	Somewhat easy	Maybe	Yes	Yes	Yes

3.3 Implementasi Model

Sistem klasifikasi kesehatan mental pada pekerja *remote* dirancang dengan arsitektur vertikal, yang menggambarkan alur proses pengolahan data secara berurutan mulai dari input data mentah hingga menghasilkan output klasifikasi akhir. Arsitektur ini dipilih karena sesuai dengan karakteristik alur kerja *machine learning* yang membutuhkan tahapan sistematis agar hasil prediksi yang diperoleh valid dan reliabel. Proses ini mencakup tahapan utama yaitu input data, *preprocessing*, *feature selection*, *training model Random Forest*, dan *clasification*.



Gambar 3.2 Implementasi Model

3.4 Input Data

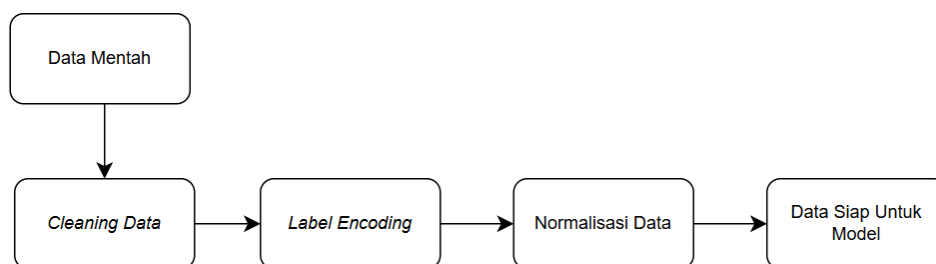
Tahap input data merupakan langkah awal dalam sistem prediksi yang bertujuan untuk memasukkan dataset mentah ke dalam lingkungan pemrosesan. Data yang digunakan pada penelitian ini berasal dari dataset publik *Mental Health*

in Tech Survey 2016 yang dirilis oleh *Open Sourcing Mental Illness (OSMI)*. Dataset ini berisi hasil survei terhadap pekerja di bidang teknologi mengenai kondisi kesehatan mental, dukungan perusahaan, serta lingkungan kerja yang mereka alami. Dari total 1.433 data mentah yang tersedia, dilakukan pemilahan terhadap responden yang bekerja secara *remote* dengan kategori “*Always*” dan “*Sometimes*”, sehingga diperoleh 1.100 data responden yang digunakan sebagai dasar penelitian. Proses ini memastikan bahwa data yang digunakan relevan dengan tujuan penelitian, yaitu untuk mengklasifikasi kebutuhan konsultasi kesehatan mental pada pekerja *remote*.

3.5 Preprocessing

Tahap *preprocessing* merupakan langkah penting dalam penelitian berbasis *machine learning* untuk menyiapkan data agar siap digunakan pada tahap pemodelan. Tujuannya adalah memastikan data berada dalam kondisi bersih, konsisten, dan dalam format yang dapat diproses oleh algoritma *Random Forest*. Dataset yang digunakan dalam penelitian ini berasal dari *OSMI Mental Health in Tech 2016*, yang terdiri dari 1.433 responden dan 63 atribut. Setelah dilakukan penyaringan terhadap pekerja *remote* dengan kategori *Always* dan *Sometimes*, diperoleh 1.100 responden dan 21 atribut yang digunakan sebagai dasar pemodelan. Dari jumlah tersebut, 557 responden menjawab “Yes” (pernah menjalani konsultasi kesehatan mental) dan 543 menjawab “No” (belum pernah), menunjukkan bahwa distribusi data tergolong seimbang dan tidak memerlukan teknik penyeimbangan seperti *SMOTE*. Karena dataset bersifat tabular dan berisi atribut numerik serta kategorikal, maka tahapan *preprocessing* difokuskan pada

pembersihan data, transformasi nilai kategorikal menjadi numerik, serta penyeragaman skala data agar dapat diproses oleh algoritma *Random Forest* secara optimal. Berikut proses yang terjadi dalam tahapan *preprocessing*.



Gambar 3.3 Preprocessing

3.5.1 Data Mentah

Tahap data mentah merupakan langkah awal dalam pengolahan dataset hasil input data yang telah difilter sesuai kriteria penelitian. Dataset berisi atribut seperti *Age*, *Gender*, *Family_history*, *Remote_work*, *Benefits*, *Work_interfere*, dan *Treatment*, yang mencerminkan kondisi kesehatan mental dan lingkungan kerja responden. Data ini masih mengandung nilai kosong dan format kategorikal yang belum siap diolah oleh algoritma *machine learning*. Oleh karena itu, tahap ini berfungsi sebagai dasar sebelum dilakukan proses pembersihan, konversi, dan normalisasi agar data siap digunakan dalam pemodelan.

3.5.2 Cleaning Data

Tahap *cleaning data* dilakukan untuk meningkatkan kualitas dataset dengan cara menghapus data duplikat, memperbaiki ketidakkonsistenan, serta menangani nilai kosong (*missing values*). Nilai kosong pada atribut kategorikal seperti *self_employed* dan *work_interfere* diisi menggunakan modus (nilai yang

paling sering muncul), sedangkan nilai kosong pada atribut numerik seperti *Age* diisi dengan median (nilai tengah) untuk menghindari pengaruh data ekstrem. Selain itu, atribut yang tidak relevan terhadap proses klasifikasi, seperti *Timestamp* dan *State*, dihapus karena tidak memberikan kontribusi terhadap variabel target. Setelah tahap ini, data menjadi lebih bersih dan siap untuk proses transformasi berikutnya.

3.5.3 *Label Encoding*

Tahap *label encoding* dilakukan untuk mengubah nilai kategorikal dalam dataset menjadi bentuk numerik agar dapat diolah oleh algoritma *machine learning*. Sebagian besar atribut pada dataset *OSMI 2016*, seperti *gender*, *remote_work*, *family_history*, *work_interfere*, dan *benefits*, memiliki nilai berbentuk teks yang tidak dapat diproses langsung oleh algoritma seperti *Random Forest*. Oleh karena itu, diperlukan proses konversi kategori menjadi nilai numerik tanpa mengubah makna atau bobot informasinya. Proses *encoding* ini menghasilkan data yang lebih terstruktur dan siap digunakan dalam tahap pemodelan. Tabel berikut menunjukkan contoh data sebelum dilakukan proses *label encoding*, di mana setiap atribut masih dalam bentuk teks atau kategori.

Tabel 3.3 Data Sebelum *Encoding*

Gender	Remote Work	Family History	Work Interfere	Benefits	Treatment
Male	Always	Yes	Sometimes	Yes	1
Female	Sometimes	No	Rarely	No	0
Male	Always	Yes	Often	Yes	1
Other	Sometimes	Yes	Never	Don't know	0
Male	Always	No	Sometimes	Yes	1

Tabel 3.4 Data Setelah *Encoding*

Gender	Remote_Work	Family_History	Work_Interfere	Benefits	Treatment
0	1	1	2	2	1
1	0	0	1	0	0
0	1	1	3	2	1
2	0	1	0	1	0
0	1	0	2	2	1

Keterangan kode:

Gender: Male = 0, Female = 1, Other = 2

Remote_Work: Sometimes = 0, Always = 1

Family_History: No = 0, Yes = 1

Work_Interfere: Never = 0, Rarely = 1, Sometimes = 2, Often = 3

Benefits: No = 0, Don't Know = 1, Yes = 2

Treatment: No = 0, Yes = 1

Perbandingan antara Tabel 3.4 dan Tabel 3.5 menunjukkan hasil transformasi data kategorikal menjadi data numerik. Misalnya, pada atribut *Remote_Work*, nilai “*Always*” dikonversi menjadi 1 dan “*Sometimes*” menjadi 0, sedangkan pada atribut *Work_Interfere*, semakin tinggi angka menunjukkan tingkat gangguan yang semakin besar terhadap pekerjaan. Proses ini dilakukan menggunakan metode *Label Encoding* karena lebih efisien untuk data kategorikal yang bersifat ordinal dan non-ordinal. Dengan demikian, hasil transformasi ini memastikan seluruh data berada dalam format numerik yang dapat diolah dengan baik oleh algoritma *Random Forest* pada tahap pemodelan selanjutnya.

3.5.4 Normalisasi Data

Tahap normalisasi data dilakukan untuk menyeragamkan skala nilai pada atribut numerik agar tidak menimbulkan bias selama proses pelatihan model. Pada penelitian ini, atribut *age* memiliki rentang nilai yang cukup besar dibandingkan atribut lain yang sudah dikonversi ke bentuk kategorikal. Perbedaan skala ini dapat menyebabkan algoritma *machine learning* memberikan bobot lebih besar

terhadap variabel dengan nilai yang lebih tinggi. Oleh karena itu, dilakukan normalisasi menggunakan metode *Min-Max Normalization*, yang mengubah rentang nilai data ke interval $[0,1]$.

Rumus normalisasi yang digunakan adalah sebagai berikut:

$$X' = \frac{X - X_{min}}{X_{max} - X_{min}} \quad 3.1$$

dengan keterangan:

X' : nilai hasil normalisasi,

X : nilai asli atribut,

X_{min} : nilai minimum atribut,

X_{max} : nilai maksimum atribut.

Sebagai contoh, jika nilai minimum usia responden adalah 18 tahun dan maksimum 65 tahun, maka untuk responden dengan usia 35 tahun, hasil normalisasi dihitung sebagai berikut:

$$X'_{age} = \frac{35 - 18}{65 - 18} = 0.36$$

Artinya, nilai usia 35 tahun dikonversi menjadi 0.36 setelah proses normalisasi. Nilai ini menunjukkan posisi relatif dari usia responden terhadap rentang keseluruhan data tanpa mengubah proporsinya. Proses *Min-Max Normalization* menghasilkan skala nilai antara 0 hingga 1 untuk seluruh atribut numerik. Dengan demikian, tidak ada atribut yang mendominasi atribut lain selama pelatihan model *Random Forest*. Tahapan ini memastikan distribusi data yang lebih seimbang dan meningkatkan stabilitas model dalam menghasilkan klasifikasi terhadap variabel *treatment* pada pekerja *remote*.

3.6 Feature Selection

Tahap *feature selection* bertujuan untuk memilih fitur yang paling relevan terhadap variabel target treatment, yaitu status apakah seorang pekerja pernah menjalani konsultasi kesehatan mental atau tidak. Proses ini dilakukan untuk meningkatkan efisiensi model dan mengurangi risiko *overfitting* akibat adanya fitur yang tidak relevan atau memiliki kontribusi rendah terhadap hasil prediksi. Metode *feature selection* pada penelitian ini menggunakan pendekatan *Feature Importance* yang diperoleh dari model *Random Forest*. Nilai *feature importance* dihitung berdasarkan kontribusi masing-masing fitur dalam menurunkan nilai ketidakmurnian (*impurity decrease*) di setiap node pada seluruh pohon keputusan (*decision tree*). Ukuran ketidakmurnian yang digunakan adalah *Gini Impurity*, dengan rumus:

$$Gini(D) = 1 - \sum_{i=1}^m p_i^2 \quad 3.2$$

keterangan:

$Gini(D)$: nilai ketidakmurnian pada dataset D,

p_i : proporsi kelas ke-i dalam data D,

m : jumlah kelas target (*Yes* dan *No*).

Fitur dengan nilai *Gini Impurity* yang tinggi berarti lebih berpengaruh dalam menentukan keputusan pada model. Berdasarkan hasil perhitungan *feature importance*, atribut yang memiliki pengaruh terbesar terhadap hasil prediksi adalah *Age*, *Gender*, *Remote_Work*, *Family_History*, *Work_Interfere*, dan *Benefits*. Keenam atribut ini digunakan dalam proses pemodelan untuk memastikan model bekerja secara optimal (Breiman, 2001). Hasil seleksi fitur yang digunakan dalam penelitian ini ditunjukkan pada Tabel 3.5 berikut.

Tabel 3.5 Hasil Akhir Seleksi Fitur yang Digunakan dalam Pemodelan

No.	Atribut	Jenis Data	Keterangan
1	<i>age</i>	Numerik	Usia responden dalam tahun
2	<i>gender</i>	Kategorikal	Jenis kelamin responden
3	<i>country</i>	Kategorikal	Negara tempat bekerja responden
4	<i>state</i>	Kategorikal	Wilayah kerja atau negara bagian (khusus responden Amerika Serikat)
5	<i>self_employed</i>	Kategorikal	Status bekerja secara mandiri
6	<i>remote_work</i>	Kategorikal	Frekuensi bekerja jarak jauh
7	<i>family_history</i>	Kategorikal	Riwayat keluarga dengan gangguan kesehatan mental
8	<i>work_interfere</i>	Kategorikal	Pengaruh gangguan kesehatan mental terhadap pekerjaan
9	<i>no_employees</i>	Kategorikal	Jumlah karyawan di perusahaan tempat responden bekerja
10	<i>tech_company</i>	Kategorikal	Status bekerja di sektor teknologi
11	<i>benefits</i>	Kategorikal	Fasilitas kesehatan mental di tempat kerja
12	<i>care_options</i>	Kategorikal	Akses layanan kesehatan mental di perusahaan
13	<i>wellness_program</i>	Kategorikal	Adanya program kesejahteraan mental di tempat kerja
14	<i>anonymity</i>	Kategorikal	Kebijakan anonimitas dalam mencari bantuan psikologis
15	<i>leave</i>	Kategorikal	Kemudahan mengambil cuti untuk alasan kesehatan mental
16	<i>mental_health_consequence</i>	Kategorikal	Dampak terhadap karier jika membicarakan isu kesehatan mental
17	<i>coworkers</i>	Kategorikal	Kenyamanan membicarakan kesehatan mental dengan rekan kerja
18	<i>supervisor</i>	Kategorikal	Kenyamanan membicarakan kesehatan mental dengan atasan
19	<i>mental_health_interview</i>	Kategorikal	Pengalaman membahas isu kesehatan mental saat wawancara kerja
20	<i>seek_help</i>	Kategorikal	Kebijakan perusahaan dalam memfasilitasi bantuan profesional
21	<i>treatment</i>	Biner (0/1)	Variabel target: 1 = membutuhkan konsultasi, 0 = tidak membutuhkan

Setelah tahap *feature selection*, fitur-fitur terpilih selanjutnya digunakan dalam proses pelatihan model *Random Forest Classifier*. Tahapan ini dilakukan untuk memastikan bahwa atribut yang digunakan benar-benar relevan dan memiliki kontribusi signifikan terhadap proses prediksi kebutuhan konsultasi kesehatan mental pada pekerja *remote*. Dengan demikian, model yang dibangun

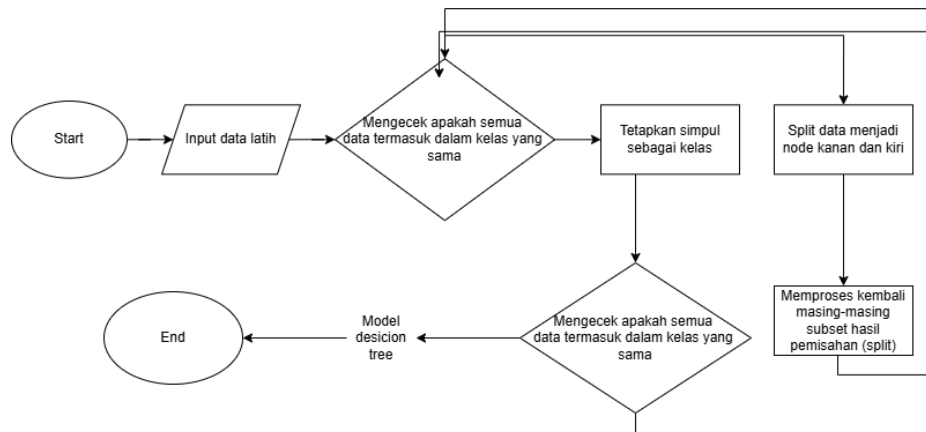
diharapkan mampu menghasilkan kinerja yang efisien, stabil, dan memiliki tingkat akurasi yang tinggi.

3.7 Model Random Forest

Random Forest merupakan algoritma pembelajaran *ensemble* yang diperkenalkan oleh Breiman (2001) sebagai pengembangan dari metode *Bootstrap Aggregating (Bagging)* dan *Decision Tree*. Algoritma ini bekerja dengan menggabungkan hasil dari sejumlah pohon keputusan (*decision trees*) yang dilatih menggunakan subset data dan subset fitur yang berbeda. Pendekatan ini menghasilkan model yang lebih stabil, akurat, dan tahan terhadap *overfitting* (Ibrahim, 2023). Kinerja *Random Forest* sangat dipengaruhi oleh pengaturan *hyperparameter* yang berfungsi mengontrol kompleksitas, variasi, dan stabilitas model. Tabel berikut menjelaskan beberapa *hyperparameter* penting yang digunakan dalam penelitian ini.

Tabel 3.6 Tuning Hyperparameter

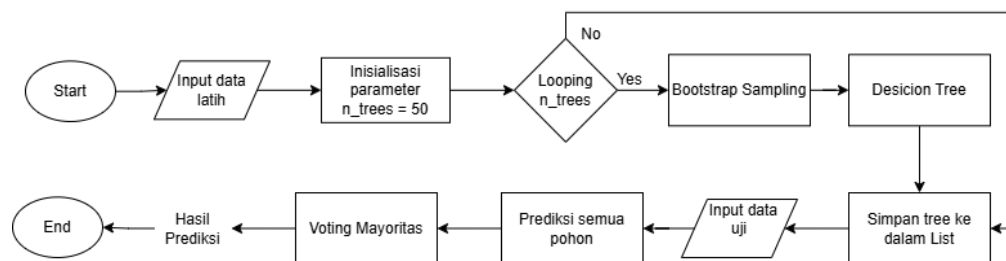
Parameter	Nilai	Deskripsi
<i>n_estimators</i>	500	Jumlah pohon keputusan (<i>trees</i>) yang dibangun dalam satu <i>forest</i> . Semakin banyak jumlah pohon, semakin stabil hasil prediksi, namun waktu komputasi meningkat.
<i>max_features</i>	4	Jumlah fitur acak yang digunakan pada setiap pemisahan node. Nilai diperoleh dari hasil perhitungan $\sqrt{21} \approx 4,58$ dan dibulatkan ke bawah menjadi 4 fitur acak per node.
<i>max_depth</i>	None (otomatis)	Kedalaman maksimum pohon keputusan. Jika None, maka pohon akan tumbuh hingga seluruh data terklasifikasi sempurna.
<i>min_samples_split</i>	2	Jumlah minimum sampel yang diperlukan untuk membagi sebuah node menjadi dua cabang.
<i>min_samples_leaf</i>	1	Jumlah minimum sampel yang harus terdapat pada setiap <i>leaf node</i> .
<i>criterion</i>	“gini”	Kriteria pemisahan node berdasarkan pengurangan nilai <i>Gini Impurity</i> .



Gambar 3.4 Flowchart Decision Tree

Pada Gambar 3.4 menggambarkan alur proses pembentukan model decision tree yang digunakan dalam penelitian ini. Proses dimulai dengan melakukan input data latih yang telah melalui tahap praproses sebelumnya. Selanjutnya, sistem akan memeriksa apakah seluruh data yang terdapat dalam simpul (node) tersebut termasuk dalam kelas yang sama. Jika semua data berada dalam satu kelas yang sama, maka simpul tersebut akan ditetapkan sebagai simpul daun (leaf node), dan proses pada cabang tersebut dihentikan. Namun, apabila data terdiri dari lebih dari satu kelas, maka dilakukan proses pemisahan (split) terhadap data berdasarkan atribut terbaik. Proses ini menghasilkan dua subset data, yaitu subset untuk node kiri dan node kanan. Masing-masing subset hasil pemisahan kemudian diproses kembali dengan langkah yang sama secara rekursif. Proses ini terus berulang hingga setiap simpul hanya mengandung data dari satu kelas atau hingga kondisi terminasi tertentu terpenuhi. Hasil akhir dari proses ini adalah model pohon keputusan yang dapat digunakan untuk melakukan prediksi terhadap data baru. Flowchart tersebut menggambarkan secara umum prinsip

kerja algoritma decision tree yang bersifat rekursif dan berbasis pada pemilihan atribut terbaik untuk membagi data agar mencapai prediksi yang optimal.



Gambar 3.5 Flowchart Random Forest

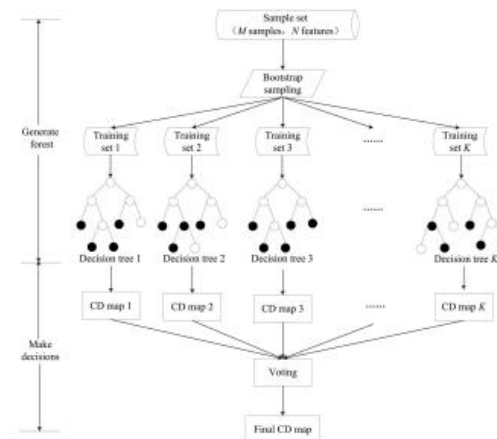
Gambar 3.5 menjelaskan alur kerja algoritma Random Forest yang digunakan dalam penelitian ini untuk memprediksi kebutuhan konsultasi kesehatan mental pekerja remote. Proses dimulai dengan memasukkan data latih sebanyak ± 1.100 responden dari dataset OSMI Mental Health in Tech 2016 yang telah melalui tahap preprocessing berupa cleaning data, label encoding, dan normalisasi. Setelah itu dilakukan inisialisasi parameter utama, yaitu jumlah pohon keputusan ($n_estimators$) sebanyak 500, kedalaman maksimum pohon (max_depth) yang tidak dibatasi, jumlah minimum sampel untuk pemisahan ($min_samples_split$) sebanyak 2, serta jumlah fitur acak per node ($max_features$) sebanyak 4 dengan kriteria Gini sebagai ukuran impuritas.

Tahap berikutnya adalah proses looping untuk membentuk pohon keputusan sesuai jumlah yang telah ditentukan. Pada setiap iterasi dilakukan pemeriksaan apakah jumlah pohon yang terbentuk masih kurang dari nilai $n_estimators$. Jika kondisi “Yes” terpenuhi, maka proses dilanjutkan ke tahap bootstrap sampling, yaitu pengambilan data latih secara acak dengan pengembalian sebanyak ± 695 data untuk melatih satu pohon keputusan. Pohon

yang terbentuk kemudian disimpan ke dalam list pohon, dan sistem kembali ke tahap pengecekan jumlah pohon. Proses ini akan terus berulang hingga jumlah pohon mencapai 500. Sebaliknya, jika kondisi “No” terpenuhi, artinya jumlah pohon telah mencapai batas maksimal, maka proses looping dihentikan.

Setelah seluruh pohon terbentuk, proses dilanjutkan ke tahap klasifikasi. Data uji sebanyak ± 405 data *Out-of-Bag* (OOB) atau melalui skenario validasi *K-Fold Cross Validation* (5-Fold: 880 data latih dan 220 data uji; 10-Fold: 990 data latih dan 110 data uji) dimasukkan ke dalam model. Masing-masing pohon memberikan prediksi terhadap data uji, kemudian hasil prediksi dari seluruh pohon dikumpulkan dan dilakukan proses voting mayoritas. Label kelas yang paling banyak dipilih oleh seluruh pohon akan menjadi hasil klasifikasi akhir, yaitu 0 = membutuhkan konsultasi dan 1 = tidak membutuhkan konsultasi.

Tahap terakhir adalah evaluasi model menggunakan metrik seperti *accuracy*, *precision*, *recall*, *F1-score*, dan *ROC-AUC*, serta validasi dengan OOB score dan *K-Fold Cross Validation* untuk memastikan performa model stabil dan representatif terhadap distribusi data nyata. Dengan demikian, flowchart ini menggambarkan secara sistematis bagaimana algoritma Random Forest bekerja dalam penelitian untuk mendeteksi kebutuhan konsultasi kesehatan mental pekerja remote.



Gambar 3.6 Flowchart Random Forest Classifier

Flowchart 3.6 *Random Forest Classifier* menggambarkan tahapan proses klasifikasi yang dilakukan oleh algoritma Random Forest, mulai dari data masukan hingga menghasilkan prediksi akhir. Proses diawali dengan sample set, yaitu kumpulan data yang terdiri atas sejumlah sampel (M samples) dan sejumlah atribut (N features) yang digunakan sebagai data pelatihan. Selanjutnya, data tersebut diproses menggunakan teknik bootstrap sampling untuk membentuk beberapa data latih (*training set*) yang berbeda melalui proses pengambilan sampel secara acak dengan pengembalian (*sampling with replacement*). Tujuan dari proses ini adalah menghasilkan variasi data latih sehingga setiap pohon keputusan memiliki karakteristik yang berbeda.

Setelah beberapa training set terbentuk, masing-masing data latih digunakan untuk membangun satu decision tree secara independen. Pada setiap node dalam pohon keputusan, algoritma Random Forest memilih sebagian atribut secara acak (*random feature selection*) untuk menentukan pemisahan data yang paling optimal. Proses tersebut dilakukan secara berulang hingga terbentuk

sejumlah decision tree, yaitu Decision Tree 1 sampai Decision Tree K, yang masing-masing menghasilkan prediksi secara terpisah. Hasil prediksi dari setiap pohon kemudian direpresentasikan sebagai CD Map (*Classification Decision Map*) yang menunjukkan keluaran klasifikasi dari masing-masing model.

Tahap terakhir adalah menggabungkan seluruh hasil prediksi yang diperoleh dari setiap decision tree melalui proses voting. Pada proses ini, kelas yang memperoleh jumlah suara terbanyak (majority voting) akan dipilih sebagai hasil klasifikasi akhir. Hasil voting tersebut menghasilkan Final CD Map, yaitu keputusan akhir yang merupakan gabungan dari seluruh prediksi setiap pohon keputusan sehingga menghasilkan model yang lebih stabil dan akurat dibandingkan penggunaan satu pohon keputusan saja. Dengan demikian, algoritma *Random Forest Classifier* mampu meningkatkan kemampuan klasifikasi serta mengurangi risiko *overfitting* karena keputusan akhir didasarkan pada hasil kolektif dari banyak decision tree.

3.7.1 Perhitungan Manual

Proses komputasi klasifikasi kebutuhan konsultasi kesehatan mental menggunakan algoritma Random Forest ditahap-awali oleh proses transformasi data (*data preprocessing*). Data mentah berupa teks jawaban kuesioner dikonversi menjadi representasi numerik menggunakan teknik *label encoding* agar dapat diolah secara matematis oleh algoritma. Pada variabel target yaitu *treatment*, jawaban kuesioner berupa "Yes" dikodekan ke dalam nilai biner 1 (yang merepresentasikan bahwa pekerja membutuhkan atau pernah menjalani perawatan kesehatan mental), sedangkan jawaban "No" dikodekan ke dalam nilai biner 0.

Fitur kontinu seperti *Age* (umur) diproses menggunakan normalisasi skala untuk menjaga keseimbangan bobot data. Setelah seluruh dataset siap, model memisahkan data menjadi data latih (*training data*) dan data uji (*testing data*) berdasarkan skenario pengujian.

Pada Skenario 1 dengan rasio pembagian data 60:40, diperoleh total data uji sebanyak 438 data yang akan dievaluasi oleh model. Langkah berikutnya adalah pembentukan sejumlah pohon keputusan (*Decision Trees*) secara simultan melalui metode *Bootstrap Aggregating* (Bagging). Algoritma mengambil sampel acak dari data latih dengan pengembalian (*with replacement*) untuk membangun setiap pohon secara independen. Di setiap percabangan (*node*) pada pohon keputusan yang terbentuk, dilakukan penentuan pembelahan cabang (*splitting node*) berbasis nilai tingkat impuritas data. Parameter fitur acak yang dievaluasi pada tiap pembelahan diatur sebesar $\text{max_features} = 4$. Kriteria yang digunakan untuk mengukur tingkat impuritas data (D) tersebut adalah kriteria *Gini Impurity*, yang dihitung menggunakan rumus matematis sebagai berikut:

$$Gini(s_i) = 1 - \sum_{j=1}^c p_{ij}^2 \quad 3.3$$

Di mana p_i merupakan nilai proporsi atau probabilitas kemunculan kelas s_i pada *node* yang bersangkutan. Sebagai ilustrasi perhitungan, apabila pada suatu *node* acak terdapat total 100 data pekerja yang terdiri atas 60 pekerja Kelas 1 (pernah *treatment*) dan 40 pekerja Kelas 0 (tidak pernah *treatment*), maka proporsi masing-masing kelas adalah $p_1 = 0,6$ dan $p_0 = 0,4$. Nilai impuritas awal pada *node* tersebut dihitung melalui persamaan $Gini = 1 - [(0,6)^2 + (0,4)^2]$

$= 1 - [0,36 + 0,16] = 0,48$. Algoritma akan terus mencari pembelahan fitur yang memberikan penurunan nilai impuritas (*Gini Gain*) terbesar secara rekursif hingga seluruh struktur pohon selesai dibangun. Proses ini diulang secara masif hingga mencapai parameter $n_estimators = 500$, menghasilkan total 500 pohon keputusan yang siap melakukan prediksi.

Tahap akhir dari algoritma ini adalah penentuan keputusan final melalui mekanisme *Majority Voting* (Voting Mayoritas) terhadap data uji. Ketika 438 data uji dialirkan ke dalam model, ke-500 pohon keputusan yang telah dilatih akan memberikan prediksi kelas secara independen untuk setiap baris data. Setiap prediksi dari satu pohon bernilai sebagai satu hak suara (*vote*). Hasil akumulasi voting terbanyak dari seluruh pohon diputuskan sebagai hasil klasifikasi final untuk data pekerja tersebut. Berdasarkan hasil pengujian riil pada Skenario 1, performa akumulatif dari mekanisme voting ini menghasilkan representasi matriks konfusi (*confusion matrix*) dengan detail: *True Positive* (TP) sebesar 231 data, *True Negative* (TN) sebesar 141 data, *False Positive* (FP) sebesar 38 data, and *False Negative* (FN) sebesar 28 data.

Untuk mengukur performa akhir dari hasil klasifikasi tersebut, dilakukan perhitungan manual metrik evaluasi menggunakan rumusan standar *machine learning*. Metrik pertama adalah akurasi yang mengukur persentase ketepatan prediksi model, dihitung melalui persamaan berikut:

$$Akurasi = \frac{TP + TN}{TP + TN + FP + FN} = \frac{231 + 141}{231 + 141 + 38 + 28}$$

Metrik kedua adalah presisi yang mengukur rasio ketepatan prediksi positif model, dihitung menggunakan rumus:

$$Presisi = \frac{TP}{TP + FP} = \frac{231}{231 + 38} = \frac{231}{269} = 0,8587 \text{ (85,87\%)}$$

Metrik ketiga adalah *recall* yang mengukur kemampuan model dalam menemukan kembali informasi kelas positif, dihitung dengan rumus:

$$Recall = \frac{TP}{TP + FN} = \frac{231}{231 + 28} = \frac{231}{259} = 0,8919 \text{ (89,19\%)}$$

Metrik terakhir adalah *F1-Score* yang mengukur nilai keseimbangan rata-rata harmonis antara presisi dan *recall*, dihitung melalui persamaan matematis sebagai berikut:

$$F1 - Score = 2 \times \frac{Presisi \times Recall}{Presisi + Recall} = 2 \times \frac{0,8587 \times 0,8919}{0,8587 + 0,8919} = 0,8750$$

Melalui visualisasi perhitungan manual di atas, dapat dibuktikan secara matematis bahwa aliran pemrosesan data pada model Random Forest menghasilkan keputusan klasifikasi yang valid dengan tingkat akurasi akhir mencapai 84,93% dan nilai *F1-score* sebesar 0,8750 pada kondisi pembagian data latih dan data uji sebesar 60:40.

3.8 Classification

Output sistem berupa klasifikasi biner, yaitu:

- a. 0 → Pekerja membutuhkan konsultasi kesehatan mental.
- b. 1 → Pekerja tidak membutuhkan konsultasi kesehatan mental.

Hasil klasifikasi yang diperoleh kemudian digunakan sebagai dasar untuk evaluasi performa model. Evaluasi dilakukan dengan metode k-Fold Cross Validation menggunakan dua variasi nilai k ($k = 5$ dan $k = 10$) untuk membandingkan tingkat keakuratan dan stabilitas model.

3.9 Skenario Pengujian

Tahap skenario pengujian dilakukan untuk menguji hasil prediksi model *Random Forest* dalam menentukan kebutuhan konsultasi kesehatan mental pada pekerja *remote*. Pengujian ini bertujuan untuk menilai seberapa baik model mampu memprediksi data pekerja ke dalam dua kategori, yaitu *membutuhkan konsultasi* (1) dan *tidak membutuhkan konsultasi* (0). Selain menguji akurasi hasil prediksi tersebut, penelitian ini juga melakukan perbandingan performa model berdasarkan variasi nilai k pada metode *k-Fold Cross Validation*, yaitu $k = 5$ dan $k = 10$, sehingga dapat diketahui konfigurasi validasi yang paling optimal terhadap dataset penelitian.

Metode *k-Fold Cross Validation* merupakan teknik validasi model yang membagi dataset menjadi k bagian (*fold*) dengan ukuran yang sama. Pada setiap iterasi, satu bagian digunakan sebagai data uji (*testing set*), sementara $(k-1)$ bagian lainnya digunakan sebagai data latih (*training set*). Proses ini diulang sebanyak k kali hingga setiap bagian pernah menjadi data uji satu kali. Nilai rata-rata hasil pengujian kemudian digunakan sebagai ukuran performa model secara keseluruhan.

Rumus perhitungan akurasi rata-rata dituliskan sebagai berikut:

$$Accuracy_{avg} = \frac{\sum_{i=1}^k Accuracy_i}{k} \quad 3.4$$

keterangan:

$Accuracy_{avg}$: nilai rata-rata akurasi model dari seluruh iterasi validasi.

$Accuracy_i$: nilai akurasi pada iterasi ke- i .

k : jumlah *fold* yang digunakan dalam proses validasi (*k-Fold Cross Validation*).

Berdasarkan hasil penelitian yang menggunakan dataset *National Health and Nutrition Examination Survey* (NHANES, 2024) serta *EEG Signal Dataset* oleh *EEG Research Group* (2025), validasi *k*-Fold 5 terbukti menghasilkan tingkat akurasi yang stabil di kisaran 88–89%. Sementara itu, penelitian oleh Rahman et al. (2023) menunjukkan bahwa penggunaan validasi *k*-Fold 10 pada algoritma *Random Forest* mampu mencapai akurasi yang lebih tinggi, yaitu sebesar 90,1%. Oleh karena itu, dalam penelitian ini digunakan kedua variasi nilai *k* tersebut, yaitu $k = 5$ dan $k = 10$, untuk membandingkan performa model dan menentukan konfigurasi validasi yang paling optimal terhadap prediksi kebutuhan konsultasi kesehatan mental pada pekerja *remote*.

Proses pengujian dilakukan dengan membagi 1.100 data pekerja *remote* menjadi 80% data latih dan 20% data uji pada setiap iterasi validasi. Dengan demikian, setiap pengujian melibatkan sekitar 880 data latih dan 220 data uji. Distribusi kelas yang seimbang, yaitu 557 *Yes* dan 543 *No*, memastikan hasil pengujian lebih stabil dan tidak bias terhadap salah satu kelas.

Metode *k*-Fold *Cross Validation* memberikan dua keuntungan utama, yaitu meningkatkan reliabilitas hasil pengujian karena setiap data digunakan baik sebagai data latih maupun data uji, serta mengurangi risiko bias akibat pembagian data acak tunggal. Nilai rata-rata akurasi, presisi, *recall*, dan *f1-score* dari setiap iterasi digunakan untuk menilai performa model secara keseluruhan. Selain menghasilkan prediksi terhadap status kebutuhan konsultasi kesehatan mental (output utama model), tahap pengujian ini juga menghasilkan perbandingan

performa model berdasarkan variasi nilai k , yang menjadi output kedua dari penelitian ini.

3.10 Evaluasi Model

Tahap evaluasi model dilakukan untuk menilai sejauh mana algoritma *Random Forest Classifier* mampu memprediksi kebutuhan konsultasi kesehatan mental pada pekerja *remote* secara akurat dan konsisten. Evaluasi ini penting untuk mengukur keandalan model dalam mengenali pola data serta meminimalkan kesalahan klasifikasi. Hasil evaluasi juga digunakan untuk menentukan nilai *k-Fold Cross Validation* terbaik antara $k = 5$ dan $k = 10$, sesuai dengan rumusan masalah penelitian. Hasil dari setiap skenario kemudian dievaluasi menggunakan metrik akurasi (*accuracy*), presisi (*precision*), *recall*, *F1-score*, dan *ROC-AUC* untuk menilai performa model secara menyeluruh. Berikut penjelasan lebih lengkapnya.

3.10.1 Strategi Validasi (*K-Fold Cross Validation* dan *Out-of-Bag Score*)

Selain pembagian data menggunakan *train-test split*, penelitian ini menerapkan *K-Fold Cross Validation* untuk memperoleh estimasi performa yang lebih reliabel. Dataset dibagi menjadi k lipatan (*fold*) berukuran hampir sama. Pada setiap iterasi, satu lipatan digunakan sebagai data uji, sedangkan $(k-1)$ lipatan lainnya digunakan sebagai data latih. Proses ini diulang sebanyak k kali hingga seluruh data minimal sekali menjadi data uji. Nilai akurasi dari setiap iterasi kemudian dirata-ratakan untuk memperoleh estimasi performa model yang lebih stabil.

Dalam penelitian ini digunakan dua konfigurasi nilai k , yaitu $k = 5$ dan $k = 10$, untuk membandingkan performa model pada kedua kondisi tersebut. Nilai $k = 5$ dipilih untuk menjaga keseimbangan antara bias dan varians, sedangkan $k = 10$ digunakan untuk mendapatkan estimasi performa yang lebih mendetail.

Selain itu, penelitian ini juga mengevaluasi *Out-of-Bag (OOB) Score*, yang merupakan fitur bawaan dari algoritma *Random Forest*. *OOB score* diperoleh dari data yang tidak digunakan untuk melatih setiap *tree* selama proses *bootstrap sampling*, sehingga dapat berfungsi sebagai validasi internal tanpa memerlukan data uji tambahan. Kombinasi antara *K-Fold Cross Validation* dan *OOB score* memberikan gambaran evaluasi yang lebih komprehensif, serta membantu memastikan reliabilitas hasil prediksi model sebelum digunakan pada data baru.

3.10.2 Confusion Matrix

Confusion Matrix merupakan metode evaluasi yang umum digunakan untuk mengukur kinerja algoritma klasifikasi. Matriks ini berbentuk tabel dua dimensi yang memperlihatkan jumlah prediksi yang benar dan salah dari model terhadap setiap kelas dalam dataset. Dalam penelitian ini, karena model digunakan untuk klasifikasi biner, confusion matrix terdiri dari empat komponen utama, yaitu:

- a. True Positive (TP): jumlah data yang benar-benar termasuk kelas positif dan berhasil diprediksi dengan benar oleh model.
- b. True Negative (TN): jumlah data yang termasuk kelas negatif dan berhasil diprediksi dengan benar sebagai negatif.

- c. False Positive (FP): jumlah data yang sebenarnya negatif tetapi salah diklasifikasikan sebagai positif oleh model (*false alarm* atau kesalahan tipe I).
- d. False Negative (FN): jumlah data yang sebenarnya positif namun salah diklasifikasikan sebagai negatif (kesalahan tipe II).

Tabel 3.6 berikut menggambarkan struktur dasar dari *Confusion Matrix* yang digunakan dalam penelitian ini

Tabel 3.7 Confusion Matrix

<i>Actual Class</i>	<i>Predicted Positive</i>	<i>Predicted Negative</i>
<i>Positive (Yes)</i>	<i>True Positive (TP)</i>	<i>False Negative (FN)</i>
<i>Negative (No)</i>	<i>False Positive (FP)</i>	<i>True Negative (TN)</i>

Nilai-nilai dari Confusion Matrix tersebut menjadi dasar dalam menghitung beberapa metrik evaluasi yang digunakan untuk menilai kinerja model.

- a) Akurasi menunjukkan tingkat ketepatan model dalam melakukan prediksi terhadap seluruh data uji. Semakin tinggi nilai akurasi, semakin baik kemampuan model dalam mengklasifikasikan data dengan benar.

Persamaan akurasi ditunjukkan sebagai berikut:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad 3.5$$

- b) Presisi digunakan untuk mengukur seberapa tepat model dalam memprediksi kelas positif. Pada konteks penelitian ini, presisi menunjukkan seberapa banyak pekerja yang diprediksi membutuhkan konsultasi kesehatan mental benar-benar membutuhkan konsultasi tersebut. Persamaannya adalah:

$$Precision = \frac{TP}{TP + FP} \quad 3.6$$

- c) Recall digunakan untuk mengukur sejauh mana model mampu menemukan semua data yang seharusnya termasuk dalam kelas positif. Nilai recall yang tinggi menunjukkan bahwa model memiliki kemampuan yang baik dalam mengenali pekerja yang benar-benar membutuhkan konsultasi kesehatan mental. Persamaannya ditunjukkan sebagai berikut:

$$Recall = \frac{TP}{TP + FN} \quad 3.7$$

- d) F1-Score merupakan metrik gabungan yang memperhitungkan keseimbangan antara nilai presisi dan recall. Metrik ini penting untuk menilai model pada kondisi data yang seimbang seperti dalam penelitian ini, di mana terdapat 557 data “Yes” dan 543 data “No”. Persamaannya adalah:

$$F1\ Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad 3.8$$

Keempat metrik tersebut memberikan gambaran menyeluruh mengenai kinerja model. Akurasi menunjukkan seberapa banyak klasifikasi yang benar secara keseluruhan, presisi mengukur ketepatan model dalam mengenali kelas positif, recall menilai kemampuan model dalam menemukan semua data positif, dan F1-Score menyeimbangkan kedua metrik tersebut. Penggunaan Confusion Matrix dan metrik evaluasi ini sangat penting dalam penelitian berbasis data riil seperti klasifikasi kondisi kesehatan mental pekerja *remote*. Dengan menganalisis nilai-nilai tersebut, dapat diketahui sejauh mana model *Random Forest* mampu

memberikan hasil klasifikasi yang akurat serta mendeteksi pola yang relevan tanpa kehilangan data penting.

BAB IV

HASIL DAN PEMBAHASAN

4.1 Data Penelitian

Dataset yang digunakan dalam penelitian ini merupakan dataset *Mental Health in Tech Survey 2016* yang diperoleh melalui platform Kaggle dan dipublikasikan oleh organisasi Open Sourcing Mental Illness (OSMI). Dataset ini digunakan untuk melakukan klasifikasi kebutuhan konsultasi kesehatan mental pada pekerja remote. Berdasarkan faktor kondisi kesehatan mental, lingkungan kerja, serta dukungan perusahaan terhadap kesehatan psikologis pekerja.

Variabel target dalam penelitian ini adalah treatment, yang menunjukkan status apakah responden pernah melakukan konsultasi atau mendapatkan penanganan dari tenaga profesional kesehatan mental. Variabel tersebut terdiri dari dua kelas, yaitu:

- a. 0 (Tidak melakukan konsultasi), menunjukkan responden belum pernah melakukan konsultasi dengan tenaga profesional kesehatan mental.
- b. 1 (Melakukan konsultasi), menunjukkan responden pernah melakukan konsultasi dengan tenaga profesional kesehatan mental.

Dengan demikian, penelitian ini termasuk dalam klasifikasi biner (*binary classification*) karena proses klasifikasi dilakukan untuk menentukan dua kategori keluaran. Dataset awal terdiri dari 1.433 responden dengan 63 atribut berupa pertanyaan survei mengenai kondisi kesehatan mental pekerja di bidang teknologi. Selanjutnya, dilakukan proses penyaringan terhadap responden yang bekerja secara remote dengan kategori *Always* dan *Sometimes*, serta pemilihan atribut

yang relevan terhadap penelitian. Fitur yang digunakan dalam proses klasifikasi terdiri dari atribut *age*, *gender*, *country*, *self_employed*, *remote_work*, *family_history*, *work_interfere*, *no_employees*, *tech_company*, *benefits*, *care_options*, *wellness_program*, *seek_help*, *anonymity*, *leave*, *mental_health_consequence*, *coworkers*, *supervisor*, dan *mental_health_interview*, dengan *treatment* sebagai variabel target.

Berdasarkan analisis awal, dataset memiliki beberapa karakteristik sebagai berikut:

- a. Dataset terdiri atas atribut numerik dan kategorikal sehingga diperlukan proses transformasi data ke dalam bentuk numerik.
- b. Dataset masih memiliki nilai kosong, variasi penulisan data, serta beberapa data yang tidak sesuai sehingga diperlukan proses pembersihan data.
- c. Dataset mengandung data yang tidak sesuai dengan fokus penelitian, seperti responden non-remote dan data duplikat sehingga dilakukan proses penyaringan data.

Karakteristik tersebut menjadi dasar dilakukannya tahap *preprocessing* yang meliputi *cleaning data*, *label encoding*, dan normalisasi menggunakan metode *Min-Max Normalization* sebelum dilakukan proses pembentukan model menggunakan algoritma *Random Forest*.

4.2 Hasil Preprocessing Data

Tahap *preprocessing* data dilakukan untuk menyiapkan dataset agar dapat digunakan pada proses pembentukan model *Random Forest*. Dataset awal masih mengandung berbagai bentuk data seperti atribut kategorikal, variasi penulisan

data, nilai kosong, serta data yang tidak sesuai dengan fokus penelitian. Oleh karena itu, dilakukan beberapa tahapan *preprocessing* yang meliputi proses *cleaning data*, *label encoding*, dan normalisasi data.

4.2.1 Cleaning Data

Tahap *cleaning* data dilakukan untuk meningkatkan kualitas data sehingga dataset yang digunakan sesuai dengan kebutuhan penelitian dan dapat diproses pada tahap pembentukan model. Proses ini meliputi penyaringan data, pemilihan atribut, standarisasi data, penanganan nilai kosong (*missing value*), perbaikan nilai yang tidak valid, serta penghapusan data duplikat. Tahap awal dilakukan dengan melakukan penyaringan data berdasarkan atribut *remote_work*.

Pada penelitian ini hanya digunakan responden yang bekerja secara *remote* dengan kategori *Always* dan *Sometimes*, sedangkan responden dengan kategori *Never* dihapus karena tidak sesuai dengan fokus penelitian. Setelah proses penyaringan dilakukan, tahap berikutnya adalah pemilihan atribut yang relevan terhadap penelitian. Dari total 63 atribut pada dataset awal, dipilih 21 atribut yang berkaitan dengan kondisi kesehatan mental dan lingkungan kerja responden. Selanjutnya, atribut *state* tidak digunakan dalam proses analisis sehingga dihapus dan menghasilkan 20 atribut yang digunakan pada tahap berikutnya.

Selanjutnya dilakukan standarisasi pada atribut *gender* untuk mengatasi ketidakkonsistenan bentuk penulisan jawaban responden. Berbagai variasi data pada atribut tersebut dikelompokkan ke dalam tiga kategori utama, yaitu *Male*, *Female*, dan *Other*, sehingga menghasilkan format data yang lebih seragam dan konsisten. Tahap berikutnya adalah penanganan nilai kosong (*missing value*) yang

terdapat pada beberapa atribut. Nilai kosong pada atribut *self_employed* dan *work_interfere* diisi menggunakan nilai modus dari masing-masing atribut.

Selain itu, atribut yang berkaitan dengan kondisi perusahaan atau pemberi kerja juga dilakukan pengisian menggunakan nilai modus agar tidak terdapat nilai kosong yang dapat mengganggu proses pengolahan data pada tahap selanjutnya. Pada atribut *age*, dilakukan pemeriksaan terhadap data usia yang kosong serta nilai usia yang berada di luar rentang yang telah ditentukan, yaitu 18 hingga 72 tahun. Data usia yang tidak valid kemudian digantikan menggunakan nilai median dari data usia yang valid. Penggunaan nilai median bertujuan untuk mengurangi pengaruh nilai ekstrem sehingga distribusi usia tetap dapat merepresentasikan karakteristik responden.

4.2.2 *Label Encoding*

Tahap *label encoding* dilakukan setelah proses *cleaning* data selesai dilakukan dan bertujuan untuk mengubah atribut yang masih berbentuk kategorikal menjadi nilai numerik agar dapat digunakan dalam proses pembentukan model menggunakan algoritma Random Forest. Pada tahap ini, setiap kategori pada suatu atribut diberikan kode numerik yang berbeda sesuai dengan karakteristik nilai yang dimiliki. Atribut yang mengalami proses *label encoding* meliputi *gender*, *country*, *remote_work*, *family_history*, *work_interfere*, *no_employees*, *benefits*, *care_options*, *wellness_program*, *seek_help*, *anonymity*, *leave*, *mental_health_consequence*, *coworkers*, *supervisor*, dan *mental_health_interview*. Sementara itu, atribut yang telah memiliki bentuk

numerik seperti *self_employed*, *tech_company*, dan *treatment* tetap dipertahankan sesuai dengan nilai aslinya.

Proses *label encoding* dilakukan dengan membuat representasi angka untuk setiap kategori pada atribut sehingga seluruh data memiliki format yang seragam dan dapat diproses oleh algoritma machine learning. Sebagai contoh, atribut *gender* dikonversi menjadi tiga kode utama yaitu *Male* menjadi 0, *Female* menjadi 1, dan *Other* menjadi 2. Selain itu, atribut *remote_work* dikonversi menjadi kode 0 untuk kategori *Sometimes* dan kode 1 untuk kategori *Always*. Hasil dataset setelah dilakukan proses *label encoding* dapat dilihat pada tabel 4.1 berikut.

Tabel 4.1 Hasil Encoding

<i>Gender</i>	<i>Country</i>	<i>Remote Work</i>	<i>Family History</i>	<i>Work Interfere</i>	<i>Treatment</i>
0	45	0	0	0	0
0	45	1	0	0	1
0	45	0	0	3	1
1	46	0	1	3	1
0	45	0	0	0	1

Keterangan kode:

Gender: 0 = *Male*, 1 = *Female*, dan 2 = *Other*.

Remote Work: 0 = *Sometimes* dan 1 = *Always*.

Family History: 0 = *No*, 1 = *Yes*, dan 2 = *I don't know*.

Work Interfere: 0 = *Not applicable to me*, 1 = *Never*, 2 = *Rarely*, 3 = *Sometimes*, dan 4 = *Often*.

Treatment: menggunakan nilai asli, yaitu 0 dan 1, dengan 0 menunjukkan responden tidak pernah melakukan konsultasi dan 1 menunjukkan responden pernah melakukan konsultasi dengan tenaga profesional kesehatan mental.

Country: dikonversi menjadi kode numerik berdasarkan nilai unik setiap negara yang terdapat pada dataset.

Berdasarkan hasil proses *label encoding*, seluruh atribut kategorikal pada dataset telah berhasil dikonversi ke dalam bentuk numerik dengan tetap mempertahankan informasi dari setiap kategori yang ada. Proses pengkodean tersebut menghasilkan struktur data yang lebih terstandarisasi sehingga dapat

diproses oleh algoritma Random Forest pada tahap pembentukan model. Selain itu, atribut *country* dikonversi menjadi kode numerik berdasarkan nilai unik negara yang terdapat pada dataset. Proses pemberian kode dilakukan secara otomatis berdasarkan daftar negara yang tersedia pada dataset, sehingga setiap negara memiliki representasi kode numerik yang berbeda. Dataset hasil label *encoding* selanjutnya digunakan pada tahap normalisasi untuk menyesuaikan rentang nilai atribut *age* sebelum dilakukan proses pemodelan.

4.2.3 Normalisasi Data

Tahap normalisasi data dilakukan setelah proses *label encoding* selesai dilakukan sebagai bagian akhir dari tahapan *preprocessing* data. Pada penelitian ini, normalisasi hanya diterapkan pada atribut *age* karena atribut tersebut memiliki rentang nilai yang lebih besar dibandingkan atribut lainnya yang telah direpresentasikan dalam bentuk kode numerik. Proses normalisasi dilakukan menggunakan metode *Min-Max Normalization* dengan mengubah nilai usia ke dalam rentang nilai antara 0 hingga 1 berdasarkan nilai minimum dan maksimum dari data setelah melalui tahap *cleaning data*. Dengan dilakukannya normalisasi, perbedaan skala pada atribut *age* dapat dikurangi sehingga data menjadi lebih seragam dan sesuai untuk digunakan pada proses pembentukan model *Random Forest*.

Hasil dari proses normalisasi menunjukkan bahwa atribut *age* yang sebelumnya berupa nilai usia asli telah berhasil dikonversi menjadi nilai desimal dengan rentang antara 0 hingga 1. Perubahan tersebut bertujuan untuk menghasilkan skala data yang lebih konsisten tanpa mengubah pola atau informasi

yang terkandung dalam atribut tersebut. Contoh hasil perubahan nilai atribut *age* sebelum dan sesudah proses normalisasi ditampilkan pada tabel berikut.

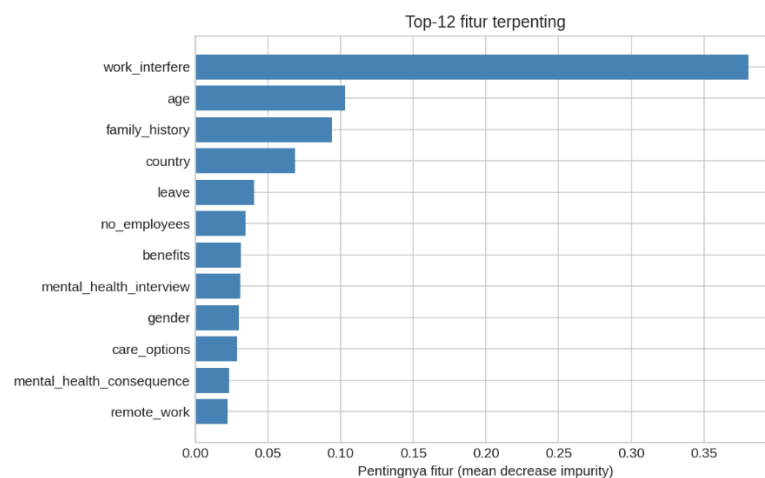
Tabel 4.2 Hasil Normalisasi Atribut Age

<i>Age</i> Sebelum Normalisasi	<i>Age</i> Setelah Normalisasi
38	0,372549
39	0,392157
42	0,450980
43	0,470588
43	0,470588

Hasil tersebut menunjukkan bahwa proses normalisasi berhasil mengubah data usia menjadi bentuk yang lebih sesuai untuk digunakan pada tahap pemodelan. Berdasarkan hasil normalisasi yang telah dilakukan, seluruh tahapan *preprocessing* pada dataset telah berhasil diselesaikan. Dataset yang dihasilkan telah memiliki atribut kategorikal dalam bentuk numerik melalui proses *label encoding* serta atribut *age* yang telah memiliki skala nilai yang seragam. Dengan demikian, dataset telah siap digunakan pada tahap selanjutnya.

4.3 Hasil Feature Selection

Tahap seleksi fitur dilakukan setelah seluruh proses *preprocessing* data selesai dilakukan untuk mengetahui tingkat pengaruh masing-masing atribut terhadap prediksi kebutuhan konsultasi kesehatan mental pada pekerja *remote*. Proses seleksi fitur menggunakan metode *feature importance* pada algoritma Random Forest yang mengukur kontribusi setiap atribut dalam proses pembentukan model. Nilai *feature importance* yang lebih tinggi menunjukkan bahwa atribut tersebut memiliki pengaruh yang lebih besar dalam membantu model melakukan klasifikasi terhadap variabel target *treatment*. Hasil seleksi fitur pada penelitian ini ditampilkan pada gambar 4.1 berikut.



Gambar 4.1 Hasil Feature Importance

Berdasarkan Gambar 4.1, atribut *work_interfere* memiliki nilai *feature importance* tertinggi dibandingkan atribut lainnya sehingga menjadi atribut yang paling berpengaruh dalam proses prediksi kebutuhan konsultasi kesehatan mental. Hal ini menunjukkan bahwa tingkat gangguan pekerjaan yang disebabkan oleh kondisi kesehatan mental memiliki hubungan yang kuat terhadap keputusan seseorang untuk melakukan konsultasi kesehatan mental. Selain itu, atribut *age*, *family_history*, dan *country* juga memiliki kontribusi yang cukup besar terhadap proses pembentukan model karena memiliki nilai *feature importance* yang relatif tinggi. Hasil tersebut menunjukkan bahwa karakteristik individu serta riwayat kesehatan mental menjadi faktor yang perlu diperhatikan dalam proses prediksi.

Atribut lainnya seperti *leave*, *no_employees*, *benefits*, *gender*, *care_options*, *mental_health_interview*, *mental_health_consequence*, dan *remote_work* memiliki nilai *feature importance* yang lebih rendah dibandingkan atribut utama, namun tetap memberikan kontribusi terhadap kinerja model *Random Forest*. Perbedaan tingkat kepentingan setiap atribut menunjukkan bahwa

tidak seluruh variabel memiliki pengaruh yang sama terhadap prediksi kebutuhan konsultasi kesehatan mental. Meskipun beberapa atribut memiliki nilai pengaruh yang kecil, atribut tersebut tetap digunakan karena masih menyimpan informasi yang dapat membantu model dalam mengenali pola pada data. Dengan demikian, hasil seleksi fitur menjadi dasar dalam proses pembentukan model *Random Forest*.

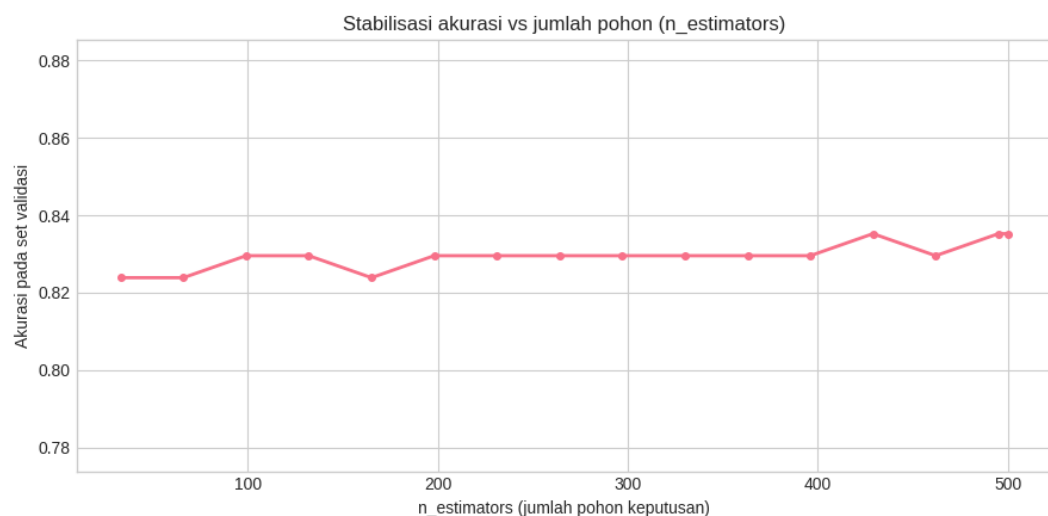
4.4 Pembentukan Model Random Forest

Pembentukan model Random Forest dilakukan menggunakan dataset hasil *preprocessing* dan hasil seleksi fitur yang telah diperoleh pada tahap sebelumnya. Dataset tersebut terdiri dari atribut prediktor sebagai variabel masukan *input* dan atribut *treatment* sebagai variabel target untuk melakukan prediksi kebutuhan konsultasi kesehatan mental pada pekerja di bidang teknologi. Algoritma Random Forest membangun sejumlah pohon keputusan (*decision tree*) melalui teknik *bootstrap sampling*, yaitu proses pengambilan sampel data secara acak dengan pengembalian dari data latih. Setiap pohon keputusan dibentuk menggunakan sebagian data dan atribut yang dipilih secara acak sehingga menghasilkan variasi pola pembelajaran yang dapat meningkatkan kemampuan generalisasi model.

Pada tahap pembentukan model, digunakan konfigurasi hyperparameter Random Forest yang telah ditentukan pada Tabel 3.6. Konfigurasi tersebut meliputi parameter jumlah pohon keputusan (*n_estimators*), jumlah fitur maksimum yang digunakan pada setiap proses pemisahan node (*max_features*), kriteria perhitungan *impurity*, serta beberapa parameter pendukung lainnya dalam proses pembentukan model. Penggunaan konfigurasi *hyperparameter* tersebut

bertujuan untuk menghasilkan model yang sesuai dengan karakteristik dataset dan memiliki performa klasifikasi yang baik. Salah satu parameter utama yang diperhatikan dalam penelitian ini adalah jumlah pohon keputusan karena parameter tersebut berpengaruh terhadap kestabilan performa model Random Forest.

Selanjutnya, dilakukan pengujian terhadap beberapa variasi jumlah pohon keputusan ($n_estimators$) untuk mengamati perubahan nilai akurasi model. Pengujian dilakukan menggunakan data validasi internal dengan membandingkan nilai akurasi yang dihasilkan dari setiap variasi jumlah pohon keputusan. Hasil pengamatan hubungan antara jumlah pohon keputusan dengan nilai akurasi model Random Forest ditampilkan pada gambar berikut.



Gambar 4.2 Stabilisasi Akurasi Jumlah Pohon Dari Proses Pelatihan Model

Berdasarkan gambar tersebut, peningkatan jumlah pohon keputusan memberikan pengaruh terhadap perubahan nilai akurasi model pada tahap awal pembentukan Random Forest. Nilai akurasi mengalami peningkatan seiring bertambahnya jumlah pohon hingga mencapai kondisi yang relatif stabil. Pada

penelitian ini, jumlah pohon keputusan sebanyak 500 dipilih sebagai konfigurasi akhir karena menghasilkan nilai akurasi yang stabil tanpa adanya peningkatan performa yang signifikan setelah jumlah pohon terus ditambahkan. Pemilihan nilai $n_estimators$ tersebut bertujuan untuk memperoleh keseimbangan antara performa klasifikasi model dan efisiensi proses komputasi.

Pada tahap akhir pembentukan model Random Forest, dilakukan proses agregasi hasil prediksi yang dikenal sebagai *Bootstrap Aggregating* atau *Bagging*. Pada proses ini, setiap pohon keputusan yang telah dibangun akan menghasilkan prediksi secara independen terhadap data masukan, kemudian hasil prediksi tersebut digabungkan menggunakan metode *majority voting* untuk menentukan kelas akhir berdasarkan jumlah suara terbanyak. Dalam penelitian ini, hasil prediksi berupa kelas 0 yang menunjukkan responden tidak pernah melakukan konsultasi kesehatan mental dan kelas 1 yang menunjukkan responden pernah melakukan konsultasi kesehatan mental. Mekanisme *bagging* dan *majority voting* memungkinkan Random Forest menghasilkan prediksi yang lebih stabil, mengurangi variansi model, serta meminimalkan risiko *overfitting* yang umum terjadi pada penggunaan satu pohon keputusan.

4.5 Evaluasi Model

Evaluasi model dilakukan untuk mengetahui kemampuan model Random Forest dalam melakukan klasifikasi kebutuhan konsultasi kesehatan mental serta memastikan kestabilan performa model yang telah dibentuk. Pada penelitian ini, evaluasi model dilakukan menggunakan metode *Stratified K-Fold Cross Validation* dan *Out-of-Bag (OOB) Score* sebagai metode validasi internal pada

algoritma Random Forest. Kedua metode tersebut digunakan untuk mengukur kemampuan generalisasi model berdasarkan data pelatihan sebelum dilakukan pengujian menggunakan data testing. Hasil evaluasi pada setiap skenario pembagian data digunakan sebagai dasar untuk mengetahui konsistensi performa model dalam proses pelatihan model.

4.5.1 Evaluasi Model Dengan *Stratified K-Fold Cross Validation*

Metode *Stratified K-Fold Cross Validation* digunakan untuk mengevaluasi kestabilan model dengan membagi data pelatihan ke dalam beberapa bagian (*fold*) dengan tetap mempertahankan proporsi kelas pada setiap pembagian data. Pada penelitian ini digunakan dua variasi jumlah *fold*, yaitu $K=5$ dan $K=10$ pada tiga skenario pembagian data, yaitu 60:40, 70:30, dan 80:20. Evaluasi dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, *F1-score*, dan *ROC-AUC* dengan menghitung nilai rata-rata (*mean*) dan standar deviasi (*standard deviation*) dari setiap proses validasi. Hasil evaluasi *Stratified K-Fold Cross Validation* dapat dilihat pada tabel berikut.

Tabel 4.3 Hasil Evaluasi Stratified K-Fold Cross Validation

Pembagian Data	K-Fold	Accuracy (Mean $\pm$ Std)	Precision (Mean $\pm$ Std)	Recall (Mean $\pm$ Std)	F1-Score (Mean $\pm$ Std)	ROC-AUC (Mean $\pm$ Std)
60:40	5	0,8432 $\pm$ 0,0232	0,8596 $\pm$ 0,0345	0,8813 $\pm$ 0,0433	0,8689 $\pm$ 0,0193	0,8835 $\pm$ 0,0288
70:30	5	0,8459 $\pm$ 0,0169	0,8590 $\pm$ 0,0109	0,8852 $\pm$ 0,0374	0,8714 $\pm$ 0,0164	0,8784 $\pm$ 0,0191
80:20	5	0,8561 $\pm$ 0,0152	0,8643 $\pm$ 0,0027	0,8977 $\pm$ 0,0281	0,8805 $\pm$ 0,0142	0,8885 $\pm$ 0,0127
60:40	10	0,8476 $\pm$ 0,0385	0,8608 $\pm$ 0,0437	0,8891 $\pm$ 0,0529	0,8731 $\pm$ 0,0315	0,8850 $\pm$ 0,0460
70:30	10	0,8433 $\pm$ 0,0297	0,8566 $\pm$ 0,0340	0,8854 $\pm$ 0,0448	0,8697 $\pm$ 0,0257	0,8758 $\pm$ 0,0367
80:20	10	0,8539 $\pm$ 0,0262	0,8616 $\pm$ 0,0255	0,8978 $\pm$ 0,0373	0,8788 $\pm$ 0,0233	0,8848 $\pm$ 0,0420

Berdasarkan tabel di atas, hasil evaluasi menggunakan *Stratified K-Fold Cross Validation* menunjukkan bahwa seluruh skenario pembagian data menghasilkan performa yang baik dengan nilai *accuracy*, *precision*, *recall*, *F1-score*, dan *ROC-AUC* yang relatif tinggi. Skenario pembagian data 80:20 dengan K-Fold 5 memperoleh performa terbaik berdasarkan nilai *accuracy* sebesar 0,8561, *F1-score* sebesar 0,8805, dan *ROC-AUC* sebesar 0,8885 dengan nilai standar deviasi yang relatif rendah. Hal tersebut menunjukkan bahwa model memiliki tingkat konsistensi yang baik selama proses validasi pada data pelatihan. Dengan demikian, hasil *Stratified K-Fold Cross Validation* menunjukkan bahwa model Random Forest memiliki kemampuan pembelajaran dan generalisasi yang baik.

4.5.2 Evaluasi Model Dengan *Out-of-Bag (OOB) Score*

Selain menggunakan *Stratified K-Fold Cross Validation*, evaluasi model juga dilakukan menggunakan metode *Out-of-Bag (OOB) Score* sebagai mekanisme validasi internal pada algoritma *Random Forest*. Nilai OOB Score diperoleh dari data yang tidak terpilih selama proses *bootstrap sampling* pada pembentukan masing-masing pohon keputusan. Data tersebut digunakan untuk mengukur kemampuan model dalam melakukan klasifikasi terhadap data yang tidak terlibat secara langsung dalam proses pelatihan. Hasil evaluasi OOB Score untuk setiap skenario pembagian data ditampilkan pada Tabel 4.4.

Tabel 4.4 Hasil Evaluasi *Out-of-Bag (OOB) Score*

Pembagian Data	K-Fold	Nilai OOB Score
60:40	5	0,8463
60:40	10	0,8463
70:30	5	0,8486
70:30	10	0,8486

Pembagian Data	K-Fold	Nilai OOB Score
80:20	5	0,8527
80:20	10	0,8527

Berdasarkan tabel di atas, hasil evaluasi menggunakan *Out-of-Bag (OOB) Score* menunjukkan bahwa seluruh skenario pembagian data memiliki kemampuan generalisasi yang baik dengan nilai OOB Score berada pada rentang 0,8463 hingga 0,8527. Skenario pembagian data 80:20 memperoleh nilai OOB Score tertinggi sebesar 0,8527, yang menunjukkan bahwa model memiliki kemampuan klasifikasi yang lebih baik terhadap data yang tidak digunakan selama proses pembentukan pohon keputusan. Perbedaan nilai OOB Score antar skenario tidak menunjukkan selisih yang signifikan sehingga dapat disimpulkan bahwa model Random Forest memiliki performa yang cukup stabil pada berbagai skenario pembagian data.

4.6 Hasil Skenario Uji Coba

Pada tahap ini dilakukan pengujian model Random Forest menggunakan tiga skenario pembagian data, yaitu 60:40, 70:30, dan 80:20 dengan variasi nilai K pada metode *Stratified K-Fold Cross Validation* sebesar 5 dan 10. Pengujian dilakukan menggunakan data testing untuk mengetahui kemampuan model dalam melakukan prediksi terhadap kebutuhan konsultasi kesehatan mental. Hasil pengujian dianalisis menggunakan nilai accuracy, precision, recall, F1-score, dan ROC-AUC, serta didukung dengan analisis Confusion Matrix dan kurva ROC pada setiap skenario pengujian.

4.6.1 Alur Pengujian Model

Pada setiap skenario, proses pengujian model dilakukan melalui beberapa tahapan sebagai berikut:

1. PembagianData

Dataset hasil preprocessing dan feature selection dibagi menjadi data latih dan data uji menggunakan tiga skenario pembagian data, yaitu 60:40, 70:30, dan 80:20.

2. PembentukanModelRandomForest

Data latih digunakan untuk membentuk model Random Forest menggunakan konfigurasi hyperparameter yang telah ditentukan sebelumnya.

3. EvaluasiModel

Model dievaluasi menggunakan metode Stratified K-Fold Cross Validation dengan variasi nilai K sebesar 5 dan 10 serta Out-of-Bag (OOB) Score untuk mengetahui kestabilan performa model.

4. PrediksiDataTesting

Model yang telah dibentuk digunakan untuk melakukan prediksi terhadap data testing pada setiap skenario pembagian data.

5. EvaluasiHasilPengujian

Hasil prediksi data testing dievaluasi menggunakan metrik accuracy, precision, recall, F1-score, ROC-AUC, Confusion Matrix, dan kurva ROC.

4.6.2 Skenario 1

Pada skenario pengujian pertama, dataset dibagi menjadi 60% data latih dan 40% data uji. Pengujian dilakukan menggunakan dua variasi *Stratified K-Fold Cross Validation*, yaitu K=5 dan K=10 untuk melihat pengaruh jumlah *fold* terhadap performa model Random Forest dalam melakukan klasifikasi kondisi kesehatan mental berdasarkan data yang telah melalui tahapan *preprocessing* dan *feature selection*.

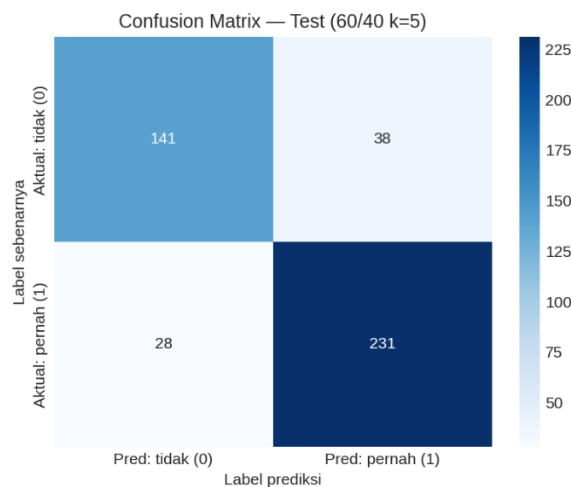
Tabel 4.5 Hasil Evaluasi Model Skenario 1

Parameter	K=5	K=10
Accuracy	0,8493	0,8493
Precision	0,8587	0,8587
Recall	0,8919	0,8919
F1-Score	0,8750	0,8750
ROC-AUC	0,8695	0,8695

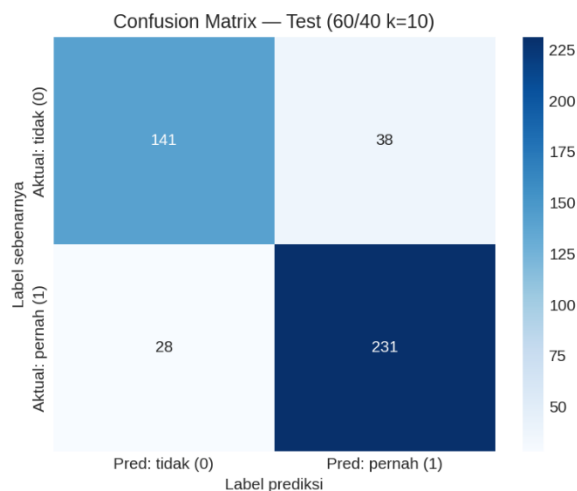
Berdasarkan Tabel 4.5, diperoleh hasil bahwa model Random Forest mampu memberikan performa klasifikasi yang baik pada kedua variasi *K-Fold* yang digunakan. Pengujian menggunakan K=5 maupun K=10 menghasilkan nilai evaluasi yang sama, yaitu accuracy sebesar 0,8493, precision sebesar 0,8587, recall sebesar 0,8919, F1-score sebesar 0,8750, dan ROC-AUC sebesar 0,8695. Nilai accuracy sebesar 84,93% menunjukkan bahwa model mampu mengklasifikasikan sebagian besar data dengan benar. Nilai precision sebesar 85,87% menunjukkan bahwa dari seluruh data yang diprediksi termasuk ke dalam kelas *pernah melakukan treatment*, sebesar 85,87% merupakan prediksi yang sesuai dengan kondisi sebenarnya.

Selain itu, nilai recall sebesar 89,19% menunjukkan bahwa model memiliki kemampuan yang baik dalam mengenali data yang benar-benar termasuk

dalam kelas *pernah melakukan treatment*. Keseimbangan antara nilai precision dan recall menghasilkan nilai F1-score sebesar 87,50%, yang mengindikasikan bahwa model memiliki performa yang baik dalam menangani klasifikasi kedua kelas. Nilai ROC-AUC sebesar 0,8695 menunjukkan bahwa model memiliki kemampuan diskriminasi yang baik dalam membedakan kelas *pernah melakukan treatment* dan *tidak pernah melakukan treatment*. Berdasarkan hasil tersebut, perubahan jumlah *fold* dari K=5 menjadi K=10 pada skenario pembagian data 60:40 tidak memberikan peningkatan maupun penurunan performa model, sehingga dapat disimpulkan bahwa model Random Forest memiliki kestabilan performa pada kedua konfigurasi validasi silang tersebut.



Gambar 4.3 Confusion Matrix (60:40 K 5)

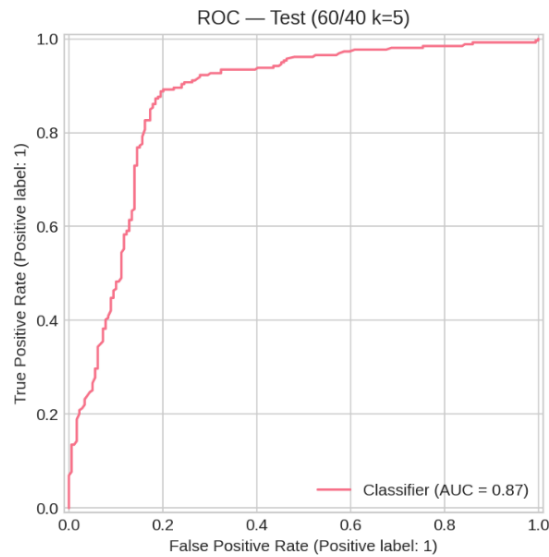


Gambar 4.4 Confusion Matrix(60:40 K 10)

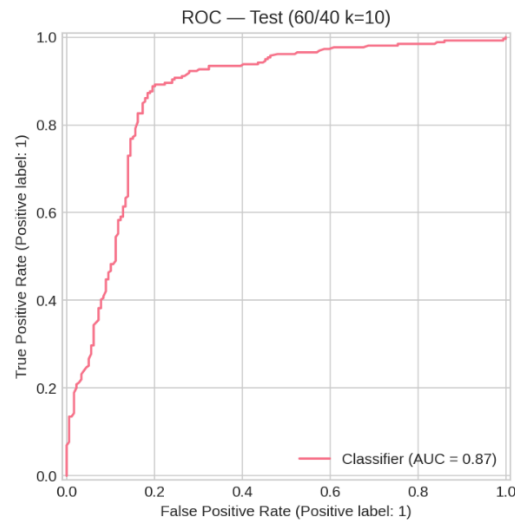
Berdasarkan hasil *confusion matrix* pada kedua pengujian, diperoleh nilai True Negative (TN) sebanyak 141 data, False Positive (FP) sebanyak 38 data, False Negative (FN) sebanyak 28 data, dan True Positive (TP) sebanyak 231 data. Hasil tersebut menunjukkan bahwa model memiliki kemampuan yang cukup baik dalam mengklasifikasikan kedua kelas, yaitu kelas tidak pernah melakukan *treatment* dan pernah melakukan *treatment*. Hal ini terlihat dari jumlah prediksi benar yang lebih besar dibandingkan jumlah kesalahan klasifikasi.

Nilai True Positive yang lebih tinggi menunjukkan bahwa model mampu mengenali sebagian besar data yang termasuk ke dalam kelas pernah melakukan *treatment*. Namun, masih terdapat kesalahan klasifikasi berupa False Positive sebanyak 38 data, yaitu data yang sebenarnya tidak pernah melakukan *treatment* tetapi diprediksi pernah melakukan *treatment*, serta False Negative sebanyak 28 data, yaitu data yang sebenarnya pernah melakukan *treatment* tetapi diprediksi tidak pernah melakukan *treatment*. Kesalahan tersebut mengindikasikan bahwa masih terdapat beberapa karakteristik antar kelas yang memiliki kemiripan

sehingga menyebabkan model belum mampu melakukan pemisahan kelas secara sempurna.



Gambar 4.5 Kurva ROC(60:40 K 5)



Gambar 4.6 Kurva ROC(60:40 K 10)

Berdasarkan kurva ROC yang diperoleh, baik pada pengujian menggunakan K=5 maupun K=10 menghasilkan nilai Area Under Curve (AUC) sebesar 0,87. Nilai tersebut menunjukkan bahwa model Random Forest memiliki kemampuan yang baik dalam membedakan antara kelas yang pernah melakukan

treatment dan tidak pernah melakukan *treatment*. Semakin tinggi nilai AUC mendekati angka 1, maka kemampuan model dalam melakukan pemisahan kedua kelas semakin baik.

Dengan diperolehnya nilai AUC sebesar 0,87, dapat disimpulkan bahwa model memiliki tingkat kemampuan diskriminasi yang baik serta mampu memberikan performa klasifikasi yang cukup optimal pada skenario pembagian data 60:40. Namun, keberadaan nilai *False Positive* dan *False Negative* pada *confusion matrix* menunjukkan bahwa masih terdapat beberapa data yang memiliki karakteristik serupa antar kelas sehingga menyebabkan terjadinya kesalahan prediksi.

4.6.3 Skenario 2

Pada skenario pengujian kedua, dataset dibagi menjadi 70% data latih dan 30% data uji. Pengujian dilakukan menggunakan dua variasi Stratified K-Fold Cross Validation, yaitu K=5 dan K=10, untuk mengetahui pengaruh jumlah *fold* terhadap performa model Random Forest dalam melakukan klasifikasi kondisi kesehatan mental berdasarkan data yang telah melalui tahapan *preprocessing* dan *feature selection*.

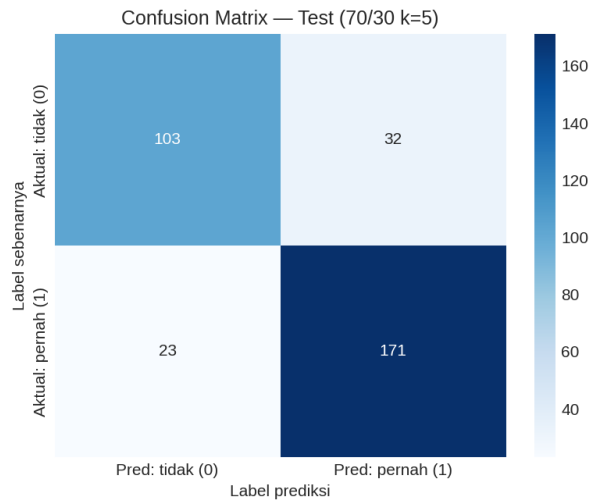
Tabel 4.6 Hasil Evaluasi Model Skenario 2

Parameter	K=5	K=10
Accuracy	0,8328	0,8328
Precision	0,8424	0,8424
Recall	0,8814	0,8814
F1-Score	0,8615	0,8615
ROC-AUC	0,8696	0,8696

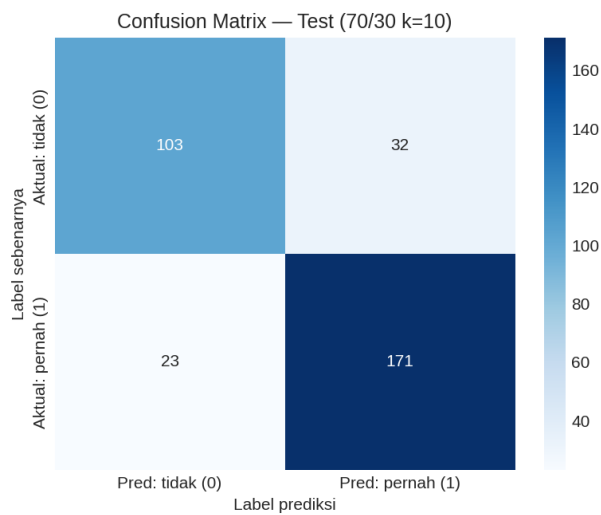
Berdasarkan Tabel 4.6, model Random Forest menunjukkan performa klasifikasi yang baik pada kedua variasi *K-Fold* yang digunakan. Pengujian

menggunakan $K=5$ maupun $K=10$ menghasilkan nilai evaluasi yang sama, yaitu accuracy sebesar 0,8328, precision sebesar 0,8424, recall sebesar 0,8814, F1-score sebesar 0,8615, dan ROC-AUC sebesar 0,8696. Nilai accuracy sebesar 83,28% menunjukkan bahwa model mampu melakukan klasifikasi dengan tingkat ketepatan yang baik terhadap data uji. Nilai precision sebesar 84,24% menunjukkan bahwa dari seluruh data yang diprediksi sebagai kelas *pernah melakukan treatment*, sebanyak 84,24% merupakan prediksi yang sesuai dengan kondisi sebenarnya.

Selain itu, nilai recall sebesar 88,14% menunjukkan bahwa model memiliki kemampuan yang baik dalam mengenali data yang benar-benar termasuk dalam kelas *pernah melakukan treatment*. Keseimbangan antara nilai *precision* dan *recall* menghasilkan nilai F1-score sebesar 86,15%, yang menunjukkan bahwa model memiliki performa yang cukup baik dalam menangani kedua kelas. Nilai ROC-AUC sebesar 0,8696 menunjukkan bahwa model memiliki kemampuan diskriminasi yang baik dalam membedakan kelas *pernah melakukan treatment* dan *tidak pernah melakukan treatment*. Berdasarkan hasil tersebut, perubahan jumlah *fold* dari $K=5$ menjadi $K=10$ tidak memberikan perubahan terhadap performa model, sehingga dapat disimpulkan bahwa model Random Forest memiliki tingkat kestabilan yang baik pada kedua konfigurasi validasi silang.



Gambar 4.7 Confusion Matrix(70:30 K 5)

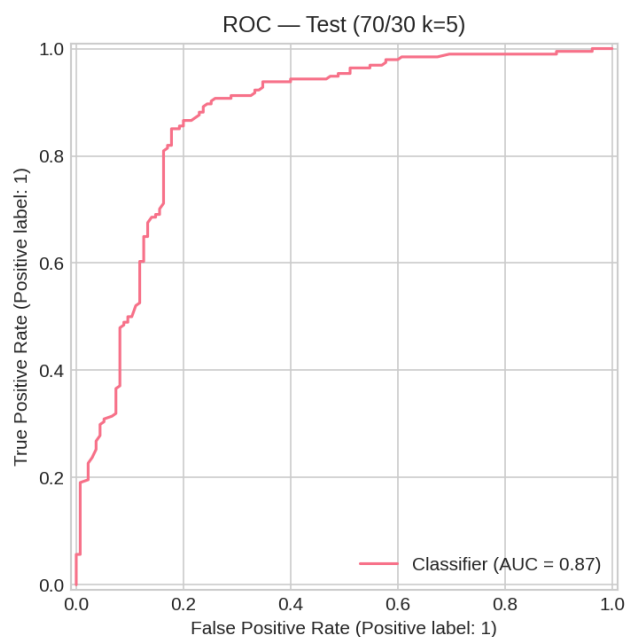


Gambar 4.8 Confusion Matrix (70:30 K 10)

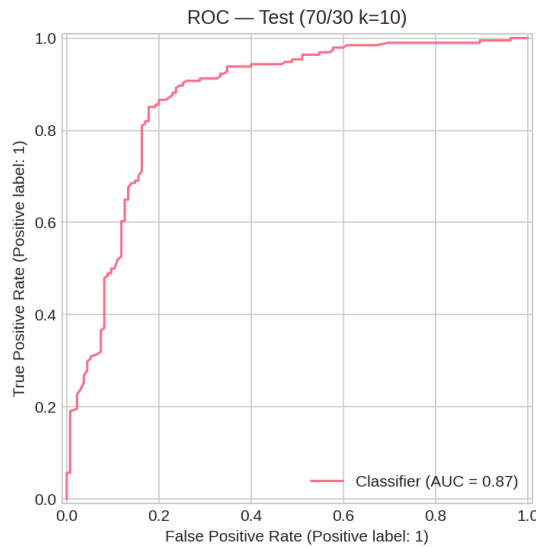
Berdasarkan hasil *confusion matrix* pada kedua pengujian, diperoleh nilai True Negative (TN) sebanyak 103 data, False Positive (FP) sebanyak 32 data, False Negative (FN) sebanyak 23 data, dan True Positive (TP) sebanyak 171 data. Hasil tersebut menunjukkan bahwa model memiliki kemampuan klasifikasi yang cukup baik karena jumlah prediksi yang benar lebih besar dibandingkan jumlah kesalahan klasifikasi. Sebanyak 103 data pada kelas tidak pernah melakukan

treatment berhasil diklasifikasikan dengan benar, sedangkan 171 data pada kelas pernah melakukan treatment juga berhasil dikenali dengan tepat oleh model.

Meskipun demikian, masih terdapat kesalahan klasifikasi berupa False Positive sebanyak 32 data, yaitu data yang sebenarnya termasuk kelas tidak pernah melakukan *treatment*, tetapi diprediksi sebagai pernah melakukan *treatment*. Selain itu, terdapat False Negative sebanyak 23 data, yaitu data yang sebenarnya pernah melakukan *treatment*, tetapi diprediksi sebagai tidak pernah melakukan *treatment*. Jumlah nilai True Positive yang jauh lebih besar dibandingkan False Negative menunjukkan bahwa model memiliki kemampuan yang baik dalam mendeteksi individu yang pernah melakukan *treatment*. Di sisi lain, keberadaan False Positive dan False Negative menunjukkan bahwa masih terdapat beberapa karakteristik yang memiliki kemiripan antar kelas sehingga menyebabkan model melakukan kesalahan prediksi pada sejumlah data.



Gambar 4.9 Kurva ROC(70:30 K 5)



Gambar 4.10 Kurva ROC(70:30 K 10)

Berdasarkan hasil kurva Receiver Operating Characteristic (ROC), diperoleh nilai Area Under Curve (AUC) sebesar 0,8696 atau dibulatkan menjadi 0,87 pada kedua pengujian, yaitu K=5 dan K=10. Kurva ROC yang berada jauh di atas garis diagonal menunjukkan bahwa model Random Forest memiliki kemampuan yang baik dalam membedakan kelas positif (*pernah melakukan treatment*) dan kelas negatif (*tidak pernah melakukan treatment*). Nilai AUC yang mendekati angka 1 mengindikasikan bahwa kemampuan diskriminasi model semakin baik, sehingga nilai AUC sebesar 0,8696 menunjukkan bahwa model mampu memberikan performa klasifikasi yang baik dan cukup optimal pada skenario pembagian data 70% data latih dan 30% data uji. Selain itu, kesamaan nilai AUC serta metrik evaluasi lainnya pada pengujian K=5 dan K=10 menunjukkan bahwa peningkatan jumlah *fold* pada proses Stratified K-Fold Cross Validation tidak memberikan pengaruh terhadap kemampuan diskriminasi maupun performa klasifikasi model pada skenario pengujian ini.

4.6.4 Skenario 3

Pada skenario pengujian ketiga, dataset dibagi menjadi 80% data latih dan 20% data uji. Pengujian dilakukan menggunakan dua variasi Stratified K-Fold Cross Validation, yaitu K=5 dan K=10, untuk mengetahui pengaruh jumlah *fold* terhadap performa model Random Forest dalam melakukan klasifikasi kondisi kesehatan mental berdasarkan data yang telah melalui tahapan *preprocessing* dan *feature selection*.

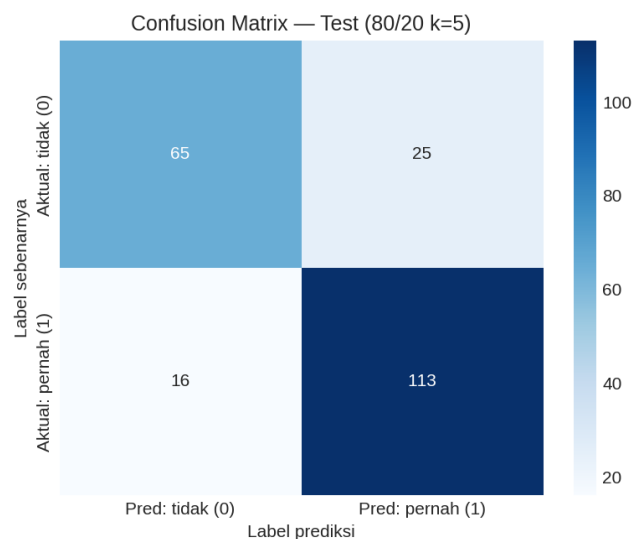
Tabel 4.7 Hasil Evaluasi Model Skenario 3

Parameter	K=5	K=10
Accuracy	0,8128	0,8128
Precision	0,8188	0,8188
Recall	0,8760	0,8760
F1-Score	0,8464	0,8464
ROC-AUC	0,8498	0,8498

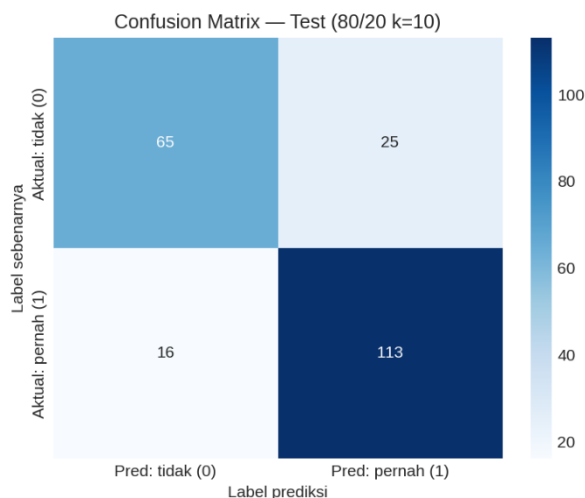
Berdasarkan Tabel 4.X, diperoleh hasil bahwa model Random Forest menunjukkan performa klasifikasi yang baik pada kedua variasi *K-Fold* yang digunakan. Pengujian menggunakan K=5 dan K=10 menghasilkan nilai evaluasi yang sama, yaitu accuracy sebesar 0,8128, precision sebesar 0,8188, recall sebesar 0,8760, F1-score sebesar 0,8464, dan ROC-AUC sebesar 0,8498. Nilai accuracy sebesar 81,28% menunjukkan bahwa model mampu melakukan klasifikasi dengan benar terhadap sebagian besar data uji. Nilai precision sebesar 81,88% menunjukkan bahwa dari seluruh data yang diprediksi termasuk dalam kelas *pernah melakukan treatment*, sebanyak 81,88% merupakan prediksi yang sesuai dengan kondisi sebenarnya.

Selain itu, nilai recall sebesar 87,60% menunjukkan bahwa model memiliki kemampuan yang baik dalam mengenali data yang benar-benar termasuk

dalam kelas *pernah melakukan treatment*. Keseimbangan antara nilai *precision* dan *recall* menghasilkan nilai F1-score sebesar 84,64%, yang menunjukkan bahwa model memiliki kemampuan klasifikasi yang cukup baik terhadap kedua kelas. Nilai ROC-AUC sebesar 0,8498 juga menunjukkan bahwa model memiliki kemampuan diskriminasi yang baik dalam membedakan kelas *pernah melakukan treatment* dan *tidak pernah melakukan treatment*. Berdasarkan hasil tersebut, perubahan jumlah *fold* dari K=5 menjadi K=10 pada proses Stratified K-Fold Cross Validation tidak memberikan peningkatan maupun penurunan terhadap performa model. Hal ini menunjukkan bahwa model Random Forest memiliki tingkat kestabilan yang baik pada kedua konfigurasi validasi silang tersebut.



Gambar 4.11 Confusion Matrix(80:20 K 5)

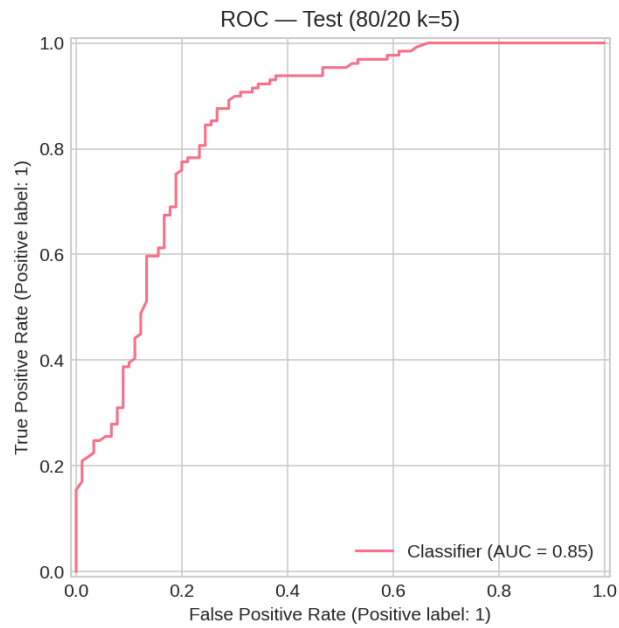


Gambar 4.12 Confusion Matrix(80:20 K 10)

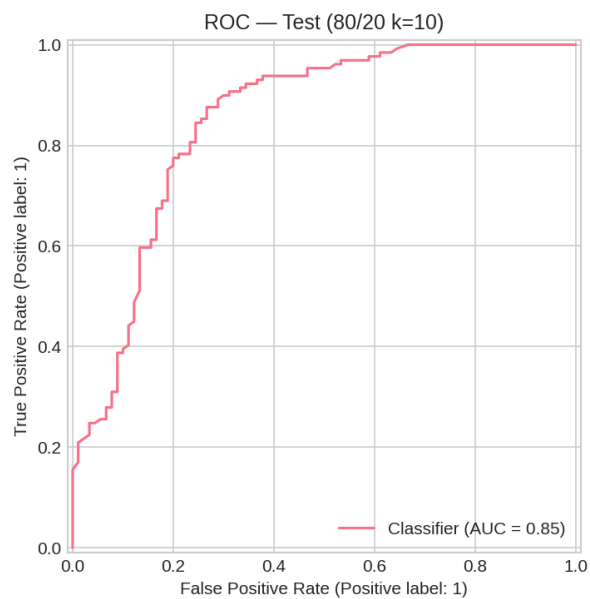
Berdasarkan hasil *confusion matrix* pada kedua pengujian, diperoleh nilai True Negative (TN) sebanyak 65 data, False Positive (FP) sebanyak 25 data, False Negative (FN) sebanyak 16 data, dan True Positive (TP) sebanyak 113 data. Hasil tersebut menunjukkan bahwa model memiliki kemampuan klasifikasi yang cukup baik karena jumlah prediksi yang benar lebih besar dibandingkan jumlah kesalahan klasifikasi. Sebanyak 65 data pada kelas tidak pernah melakukan *treatment* berhasil diklasifikasikan dengan benar, sedangkan 113 data pada kelas pernah melakukan *treatment* berhasil dikenali dengan tepat oleh model.

Meskipun demikian, masih terdapat kesalahan klasifikasi berupa False Positive sebanyak 25 data, yaitu data yang sebenarnya termasuk dalam kelas tidak pernah melakukan *treatment*, tetapi diprediksi sebagai pernah melakukan *treatment*. Selain itu, terdapat False Negative sebanyak 16 data, yaitu data yang sebenarnya pernah melakukan *treatment*, tetapi diprediksi sebagai tidak pernah melakukan *treatment*. Jumlah nilai True Positive yang lebih besar dibandingkan False Negative menunjukkan bahwa model memiliki kemampuan yang baik

dalam mendeteksi individu yang pernah melakukan *treatment*. Namun, keberadaan nilai *False Positive* dan *False Negative* menunjukkan bahwa masih terdapat beberapa karakteristik yang memiliki kemiripan antar kelas sehingga menyebabkan model melakukan kesalahan prediksi pada sejumlah data.



Gambar 4.13 Kurva ROC (80:20 K 5)



Gambar 4.14 Kurva ROC (80:20 K 10)

Berdasarkan hasil kurva ROC, diperoleh nilai Area Under Curve (AUC) sebesar 0,8498 atau dibulatkan menjadi 0,85 pada kedua pengujian. Kurva ROC yang berada jauh di atas garis diagonal menunjukkan bahwa model Random Forest memiliki kemampuan yang baik dalam membedakan kelas positif (*pernah melakukan treatment*) dan kelas negatif (*tidak pernah melakukan treatment*). Nilai AUC yang mendekati angka 1 mengindikasikan bahwa kemampuan diskriminasi model semakin baik, sehingga nilai AUC sebesar 0,8498 menunjukkan bahwa model mampu memberikan performa klasifikasi yang baik pada skenario pembagian data 80% data latih dan 20% data uji. Selain itu, kesamaan nilai AUC serta seluruh metrik evaluasi pada pengujian K=5 dan K=10 menunjukkan bahwa peningkatan jumlah *fold* pada proses Stratified K-Fold Cross Validation tidak memberikan pengaruh terhadap kemampuan diskriminasi maupun performa klasifikasi model pada skenario pengujian ini.

4.7 Pembahasan

Pengujian model pada penelitian ini dilakukan menggunakan data yang telah melalui beberapa tahapan *preprocessing*, meliputi pembersihan data, penanganan nilai yang tidak sesuai, normalisasi data menggunakan metode Min-Max Scaling, serta proses *feature selection* untuk memperoleh fitur yang paling relevan dalam proses klasifikasi kondisi kesehatan mental. Setelah melalui tahapan tersebut, model Random Forest kemudian dilatih dan diuji menggunakan beberapa skenario pembagian data, yaitu 60% data latih dan 40% data uji, 70% data latih dan 30% data uji, serta 80% data latih dan 20% data uji, dengan dua variasi Stratified K-Fold Cross Validation, yaitu K=5 dan K=10.

4.7.1 Analisis Performa Model

Berdasarkan hasil pengujian pada seluruh skenario, model Random Forest mampu menghasilkan performa klasifikasi yang cukup baik dalam mengklasifikasikan kondisi kesehatan mental berdasarkan riwayat *treatment*. Hal tersebut dapat dilihat dari nilai accuracy, precision, recall, F1-score, dan ROC-AUC yang berada pada kategori baik pada setiap skenario pengujian. Pada skenario pertama dengan pembagian data 60:40, model menghasilkan performa terbaik dibandingkan skenario lainnya dengan nilai accuracy sebesar 84,93%, precision sebesar 85,87%, recall sebesar 89,19%, F1-score sebesar 87,50%, serta ROC-AUC sebesar 0,8695. Hasil tersebut menunjukkan bahwa model mampu mengenali sebagian besar data yang termasuk dalam kelas *pernah melakukan treatment* serta memiliki kemampuan yang baik dalam membedakan kedua kelas.

Pada skenario kedua dengan pembagian data 70:30, model memperoleh nilai accuracy sebesar 83,28%, precision sebesar 84,24%, recall sebesar 88,14%, F1-score sebesar 86,15%, dan ROC-AUC sebesar 0,8696. Meskipun mengalami sedikit penurunan pada beberapa metrik dibandingkan skenario pertama, hasil tersebut masih menunjukkan bahwa model memiliki kemampuan klasifikasi yang baik dan stabil. Sementara itu, pada skenario ketiga dengan pembagian data 80:20, model menghasilkan accuracy sebesar 81,28%, precision sebesar 81,88%, recall sebesar 87,60%, F1-score sebesar 84,64%, dan ROC-AUC sebesar 0,8498. Penurunan nilai evaluasi dibandingkan skenario sebelumnya menunjukkan bahwa peningkatan proporsi data latih tidak selalu menghasilkan performa model yang lebih tinggi.

Hal ini dapat disebabkan oleh karakteristik distribusi data uji dan variasi pola data yang digunakan dalam proses pengujian. Secara keseluruhan, hasil pengujian menunjukkan bahwa pembagian data 60:40 memberikan performa terbaik pada penelitian ini. Selain itu, penggunaan jumlah *fold* $K=5$ dan $K=10$ tidak menunjukkan perbedaan performa yang signifikan karena seluruh metrik evaluasi pada kedua konfigurasi menghasilkan nilai yang sama pada masing-masing skenario.

4.7.2 Analisis *Confusion Matrix*

Analisis *confusion matrix* dilakukan untuk mengetahui kemampuan model dalam mengklasifikasikan setiap kelas secara lebih rinci, yaitu kelas tidak pernah melakukan treatment dan pernah melakukan treatment. Pada skenario pembagian data 60:40, model menghasilkan nilai True Negative (TN) sebanyak 141 data, False Positive (FP) sebanyak 38 data, False Negative (FN) sebanyak 28 data, dan True Positive (TP) sebanyak 231 data. Hasil tersebut menunjukkan bahwa model lebih banyak menghasilkan prediksi yang benar dibandingkan kesalahan klasifikasi, terutama dalam mengenali data yang pernah melakukan *treatment*.

Pada skenario pembagian data 70:30, diperoleh nilai TN sebanyak 103 data, FP sebanyak 32 data, FN sebanyak 23 data, dan TP sebanyak 171 data. Jumlah *True Positive* yang lebih tinggi dibandingkan *False Negative* menunjukkan bahwa model memiliki kemampuan yang baik dalam mendeteksi data yang termasuk ke dalam kelas pernah melakukan *treatment*. Pada skenario pembagian data 80:20, diperoleh nilai TN sebanyak 65 data, FP sebanyak 25 data, FN sebanyak 16 data, dan TP sebanyak 113 data.

Meskipun jumlah data yang digunakan dalam pengujian lebih sedikit, model tetap mampu mempertahankan kemampuan klasifikasi dengan jumlah prediksi benar yang lebih besar dibandingkan jumlah kesalahan prediksi. Secara keseluruhan, hasil *confusion matrix* menunjukkan bahwa model Random Forest memiliki kemampuan yang baik dalam mengenali kedua kelas. Namun, masih terdapat kesalahan klasifikasi berupa *False Positive* dan *False Negative* yang menunjukkan bahwa beberapa karakteristik data antara kelas pernah melakukan *treatment* dan tidak pernah melakukan *treatment* masih memiliki kemiripan sehingga menyebabkan model mengalami kesalahan dalam proses prediksi.

4.7.3 Analisis ROC-AUC

Kurva *Receiver Operating Characteristic (ROC)* dan nilai *Area Under Curve (AUC)* digunakan untuk mengukur kemampuan model dalam membedakan antara kelas positif dan negatif secara keseluruhan. Berdasarkan hasil pengujian pada seluruh skenario, nilai ROC-AUC yang diperoleh berada pada rentang 0,8498 hingga 0,8696, yang menunjukkan bahwa model Random Forest memiliki kemampuan diskriminasi yang baik dalam membedakan kelas *pernah melakukan treatment* dan *tidak pernah melakukan treatment*. Kurva ROC pada seluruh skenario juga berada jauh di atas garis diagonal, yang menandakan bahwa model mampu mencapai nilai *True Positive Rate* yang tinggi dengan *False Positive Rate* yang relatif rendah.

Skenario pembagian data 70:30 memperoleh nilai ROC-AUC tertinggi sebesar 0,8696, diikuti oleh skenario 60:40 sebesar 0,8695, sedangkan skenario 80:20 memperoleh nilai ROC-AUC sebesar 0,8498. Meskipun terdapat perbedaan

nilai antar skenario, selisih tersebut tidak terlalu besar sehingga menunjukkan bahwa model Random Forest memiliki performa yang relatif konsisten pada berbagai variasi pembagian data. Selain itu, penggunaan *Stratified K-Fold Cross Validation* dengan jumlah *fold* $K=5$ dan $K=10$ tidak memberikan perbedaan nilai ROC-AUC pada setiap skenario pengujian. Hal ini menunjukkan bahwa model memiliki tingkat kestabilan yang baik dan tidak terlalu dipengaruhi oleh perubahan jumlah *fold* pada proses validasi.

4.7.4 Analisis Pengaruh Pembagian Data dan Jumlah *Fold*

Jika dibandingkan berdasarkan variasi pembagian data, skenario 60:40 memberikan hasil performa terbaik berdasarkan nilai accuracy, precision, recall, dan F1-score. Sementara itu, nilai ROC-AUC tertinggi diperoleh pada skenario 70:30 dengan selisih yang sangat kecil dibandingkan skenario 60:40. Hal ini menunjukkan bahwa kedua skenario tersebut mampu menghasilkan performa model yang hampir setara. Di sisi lain, pembagian data 80:20 menghasilkan penurunan pada sebagian besar metrik evaluasi. Hal ini menunjukkan bahwa peningkatan proporsi data latih tidak selalu memberikan peningkatan performa model karena hasil prediksi juga dipengaruhi oleh distribusi data, karakteristik fitur, dan variasi data yang digunakan dalam proses pengujian.

Secara keseluruhan, hasil pengujian menunjukkan bahwa model Random Forest mampu memberikan performa klasifikasi yang baik pada berbagai skenario pembagian data, yaitu 60:40, 70:30, dan 80:20. Penerapan tahapan *preprocessing*, normalisasi menggunakan metode *Min-Max Scaling*, serta *feature selection* membantu model dalam mengidentifikasi pola pada data sehingga menghasilkan

nilai accuracy, precision, recall, F1-score, dan ROC-AUC yang cukup baik. Selain itu, variasi jumlah *fold* pada metode *Stratified K-Fold Cross Validation* antara K=5 dan K=10 tidak menunjukkan perbedaan performa yang signifikan, yang menunjukkan bahwa model memiliki tingkat kestabilan yang baik pada kedua konfigurasi pengujian.

4.8 Perbandingan Kinerja Model

Hasil pengujian menunjukkan bahwa perbedaan skenario pembagian data latih dan data uji memberikan variasi terhadap performa model *Random Forest*. Pengujian dilakukan pada tiga skenario pembagian data, yaitu 60% data latih dan 40% data uji, 70% data latih dan 30% data uji, serta 80% data latih dan 20% data uji menggunakan metode *Stratified K-Fold Cross Validation* dengan variasi K=5 dan K=10. Berdasarkan hasil pengujian pada seluruh skenario, diperoleh bahwa perubahan jumlah *fold* dari K=5 menjadi K=10 tidak memberikan perbedaan terhadap hasil evaluasi model. Oleh karena itu, perbandingan kinerja model dilakukan berdasarkan hasil evaluasi dari masing-masing skenario pembagian data.

Tabel 4.8 Perbandingan Kinerja Model

Skenario	Pembagian Data	K-Fold	Accuracy	Precision	Recall	F1-Score	ROC-AUC
Skenario 1	60% Data Latih : 40% Data Uji	K=5	0,8493	0,8587	0,8919	0,8750	0,8695
Skenario 1	60% Data Latih : 40% Data Uji	K=10	0,8493	0,8587	0,8919	0,8750	0,8695
Skenario 2	70% Data Latih : 30% Data Uji	K=5	0,8328	0,8424	0,8814	0,8615	0,8696
Skenario 2	70% Data Latih : 30% Data Uji	K=10	0,8328	0,8424	0,8814	0,8615	0,8696
Skenario	80% Data	K=5	0,8128	0,8188	0,8760	0,8464	0,8498

Skenario	Pembagian Data	K-Fold	Accuracy	Precision	Recall	F1-Score	ROC-AUC
3	Latih : 20% Data Uji						
Skenario 3	80% Data Latih : 20% Data Uji	K=10	0,8128	0,8188	0,8760	0,8464	0,8498

Berdasarkan tabel di atas, dapat dilihat bahwa ketiga skenario pengujian menghasilkan performa klasifikasi yang cukup baik dengan nilai evaluasi yang relatif tinggi. Namun, terdapat perbedaan performa pada setiap pembagian data yang digunakan.

Pada skenario pertama dengan pembagian data 60:40, model Random Forest menghasilkan nilai accuracy, precision, recall, dan F1-score tertinggi, yaitu masing-masing sebesar 0,8493, 0,8587, 0,8919, dan 0,8750. Hasil tersebut menunjukkan bahwa pembagian data 60:40 memberikan kemampuan klasifikasi yang lebih baik dalam mengenali pola data pada penelitian ini.

Pada skenario kedua dengan pembagian data 70:30, model memperoleh nilai ROC-AUC tertinggi sebesar 0,8696, yang menunjukkan kemampuan diskriminasi model yang sedikit lebih baik dibandingkan skenario lainnya. Meskipun demikian, perbedaan nilai ROC-AUC antara skenario 60:40 dan 70:30 sangat kecil, yaitu hanya sebesar 0,0001, sehingga kedua skenario dapat dikatakan memiliki kemampuan klasifikasi yang hampir setara.

Sementara itu, pada skenario ketiga dengan pembagian data 80:20, seluruh metrik evaluasi mengalami penurunan dibandingkan dua skenario sebelumnya. Model menghasilkan nilai accuracy sebesar 0,8128, precision sebesar 0,8188, recall sebesar 0,8760, F1-score sebesar 0,8464, dan ROC-AUC sebesar 0,8498.

Hal ini menunjukkan bahwa peningkatan proporsi data latih menjadi 80% tidak selalu menghasilkan performa yang lebih baik, karena performa model juga dipengaruhi oleh karakteristik distribusi data dan variasi data yang digunakan pada proses pengujian.

Secara keseluruhan, hasil pengujian menunjukkan bahwa pembagian data 60% data latih dan 40% data uji memberikan performa terbaik berdasarkan sebagian besar metrik evaluasi, yaitu accuracy, precision, recall, dan F1-score. Sementara itu, pembagian data 70% data latih dan 30% data uji memberikan nilai ROC-AUC tertinggi, meskipun selisihnya sangat kecil dibandingkan skenario 60:40. Selain itu, penggunaan jumlah *fold* $K=5$ dan $K=10$ pada metode Stratified K-Fold Cross Validation tidak memberikan pengaruh yang signifikan terhadap performa model karena menghasilkan nilai evaluasi yang sama pada setiap skenario pengujian. Dengan demikian, model memiliki tingkat kestabilan yang baik pada kedua konfigurasi *K-Fold* yang digunakan.

BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berdasarkan hasil pengujian yang telah dilakukan, model *Random Forest* mampu melakukan klasifikasi kondisi kesehatan mental berdasarkan riwayat *treatment* dengan performa yang baik. Proses klasifikasi dilakukan menggunakan data yang telah melalui tahapan *preprocessing*, normalisasi menggunakan metode *Min-Max Scaling*, dan proses *feature selection* untuk memperoleh fitur yang paling relevan dalam membangun model klasifikasi.

Berdasarkan tiga skenario pembagian data yang telah dilakukan, yaitu 60% data latih dan 40% data uji, 70% data latih dan 30% data uji, serta 80% data latih dan 20% data uji, diperoleh bahwa seluruh skenario menghasilkan performa klasifikasi yang baik. Skenario pertama dengan pembagian data 60:40 memberikan hasil terbaik berdasarkan nilai *accuracy*, *precision*, *recall*, dan *F1-score*, yaitu masing-masing sebesar 0,8493, 0,8587, 0,8919, dan 0,8750. Sementara itu, nilai ROC-AUC tertinggi diperoleh pada skenario kedua (70:30) sebesar 0,8696, meskipun selisihnya sangat kecil dibandingkan skenario pertama dengan nilai ROC-AUC sebesar 0,8695.

Hasil pengujian menggunakan metode *Stratified K-Fold Cross Validation* dengan variasi K=5 dan K=10 menunjukkan bahwa perubahan jumlah *fold* tidak memberikan perbedaan yang signifikan terhadap performa model. Hal tersebut ditunjukkan oleh nilai *accuracy*, *precision*, *recall*, *F1-score*, dan *ROC-AUC* yang sama pada kedua variasi K-Fold untuk setiap skenario pengujian, sehingga dapat

disimpulkan bahwa model Random Forest memiliki tingkat kestabilan yang baik pada kedua konfigurasi validasi tersebut.

Secara keseluruhan, model *Random Forest* yang dibangun mampu memberikan kemampuan klasifikasi yang baik dalam membedakan kondisi kesehatan mental berdasarkan riwayat *treatment*. Hal ini dibuktikan dengan nilai ROC-AUC yang berada pada rentang 0,8498 hingga 0,8696, yang menunjukkan bahwa model memiliki kemampuan diskriminasi yang baik dalam membedakan kelas *pernah melakukan treatment* dan *tidak pernah melakukan treatment*.

5.2 Saran

Berdasarkan hasil penelitian yang telah dilakukan, terdapat beberapa saran yang dapat dipertimbangkan untuk pengembangan penelitian selanjutnya. Pertama, penelitian selanjutnya dapat melakukan pengembangan pada tahap *feature engineering* atau mencoba metode *feature selection* lainnya untuk memperoleh kombinasi fitur yang lebih optimal sehingga dapat meningkatkan kemampuan model dalam mengenali pola kondisi kesehatan mental.

Kedua, penelitian selanjutnya dapat melakukan perbandingan dengan algoritma klasifikasi lainnya, seperti *Support Vector Machine (SVM)*, *Decision Tree*, *K-Nearest Neighbor (KNN)*, *Logistic Regression*, maupun algoritma berbasis *boosting*, sehingga dapat diketahui metode yang memiliki performa terbaik dalam melakukan klasifikasi kondisi kesehatan mental.

Ketiga, penelitian selanjutnya dapat mencoba variasi teknik optimasi parameter (*hyperparameter tuning*) yang lebih luas atau menggunakan metode

optimasi seperti Grid Search maupun Random Search untuk menemukan konfigurasi parameter Random Forest yang paling optimal.

Keempat, penelitian selanjutnya dapat melakukan pengujian dengan variasi pembagian data dan teknik validasi lainnya, serta menggunakan dataset yang lebih beragam agar dapat mengetahui tingkat generalisasi model terhadap data yang berbeda.

DAFTAR PUSTAKA

- Breiman, L. (2001). *Random Forests*. *Machine Learning*, 45(1), 5–32.
<https://doi.org/https://doi.org/10.1023/A:1010933404324>
- Choirunnisa, N. (2025). *Random Forest* for large-scale mental disorder identification. *Universitas Islam Negeri Maulana Malik Ibrahim Malang*.
- Fahmi, R., Yusuf, A., & Adhi, D. (2025). Peningkatan prediksi kesehatan mental pekerja IT *remote* menggunakan *Random Forest* dan feature engineering. *Journal of Intelligent Systems Research*, 12(1), 33–42.
- Faisti, R., Nurdiansyah, N., & Baihaqi, W. M. (2025). Mental health classification using *Random Forest* for predictive decision support. *Malcom: Indonesian Journal of Machine Learning and Computer Science*, 5(1), 1–9.
- Febrian, R., Sari, D., & Nugroho, A. (2025). Employee mental health and productivity in *remote* working systems. *Journal of Organizational Psychology*, 14(2), 102–115.
- Friedman, J., Hastie, T., & Tibshirani, R. (2009). The elements of statistical learning: Data mining, inference, and prediction (2nd ed.). *Springer*.
- Gamage, A., Retnaswamy, R., & Kumar, S. (2025). Predicting work stress among *remote* employees using *Random Forest*. *Asian Journal of Computer Science*, 11(3), 188–200.
- Han, J., Kamber, M., & Pei, J. (2012). Data mining: Concepts and techniques (3rd ed.). *Morgan Kaufmann*.
- Hendra, H., Deris, M. M., & Windiarti, I. S. (2025). Mental health classification using machine learning with PCA and logistic regression approaches for decision making. *Engineering Proceedings*, 84(1), 47.
<https://doi.org/https://doi.org/10.3390/engproc2025084047>
- Kannan, P., Devi, R., & Suresh, K. (2024). Advances in machine learning for mental disorder diagnosis and management. *Journal of Computational Neuroscience*, 18(2), 145–159.
- Khan, I., Wu, C., & Ropponen, A. (2025). Psychological impacts of *remote* working: A meta-analysis. *International Journal of Mental Health Studies*, 9(1), 55–70.
- Kunhayati, U. (2025). Klasifikasi kesehatan mental pada imbalanced data menggunakan *Random Forest* dan SMOTE. *Universitas Islam Negeri Maulana Malik Ibrahim Malang*.
- Kusumarini, T., Rahmawati, D., & Santoso, B. (2021). Decision tree-based classification models for mental health data. *Journal of Data Science*, 19(4), 567–578.

- Maramis, W. F. (2022). Ilmu kedokteran jiwa. *Airlangga University Press*.
- Marsuhandi, R., Nurdin, D., & Hartono, R. (2020). Implementation of *Random Forest* algorithm for mental health prediction. *Indonesian Journal of Computing and Data Science*, 3(2), 55–62.
- Ndikumana, L., Hakizimana, E., & Uwimana, J. (2025). Mental health prediction among adolescents using *Random Forest*. *African Journal of Information Systems*, 14(1), 22–34.
- Nurdiansyah, N., Febriyan, F. S., Amanta, Z. G. D., Saputra, D. A., & Baihaqi, W. M. (2025). Analisis kesehatan mental untuk mencegah gangguan mental pada mahasiswa menggunakan algoritma K-Nearest Neighbor (K-NN) dan *Random Forest*. *Malcom: Indonesian Journal of Machine Learning and Computer Science*, 5(1), 1–9. <https://doi.org/https://doi.org/10.57152/malcom.v5i1.1537>
- Oktaviani, S., Rahman, F., & Wibowo, A. (2024). *Random Forest* model performance in predicting mental health stress factors. *Journal of Artificial Intelligence Research*, 17(4), 312–328.
- OSMI. (2016). Mental Health in Tech Survey 2016 dataset [Data set]. *Kaggle*. <https://www.kaggle.com/datasets/osmi/mental-health-in-tech-2016>
- Priyono, A., Hartanto, M., & Setiawan, D. (2024). Klasifikasi diagnosis penyakit mental menggunakan metode *Random Forest* dengan data survei kesehatan mental. *Journal of Applied Data Science*, 13(1), 44–57.
- Rahman, M. F., Wijaya, R., & Pratama, D. (2023). Exploratory data analysis and feature selection in mental health prediction using *Random Forest*. *Procedia Computer Science*, 207, 433–440.
- Retnaswamy, R., Gamage, A., & Kumar, S. (2025). Comparative analysis of *Random Forest*, Decision Tree, and SVM for mental health treatment prediction. *Asian Journal of Computer Science*, 11(3), 211–225.
- Ropponen, A. (2025). Remote working and psychological well-being: Evidence from global employment trends. *Occupational Psychology Review*, 12(1), 34–47.
- Rosenberg, E., Davis, T., & Mitchell, R. (2021). Evaluation of *Random Forest* and Decision Tree in predicting mental well-being. *Computational Behavioral Science*, 10(2), 98–110.
- Sebayang, H., Wijaya, R., & Pratama, D. (2023). Performance comparison of *Random Forest* and Logistic Regression in mental health prediction. *Procedia Computer Science*, 207, 433–440.
- Sun, Y. (2024). Evaluating machine learning models for mental health prediction using *Random Forest*. *International Journal of Intelligent Systems*, 39(2), 165–180.

WHO. (2022a). *Mental health: Strengthening our response*.

WHO. (2022b). *World Mental Health Report: Transforming Mental Health for All (1st ed.)*.

Wu, X., Zhang, H., & Li, M. (2024). The psychological effects of long-term remote working: An empirical study. *Journal of Behavioral Science*, 15(2), 99–112.