

BAB I

PENDAHULUAN

1.1 Latar Belakang

Diabetes adalah kondisi di mana tubuh mengalami gangguan dalam memproduksi insulin, yang dapat menyebabkan kerusakan pada organ lain dalam tubuh. Diabetes merupakan salah satu penyakit yang umum terjadi dan mempengaruhi banyak individu di seluruh dunia (Ahmed dkk., 2021). Berdasarkan data Federasi Diabetes Internasional, kematian akibat diabetes mencapai 6,7 juta jiwa pada tahun 2021. Sebanyak 537 juta orang berusia 20-79 tahun hidup dengan diabetes, dan jumlah ini diperkirakan meningkat menjadi 643 juta pada tahun 2030 serta 783 juta pada tahun 2045 (Thotad dkk., 2023). Menurut laporan Kementerian Kesehatan Republik Indonesia tahun 2020, Indonesia termasuk dalam sepuluh negara dengan tingkat diabetes tertinggi pada tahun 2019.

Deteksi dini diabetes sangat penting untuk mencegah kondisi kesehatan yang semakin memburuk. Keahlian tenaga medis sangat diperlukan dalam menentukan apakah seorang pasien menderita diabetes atau tidak. Namun, banyak individu terlambat menyadari kondisi tersebut, sehingga kondisi ini sudah berada pada tahap yang serius saat didiagnosis (Febrian dkk., 2022). Untuk mendukung upaya deteksi dini ini, diperlukan pengembangan model prediksi yang akurat dan andal. Salah satu pendekatan yang menjanjikan adalah pemanfaatan teknologi berbasis *machine learning* untuk menganalisis data kesehatan yang tersedia dan memprediksi risiko diabetes secara efektif. Beberapa algoritma *machine learning* konvensional, seperti *logistic regression*, *k-nearest neighbors*, *decision trees*, *random forest*, dan *support vector machine*, sering digunakan untuk membangun model prediksi (S dkk., 2024). Namun, algoritma tersebut memiliki keterbatasan dalam situasi tertentu. Untuk mengatasi keterbatasan ini, pendekatan *ensemble learning* menjadi solusi yang efektif karena menggabungkan prediksi dari beberapa model untuk meningkatkan akurasi dan stabilitas hasil (Feng dkk., 2023). Penelitian yang dilakukan untuk membuktikan bahwa metode *ensemble learning* dapat menjadi solusi yang efektif untuk kasus prediksi. Hasil pengujian menunjukkan

bahwa *ensemble learning* mampu meningkatkan akurasi dibandingkan metode individu lainnya (Ngo dkk., 2022).

Salah satu metode *ensemble learning* yang sering digunakan adalah *bootstrap aggregating (bagging)*. Penerapan metode *bagging* pada diagnosis diabetes dengan dataset UCI Diabetes menghasilkan performa yang sangat baik berdasarkan metrik seperti akurasi, presisi, *recall*, dan F1-score (Yadav dan Pal, 2021). Metode berbasis *bagging* yang populer adalah *random forest*. *Random forest* sering digunakan dalam kasus prediksi diabetes karena kemampuannya menangani data yang kompleks, dengan banyak variabel dan interaksi antar fitur, tanpa memerlukan asumsi linearitas. Model ini tahan terhadap *overfitting* karena menggabungkan prediksi dari banyak *decision tree* independen, sehingga lebih stabil dibandingkan model individual seperti *decision tree* (Tuppad dan Devi Patil, 2024). Dalam penelitian sebelumnya, penggunaan algoritma *random forest* telah banyak digunakan dalam klasifikasi diabetes. Penelitian diabetes di Distrik Haizhu, Guangzhou, menunjukkan bahwa model *random forest* memiliki performa yang sangat baik dalam memprediksi diabetes dengan akurasi mencapai 91,24%, meskipun demikian, penelitian ini tidak mengatasi masalah ketidakseimbangan kelas data (Jiang dkk., 2023).

Dalam memprediksi menggunakan klasifikasi untuk kasus diabetes, umumnya jumlah data pasien yang menderita diabetes jauh lebih sedikit dibandingkan dengan pasien non-diabetes. Ketidakseimbangan kelas ini menjadi tantangan dalam pembuatan model prediksi karena algoritma *machine learning* cenderung lebih memprioritaskan kelas dengan jumlah data yang lebih banyak, sehingga berpotensi mengabaikan kelas dengan data yang lebih sedikit dan mempengaruhi akurasi hasil prediksi. Permasalahan ini dapat diatasi dengan menyeimbangkan jumlah data pada masing-masing kelas, baik dengan meningkatkan jumlah data pada kelas minoritas maupun mengurangi jumlah data pada kelas mayoritas, menggunakan metode tertentu untuk memastikan model dapat mempelajari pola secara seimbang dan menghasilkan prediksi yang lebih akurat (Ramadhan, 2021). Untuk mengatasi masalah ketidakseimbangan data, teknik *Adaptive Synthetic Sampling* (ADASYN) dapat digunakan. Teknik ini sangat

relevan dalam *ensemble machine learning* seperti *random forest*. Dengan meningkatkan jumlah sampel pada kelas minoritas secara adaptif berdasarkan distribusi data, ADASYN bertujuan untuk mengurangi bias yang muncul akibat data yang tidak seimbang. Pendekatan ini membantu menghasilkan model prediktif yang lebih akurat dan andal dalam mempelajari pola pada kelas minoritas, sehingga memungkinkan prediksi yang lebih seimbang dan representatif pada kedua kelas, baik pasien diabetes maupun non-diabetes (Muflikhah dkk., 2024). Penelitian menggunakan kombinasi *random forest* dan ADASYN pada dataset CICIDS 2017 menunjukkan kinerja yang lebih baik dibandingkan pendekatan lainnya. Teknik ini mampu meningkatkan performa prediksi secara signifikan, dengan hasil metrik seperti presisi, *recall*, F1-score, dan AUC yang lebih baik (Chen dkk., 2021).

Selain meningkatkan akurasi prediksi, interpretabilitas hasil *machine learning* juga menjadi perhatian penting, terutama di bidang medis. Interpretabilitas memungkinkan tenaga medis memahami faktor-faktor yang mempengaruhi keputusan model. Salah satu metode yang mendukung hal ini adalah *SHapley Additive exPlanations* (SHAP). SHAP mampu menjelaskan kontribusi setiap fitur, seperti kadar glukosa darah, tekanan darah, dan BMI, terhadap hasil prediksi. Dengan transparansi ini, tenaga medis dapat menggunakan model secara lebih efektif dalam mendukung diagnosis dan perencanaan pengobatan (Fu dkk., 2024).

Penelitian ini menawarkan pendekatan integratif dengan mengkombinasikan algoritma *random forest*, teknik penyeimbang data ADASYN, dan metode interpretabilitas model SHAP. Pendekatan ini belum banyak diterapkan dalam studi sebelumnya, khususnya pada dataset PIMA Indians Diabetes. Keunggulan model tidak hanya terletak pada performa klasifikasi yang baik, tetapi juga pada kemampuannya memberikan interpretasi yang lebih transparan mengenai kontribusi masing-masing fitur terhadap hasil prediksi. Penelitian ini juga mengeksplorasi hasil klasifikasi berdasarkan subpopulasi pasien, seperti kelompok usia dan BMI, yang dapat memiliki karakteristik risiko diabetes yang berbeda. Hal ini menimbulkan kebutuhan akan pendekatan baru yang tidak hanya menghasilkan model yang akurat, tetapi juga dapat diinterpretasikan secara klinis.

Untuk membatasi ruang lingkup penelitian, studi ini hanya menggunakan dataset PIMA Indians yang tersedia secara publik, yang secara khusus merepresentasikan data klinis perempuan keturunan Indian berusia di atas 21 tahun. Fokus dari penelitian ini adalah mengintegrasikan metode klasifikasi dan interpretabilitas yang telah terbukti efektif secara individual, yaitu *random forest*, ADASYN, dan SHAP. Integrasi ketiganya dilakukan melalui proses pengujian model berbasis *pipeline machine learning*, yang bertujuan untuk menghasilkan model prediksi diabetes yang tidak hanya akurat, tetapi juga transparan dalam menjelaskan kontribusi masing-masing fitur terhadap hasil prediksi.

1.2 Tujuan Penelitian

Tujuan dari penelitian ini adalah membangun dan mengevaluasi model prediksi diabetes yang terintegrasi, dengan menggabungkan algoritma *random forest* sebagai metode klasifikasi, ADASYN untuk menangani ketidakseimbangan kelas, serta metode interpretabilitas SHAP untuk menjelaskan kontribusi masing-masing fitur terhadap hasil prediksi. Selain itu, penelitian ini juga mengevaluasi pengaruh penanganan *outlier* terhadap kinerja model, guna memperoleh pendekatan prapemrosesan data yang optimal. Model ini diharapkan tidak hanya memiliki performa prediksi yang baik, tetapi juga mampu memberikan interpretasi yang jelas dan dapat digunakan sebagai alat bantu dalam pengambilan keputusan klinis berbasis data. Model integratif tersebut diimplementasikan untuk data dari *National Institute of Diabetes and Digestive and Kidney Diseases* untuk mengevaluasi performa prediksinya.

1.3 Manfaat Penelitian

Penelitian ini diharapkan memberikan manfaat sebagai berikut:

1. Penelitian ini diharapkan dapat memberikan kontribusi terhadap pengembangan kajian di bidang *data mining* dan *machine learning*, khususnya dalam penerapan model klasifikasi yang terintegrasi dengan teknik penyeimbang data dan metode interpretabilitas. Integrasi antara *random forest*, ADASYN, dan SHAP dapat menjadi referensi bagi penelitian-penelitian selanjutnya yang ingin

membangun model prediktif yang tidak hanya akurat, tetapi juga dapat dijelaskan secara transparan.

2. Penelitian ini mengeksplorasi perbedaan performa dan interpretasi model pada subpopulasi pasien, sehingga memberikan wawasan yang lebih spesifik dan personal dalam konteks prediksi risiko diabetes. Pendekatan ini memungkinkan pengambilan keputusan klinis menjadi lebih kontekstual dan tepat sasaran.
3. Menyelaraskan hasil penelitian dengan literatur medis yang ada sebagai bentuk validasi tidak langsung terhadap model yang dikembangkan, guna memastikan hasil model tidak hanya signifikan secara statistik, tetapi juga bermakna secara klinis.



SEKOLAH PASCASARJANA