

**PREDIKSI RISIKO GERD BERDASARKAN POLA MAKAN  
MENGUNAKAN METODE *RANDOM FOREST***

**SKRIPSI**

**Oleh:**  
**NASYWA QONIATUL HAYAH**  
**NIM. 210605110112**



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

**PREDIKSI RISIKO GERD BERDASARKAN POLA MAKAN  
MENGUNAKAN 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 :  
**NASYWA QONIATUL HAYAH**  
**NIM. 210605110112**

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

**HALAMAN PERSETUJUAN**

**PREDIKSI RISIKO GERD BERDASARKAN POLA MAKAN  
MENGUNAKAN METODE *RANDOM FOREST***

**SKRIPSI**

**Oleh :**

**NASYWA QONIATUL HAYAH**

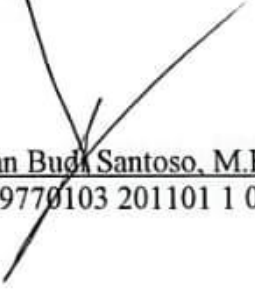
**NIM. 210605110112**

Telah Diperiksa dan Disetujui untuk Diuji:  
Tanggal: 19 Juni 2026

Pembimbing I,

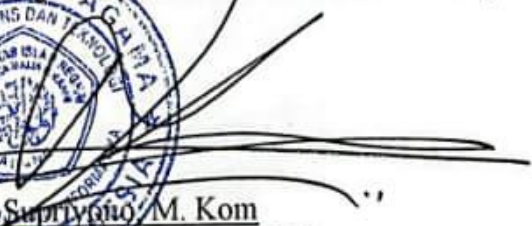
  
Dr. Cahyo Crysdian, M.Cs  
NIP. 19740424 200901 1 008

Pembimbing II,

  
Dr. Irwan Budi Santoso, M.Kom  
NIP. 19770103 201101 1 004

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

### PREDIKSI RISIKO GERD BERDASARKAN POLA MAKAN MENGUNAKAN METODE *RANDOM FOREST*

#### SKRIPSI

Oleh :  
**NASYWA QONIATUL HAYAH**  
NIM. 210605110112

Telah Dipertahankan di Depan Dewan Penguji Skripsi  
dan Dinyatakan Diterima Sebagai Salah Satu Persyaratan  
Untuk Memperoleh Gelar Sarjana Komputer ( S.Kom )  
Tanggal: 26 Juni 2026

#### Susunan Dewan Penguji

|                     |                                                                      |
|---------------------|----------------------------------------------------------------------|
| Ketua Penguji       | : <u>Dr. M. Amin Hariyadi, M.T</u><br>NIP. 19670118 200501 1 001     |
| Anggota Penguji I   | : <u>Okta Qomaruddin Aziz, M.Kom</u><br>NIP. 19911019 201903 1 013   |
| Anggota Penguji II  | : <u>Dr. Cahyo Crysdian, M.Cs</u><br>NIP. 19740424 200901 1 008      |
| Anggota Penguji III | : <u>Dr. Irwan Budi Santoso, M.Kom</u><br>NIP. 19770103 201101 1 004 |

(  )  
(  )  
(  )  
(  )

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 : Nasywa Qoniatul Hayah  
NIM : 210605110112  
Fakultas / Program Studi : Sains dan Teknologi / Teknik Informatika  
Judul Skripsi : Prediksi Risiko GERD Berdasarkan Pola Makan  
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,



Nasywa Qoniatul Hayah

NIM. 210605110112

## **MOTTO**

*"Allah tidak membebani seseorang melainkan sesuai dengan kesanggupannya."  
(QS. Al-Baqarah: 286)*

*"Barang siapa menempuh jalan untuk mencari ilmu, maka Allah akan mudahkan baginya jalan menuju surga."  
(HR. Muslim)*

*"Every action you take is a vote for the person you wish to become."  
(James Clear, Atomic Habits)*

## **HALAMAN PERSEMBAHAN**

Dengan mengucapkan rasa syukur kepada Allah SWT atas segala rahmat, nikmat, dan kemudahan yang telah diberikan, karya sederhana ini penulis persembahkan kepada:

Kedua orang tua tercinta, yang senantiasa memberikan doa, kasih sayang, serta dukungan tanpa henti dalam setiap langkah kehidupan penulis.

Kakak-kakak saya, yang telah memberikan dukungan, pengorbanan, dan semangat sehingga penulis dapat menyelesaikan pendidikan hingga tahap ini.

Seluruh teman semasa kuliah yang selalu menjadi tempat berbagi cerita, memberikan motivasi, serta menemani proses perjuangan selama masa perkuliahan dan penyusunan skripsi.

Seluruh dosen dan guru yang telah memberikan ilmu, arahan, serta inspirasi yang menjadi bekal berharga bagi penulis.

Sahabat dan teman-teman seperjuangan yang telah kebersamai, membantu, serta memberikan semangat dalam berbagai proses yang telah dilalui.

Dan terakhir, untuk diri sendiri yang telah berusaha bertahan, belajar, dan terus melangkah hingga mampu menyelesaikan salah satu amanah besar dalam kehidupan ini.

## KATA PENGANTAR

*Alhamdulillahillāhi rabbil ‘ālamīn*, puji syukur kehadiran Allah SWT atas segala rahmat, nikmat, taufik, serta hidayah-Nya sehingga penulis dapat menyelesaikan skripsi yang berjudul “*Prediksi Risiko GERD Berdasarkan Pola Makan Menggunakan Metode Random Forest*” dengan baik. Shalawat serta salam senantiasa tercurahkan kepada Nabi Muhammad SAW yang telah membawa umat manusia dari zaman kegelapan menuju zaman yang penuh dengan ilmu pengetahuan dan peradaban.

Skripsi ini disusun sebagai salah satu syarat untuk memperoleh gelar Sarjana Komputer pada Program Studi Teknik Informatika, Fakultas Sains dan Teknologi, Universitas Islam Negeri Maulana Malik Ibrahim Malang. Dalam proses penyusunan skripsi ini, penulis menyadari bahwa keberhasilan yang dicapai tidak terlepas dari bantuan, dukungan, arahan, serta doa dari berbagai pihak. Oleh karena itu, pada kesempatan ini penulis menyampaikan ucapan terima kasih kepada:

1. Prof. Dr. Hj. Ilfi Nur Diana, M.Si, selaku Rektor Universitas Islam Negeri Malang.
2. Dr. Agus Mulyono, M.Kes, selaku Dekan Fakultas Sains dan Teknologi Universitas Islam Negeri Malang.
3. Supriyono, M. Kom, selaku Ketua Program Studi Teknik Informatika Universitas Islam Negeri Maulana Malik Ibrahim Malang.

4. Dr. Cahyo Crysdiان, M.Cs. selaku dosen pembimbing I dan Dr. Irwan Budi Santoso, M.Kom. selaku dosen pembimbing II yang telah memberikan bantuan dan arahan kepada penulis, sehingga bisa menuntaskan skripsi ini.
5. Dr. M. Amin Hariyadi, M.T. selaku ketua penguji dan Okta Qomaruddin Aziz, M.Kom. selaku anggota penguji I yang telah menguji serta memberikan masukan sehingga penulis dapat menuntaskan skripsi dengan baik.
6. Nia Faricha S, Si., selaku admin Program Studi Teknik Informatika yang selalu sabar memberikan informasi, membantu, dan memberikan arahan selama perkuliahan dan proses penulisan skripsi ini.
7. Segenap dosen, laboran, dan jajaran staf Program Studi Teknik Informatika yang telah memberikan ilmu, pengetahuan, dan dukungan selama penulis menjalani studi hingga selesainya skripsi ini.
8. Kedua orang tua tercinta, yang senantiasa memberikan doa, dukungan, kasih sayang, serta pengorbanan yang tidak ternilai selama proses pendidikan penulis.
9. Saudara dan keluarga yang selalu menjadi sumber inspirasi, motivasi, dan dukungan selama perjalanan ini. Terima kasih atas semangat, nasihat, serta perhatian yang tak henti diberikan kepada penulis.
10. Teman dekat penulis, Aninda Rizky Hartanti yang senantiasa memberikan dukungan, doa, semangat, serta menjadi tempat berbagi cerita selama proses penyusunan skripsi ini.

11. Keluarga besar Kesatuan Aksi Mahasiswa Muslim Indonesia (KAMMI), baik para senior, teman seperjuangan, maupun adik tingkat, yang telah memberikan motivasi, pengalaman, kebersamaan, serta lingkungan yang mendukung penulis untuk terus bertumbuh selama masa perkuliahan.
12. Penghuni Kontrakan Kautsar, baik kakak tingkat, teman, maupun adik tingkat, yang telah menjadi keluarga kedua bagi penulis serta memberikan dukungan, kebersamaan, dan suasana yang penuh kekeluargaan selama menempuh studi.
13. Seluruh warga Teknik Informatika khususnya angkatan 2021 “Aster” yang telah menjadi rekan belajar, berbagi pengalaman, serta memberikan semangat selama menjalani proses perkuliahan hingga penyusunan skripsi ini.
14. Seluruh pihak yang tidak dapat disebutkan satu per satu yang telah membantu dalam penyusunan skripsi ini.

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       |
| البحث مستخلص.....                                         | xvii      |
| <b>BAB I PENDAHULUAN.....</b>                             | <b>1</b>  |
| 1.1 Latar Belakang .....                                  | 1         |
| 1.2 Pernyataan Masalah .....                              | 6         |
| 1.3 Batasan Masalah .....                                 | 6         |
| 1.4 Tujuan Penelitian .....                               | 6         |
| 1.5 Manfaat Penelitian .....                              | 6         |
| <b>BAB II STUDI PUSTAKA .....</b>                         | <b>8</b>  |
| <b>BAB III DESAIN DAN IMPLEMENTASI SISTEM .....</b>       | <b>17</b> |
| 3.1 Persiapan Data .....                                  | 17        |
| 3.2 Desain Sistem.....                                    | 22        |
| 3.2.1 <i>Input Data</i> .....                             | 24        |
| 3.2.2 Normalisasi Data .....                              | 24        |
| 3.2.3 Klasifikasi dengan <i>Random Forest</i> .....       | 25        |
| 3.2.4 Hasil Prediksi Risiko GERD .....                    | 29        |
| 3.3 Implementasi Sistem.....                              | 30        |
| 3.3.1 <i>Import Library</i> dan <i>Load Dataset</i> ..... | 30        |
| 3.3.2 Pemisahan Fitur dan Target.....                     | 31        |
| 3.3.3 Pembagian Data dan Normalisasi Data .....           | 31        |
| 3.3.4 Model <i>Random Forest</i> .....                    | 32        |
| <b>BAB IV UJI COBA DAN PEMBAHASAN .....</b>               | <b>34</b> |
| 4.1 Skenario Pengujian .....                              | 34        |
| 4.2 Hasil Pengujian .....                                 | 38        |
| 4.2.1 Hasil Pengujian Skenario A1 .....                   | 38        |
| 4.2.2 Hasil Pengujian Skenario A2 .....                   | 39        |
| 4.2.3 Hasil Pengujian Skenario A3 .....                   | 39        |
| 4.2.4 Hasil Pengujian Skenario A4 .....                   | 40        |
| 4.2.5 Hasil Pengujian Skenario A5 .....                   | 40        |
| 4.2.6 Hasil Pengujian Skenario A6 .....                   | 42        |
| 4.2.7 Hasil Pengujian Skenario A7 .....                   | 42        |
| 4.2.8 Hasil Pengujian Skenario A8 .....                   | 43        |

|                                                                                             |           |
|---------------------------------------------------------------------------------------------|-----------|
| 4.2.9 Hasil Pengujian Skenario A9 .....                                                     | 43        |
| 4.3 Analisis dan Pembahasan.....                                                            | 44        |
| 4.3.1 Analisis Performa Model pada Berbagai Skenario Pengujian .....                        | 44        |
| 4.3.2 Analisis Pengaruh Parameter <i>Random Forest</i> terhadap Performa Model.             | 47        |
| 4.3.3 Visualisasi <i>Decision Tree Random Forest</i> .....                                  | 48        |
| 4.3.4 Analisis <i>Feature Importance</i> .....                                              | 50        |
| 4.3.5 Pengujian Tambahan Menggunakan <i>Random UnderSampling</i> dari Model<br>Terbaik..... | 52        |
| <b>BAB V KESIMPULAN DAN SARAN .....</b>                                                     | <b>59</b> |
| 5.1 Kesimpulan .....                                                                        | 59        |
| 5.2 Saran .....                                                                             | 59        |
| <b>DAFTAR PUSTAKA</b>                                                                       |           |

## DAFTAR GAMBAR

|                                                                           |    |
|---------------------------------------------------------------------------|----|
| Gambar 3.1 Desain Sistem Prediksi Risiko GERD .....                       | 23 |
| Gambar 3.2 Desain Model <i>Random Forest</i> .....                        | 26 |
| Gambar 4.1 Diagram Alir Proses Pengujian .....                            | 35 |
| Gambar 4.2 Perbandingan Pengaruh Parameter Terhadap <i>F1-Score</i> ..... | 47 |
| Gambar 4.3 Contoh Decision Tree dari Model <i>Random Forest</i> .....     | 49 |
| Gambar 4.4 Diagram <i>Feature Importance Random Forest</i> .....          | 51 |
| Gambar 4.5 Confusion Matrix Skenario A3 dengan UnderSampling .....        | 54 |

## DAFTAR TABEL

|                                                                                       |    |
|---------------------------------------------------------------------------------------|----|
| Tabel 3.1 Rincian Fitur Dataset .....                                                 | 18 |
| Tabel 3.2 Dataset Sebelum Normalisasi .....                                           | 20 |
| Tabel 3.3 Dataset Setelah Normalisasi.....                                            | 20 |
| Tabel 3.4 Distribusi Data Sebelum dan Sesudah SMOTE.....                              | 21 |
| Tabel 4.1 Variasi Parameter Pengujian.....                                            | 36 |
| Tabel 4.2 Hasil Pengujian Skenario A1 .....                                           | 38 |
| Tabel 4.3 Hasil Pengujian Skenario A2 .....                                           | 39 |
| Tabel 4.4 Hasil Pengujian Skenario A3 .....                                           | 39 |
| Tabel 4.5 Hasil Pengujian Skenario A4 .....                                           | 40 |
| Tabel 4.6 Hasil Pengujian Skenario A5 .....                                           | 40 |
| Tabel 4.7 Hasil Pengujian Skenario A6 .....                                           | 42 |
| Tabel 4.8 Hasil Pengujian Skenario A7 .....                                           | 42 |
| Tabel 4.9 Hasil Pengujian Skenario A8 .....                                           | 43 |
| Tabel 4.10 Hasil Pengujian Skenario A9 .....                                          | 43 |
| Tabel 4.11 Rekapitulasi Performa Seluruh Skenario .....                               | 45 |
| Tabel 4.12 Perbandingan Performa Model Berdasarkan F1-Score.....                      | 45 |
| Tabel 4.13 Nilai Feature Importance Random Forest.....                                | 50 |
| Tabel 4.14 Distribusi Data Sebelum dan Sesudah Random UnderSampling .....             | 53 |
| Tabel 4.15 Perbandingan Performa Model Sebelum dan Sesudah Random UnderSampling ..... | 53 |

## ABSTRAK

Hayah, Nasywa. 2026. **Prediksi Risiko GERD Berdasarkan Pola Makan Menggunakan Metode *Random Forest***. Skripsi. Program Studi Teknik Informatika Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang. Pembimbing: (I) Dr. Cahyo Crys dian, M.Cs (II) Dr. Irwan Budi Santoso, M.Kom.

**Kata kunci:** GERD, *Random Forest*, Prediksi Risiko, Pola Makan, *Machine Learning*

*Gastroesophageal Reflux Disease* (GERD) merupakan gangguan pencernaan yang terjadi akibat naiknya asam lambung ke kerongkongan dan dapat dipengaruhi oleh pola makan. Keterlambatan dalam mengenali faktor risiko GERD dapat menyebabkan penurunan kualitas hidup dan komplikasi kesehatan. Penelitian ini bertujuan untuk membangun dan mengevaluasi model prediksi risiko GERD berdasarkan pola makan menggunakan algoritma *Random Forest*. *Dataset* yang digunakan merupakan *Acid Reflux Dataset for Classification* yang diperoleh dari Kaggle dengan 13 variabel independen yang berkaitan dengan pola makan serta satu variabel target berupa risiko GERD. Tahapan penelitian meliputi persiapan data, normalisasi menggunakan *Min-Max Scaling*, penanganan ketidakseimbangan data menggunakan *Synthetic Minority Oversampling Technique* (SMOTE), pelatihan model *Random Forest*, dan evaluasi model menggunakan *confusion matrix*, *accuracy*, *precision*, *recall*, dan *F1-Score*. Pengujian dilakukan pada sembilan skenario kombinasi parameter *n\_estimators* dan *max\_depth*. Hasil penelitian menunjukkan bahwa model terbaik diperoleh pada skenario A3 dengan parameter *n\_estimators* = 50 dan *max\_depth* = *None* yang menghasilkan *accuracy* sebesar 89,60%, *precision* sebesar 74,04%, *recall* sebesar 51,68%, dan *F1-Score* sebesar 60,87%. Pengujian tambahan menggunakan *Random UnderSampling* menghasilkan *accuracy* sebesar 75,83%, *precision* sebesar 79,84%, *recall* sebesar 69,12%, dan *F1-Score* sebesar 74,10%, yang menunjukkan bahwa model mampu memberikan performa yang baik pada *dataset* dengan distribusi kelas yang seimbang. Berdasarkan hasil tersebut, algoritma *Random Forest* mampu digunakan untuk memprediksi risiko GERD berdasarkan pola makan.

## ABSTRACT

Hayah, Nasywa. 2026. **GERD Risk Prediction Based on Dietary Patterns Using the *Random Forest* Method.** Undergraduate Thesis. Department of Informatics Engineering, Faculty of Science and Technology, State Islamic University of Maulana Malik Ibrahim Malang. Advisors: (I) Dr. Cahyo Crysdian, M.Cs. (II) Dr. Irwan Budi Santoso, M.Kom.

**Keywords:** GERD, *Random Forest*, Risk Prediction, Dietary Patterns, Machine Learning

Gastroesophageal Reflux Disease (GERD) is a digestive disorder caused by the backflow of stomach acid into the esophagus and may be influenced by dietary patterns. Delays in identifying GERD risk factors can reduce patients' quality of life and lead to health complications. This study aims to develop and evaluate a GERD risk prediction model based on dietary patterns using the *Random Forest* algorithm. The dataset used was the Acid Reflux Dataset for Classification obtained from Kaggle, consisting of 13 independent variables related to dietary patterns and one target variable representing GERD risk. The research stages included data preparation, data normalization using Min-Max Scaling, handling class imbalance using the Synthetic Minority Oversampling Technique (SMOTE), *Random Forest* model training, and model evaluation using a confusion matrix, accuracy, precision, recall, and *F1-Score*. Model evaluation was conducted using nine combinations of *n\_estimators* and *max\_depth* parameters. The results showed that the best model was obtained in scenario A3 with *n\_estimators* = 50 and *max\_depth* = None, achieving an accuracy of 89.60%, precision of 74.04%, recall of 51.68%, and *F1-Score* of 60.87%. An additional experiment using *Random UnderSampling* achieved an accuracy of 75.83%, precision of 79.84%, recall of 69.12%, and *F1-Score* of 74.10%, indicating that the model performed well on a balanced dataset. Based on these findings, the *Random Forest* algorithm can be used to predict GERD risk based on dietary patterns.

## البحث مستخلص

حياة، نسوي قُرّة ١. ٢٠٢٦. التنبؤ بمخاطر مرض الارتجاع المعدي المريئي (GERD) بناءً على النمط الغذائي باستخدام خوارزمية الغابة العشوائية (*Random Forest*) رسالة جامعية. قسم تقنية المعلومات، كلية العلوم والتكنولوجيا، جامعة مولانا مالك إبراهيم الإسلامية الحكومية مالانج. المشرف الأول: الدكتور كاهيو كريسديان، الماجستير. المشرف الثاني: الدكتور إروان بودي سانتوسو، الماجستير.

**الكلمات المفتاحية:** الارتجاع المعدي المريئي (GERD)، الغابة العشوائية (*Random Forest*)، التنبؤ بالمخاطر، النمط الغذائي، تعلم الآلة

يُعد مرض الارتجاع المعدي المريئي (GERD) أحد اضطرابات الجهاز الهضمي الناتجة عن ارتجاع أحماض المعدة إلى المريء، وقد يتأثر بالنمط الغذائي للفرد. إن التأخر في التعرف على عوامل خطر الإصابة بهذا المرض قد يؤدي إلى انخفاض جودة الحياة وظهور مضاعفات صحية. تهدف هذه الدراسة إلى بناء وتقييم نموذج للتنبؤ بمخاطر الإصابة بمرض الارتجاع المعدي المريئي اعتماداً على النمط الغذائي باستخدام خوارزمية الغابة العشوائية (*Random Forest*). استُخدمت في هذه الدراسة مجموعة بيانات *Acid Reflux Dataset for Classification* المأخوذة من منصة *Kaggle*، والتي تتكون من ثلاثة عشر متغيراً مستقلاً يتعلق بالنمط الغذائي ومتغير هدف واحد يمثل خطر الإصابة بالمرض. شملت مراحل البحث إعداد البيانات، وتطبيعها باستخدام *Min-Max Scaling*، ومعالجة عدم توازن البيانات باستخدام *Synthetic Minority Oversampling Technique (SMOTE)*، ثم تدريب نموذج الغابة العشوائية وتقييمه باستخدام مصفوفة الالتباس (*Confusion Matrix*) ومقاييس *Precision* و *Recall* و *F1-Score*. أجريت عملية التقييم من خلال تسعة سيناريوهات مختلفة تجمع بين معاملي *n\_estimators* و *max\_depth*. وأظهرت النتائج أن أفضل نموذج تحقق في السيناريو A3 باستخدام *n\_estimators = 50* و *max\_depth = None*، حيث حقق دقة (Accuracy) بلغت 89.60%، و *Precision* بنسبة 74.04%، و *Recall* بنسبة 51.68%، و *F1-Score* بنسبة 60.87% كما أُجري اختبار إضافي باستخدام *Random UnderSampling*، وحققت *Accuracy* بنسبة 75.83%، و *Precision* بنسبة 79.84%، و *Recall* بنسبة 69.12%، و *F1-Score* بنسبة 74.10%، مما يدل على أن النموذج أظهر أداءً جيداً عند استخدام بيانات متوازنة الفئات. وبناءً على هذه النتائج، يمكن استخدام خوارزمية الغابة العشوائية للتنبؤ بمخاطر الإصابة بمرض الارتجاع المعدي المريئي اعتماداً على النمط الغذائي.

# **BAB I**

## **PENDAHULUAN**

### **1.1 Latar Belakang**

Deteksi dini penyakit menjadi salah satu pilar fundamental dalam upaya pencegahan dan pengendalian risiko kesehatan. Pendekatan diagnostik konvensional yang banyak digunakan hingga saat ini masih mengandalkan respons terhadap gejala yang sudah muncul, sehingga seringkali penyakit terdeteksi pada fase lanjut. Sebaliknya, perkembangan teknologi mutakhir, khususnya dalam bidang machine learning, membuka peluang untuk melakukan prediksi risiko secara antisipatif melalui analisis data kesehatan individu.

Penerapan prediksi dalam dunia kesehatan dapat diwujudkan melalui identifikasi risiko penyakit berdasarkan berbagai faktor pemicunya, salah satunya adalah kebiasaan makan. Pendekatan berbasis pola makan ini menawarkan peluang untuk melakukan intervensi dini sebelum gejala klinis muncul. Salah satu gangguan kesehatan yang memiliki keterkaitan erat dengan faktor diet adalah *Gastroesophageal Reflux Disease* (GERD), yakni kondisi refluks asam lambung ke esofagus yang dapat dipicu oleh konsumsi makanan tertentu, pola makan yang tidak teratur, serta gaya hidup yang kurang sehat.

Prevalensi GERD di seluruh dunia terus menunjukkan tren peningkatan dari waktu ke waktu. Berdasarkan kajian sistematis dan meta-analisis oleh Nirwan dkk. (2020), tingkat prevalensi global GERD tercatat sebesar 13,98% dari keseluruhan populasi dunia, dengan disparitas yang cukup mencolok antar kawasan geografis. Data dari kawasan Timur Tengah pun menunjukkan temuan serupa, di mana prevalensi *erosive esophagitis*—salah satu bentuk komplikasi GERD—mencapai 23,8%, dengan faktor risiko yang teridentifikasi meliputi hiatal hernia, peningkatan indeks massa tubuh (BMI), dan kebiasaan merokok (Alsahafi dkk., 2024). Di Indonesia, angka prevalensi juga tergolong tinggi, yaitu sebesar 55,6% pada populasi dewasa di RSUD Karawang, yang turut dipengaruhi oleh faktor pola makan seperti konsumsi makanan pedas serta penggunaan obat-obatan tertentu (Fatimah dkk., 2024).

Angka kejadian yang tinggi tersebut mengindikasikan bahwa GERD merupakan masalah kesehatan yang cukup serius, sehingga upaya pencegahan melalui identifikasi risiko secara dini menjadi sangat penting. Deteksi dini memungkinkan intervensi dilakukan sebelum gejala klinis muncul dan kondisi memburuk. Studi terbaru menunjukkan korelasi yang signifikan antara GERD dan peningkatan indeks massa tubuh (BMI). Temuan tersebut mengindikasikan bahwa individu dengan GERD rata-rata memiliki nilai BMI yang lebih besar secara bermakna dibandingkan dengan kelompok non-GERD (Tang dkk., 2026). Hal ini semakin menegaskan bahwa faktor gaya hidup, terutama pola makan dan komposisi tubuh, memegang peranan signifikan dalam risiko GERD.

Analisis berbasis data Global Burden of Disease (GBD) 2023 mengidentifikasi bahwa konsumsi minuman berpemanis (sugar-sweetened beverages) memiliki asosiasi positif terhadap prevalensi GERD, dengan nilai Risk Ratio (RR) sebesar 1,51 (rentang kepercayaan 95%: 1,25–1,82) (Deng dkk., 2026). Temuan ini menguatkan bukti bahwa faktor diet, terutama konsumsi jenis minuman dan makanan tertentu, merupakan kontributor signifikan terhadap beban penyakit GERD secara global. Penelitian yang melibatkan mahasiswa di Tiongkok juga mengungkapkan bahwa konsumsi makanan gorengan dan kebiasaan makan di atas jam 9 malam berperan sebagai faktor risiko yang cukup kuat terhadap munculnya gejala refluks laringofaring (LPR), yang memiliki kaitan erat dengan GERD. Nilai Odds Ratio (OR) untuk kedua faktor tersebut berturut-turut adalah 1,89 dan 2,15 (Li dkk., 2026). Hal ini mengindikasikan bahwa perilaku makan dan kebiasaan hidup, terutama pada kelompok usia produktif, memiliki kontribusi penting terhadap perkembangan gejala refluks.

Dalam pandangan Islam, menjaga kesehatan merupakan bagian dari upaya melindungi diri dari hal-hal yang dapat menimbulkan mudarat. Allah SWT berfirman dalam QS. Al-Baqarah ayat 195:

وَأَنْفِقُوا فِي سَبِيلِ اللَّهِ وَلَا تُلْقُوا بِأَيْدِيكُمْ إِلَى التَّهْلُكَةِ وَأَحْسِنُوا إِنَّ اللَّهَ يُحِبُّ الْمُحْسِنِينَ ١٩٥

*“Dan infakkanlah (hartamu) di jalan Allah, dan janganlah kamu jatuhkan (diri sendiri) ke dalam kebinasaan dengan tangan sendiri, dan berbuat baiklah. Sungguh, Allah menyukai orang-orang yang berbuat baik.” (QS. Al-Baqarah : 195)*

Ayat di atas mengandung pesan bahwa setiap individu dianjurkan untuk menjaga diri dari hal-hal yang berpotensi membahayakan, termasuk dalam aspek

kesehatan. Pemeliharaan kesehatan tidak hanya dilakukan saat penyakit sudah menyerang, tetapi juga melalui tindakan antisipatif dan kewaspadaan dini.

Allah SWT juga berfirman dalam QS. Al-Hasyr ayat 18:

يَا أَيُّهَا الَّذِينَ ءَامَنُوا اتَّقُوا اللَّهَ وَلْتَنظُرْ نَفْسٌ مَّا قَدَّمَتْ لِغَدٍ وَاتَّقُوا اللَّهَ ۚ إِنَّ اللَّهَ خَبِيرٌ بِمَا تَعْمَلُونَ ١٨

*“Wahai orang-orang yang beriman! Bertakwalah kepada Allah dan hendaklah setiap orang memperhatikan apa yang telah diperbuatnya untuk hari esok, dan bertakwalah kepada Allah. Sungguh, Allah Mahateliti terhadap apa yang kamu kerjakan.” (QS. Al-Hasyr : 18)*

Ayat ini menekankan pentingnya introspeksi diri dan persiapan untuk masa depan. Nilai-nilai tersebut dapat diaplikasikan dalam berbagai aspek kehidupan, termasuk dalam menjaga kesehatan melalui deteksi dini dan pencegahan risiko penyakit.

Berbagai studi sebelumnya telah berupaya mengembangkan sistem deteksi untuk penyakit lambung, termasuk GERD. Sebagian besar pendekatan yang digunakan masih bertumpu pada analisis gejala klinis dan mengandalkan metode seperti forward chaining serta Naïve Bayes. Dewantika dkk. (2022) dan Sari dkk. (2025) merancang sistem pakar berbasis forward chaining yang mengadopsi aturan-aturan dari para ahli untuk mendiagnosis penyakit lambung. Di sisi lain, Ratnasari dkk. (2025) mengombinasikan metode *forward chaining* dan *Naïve Bayes* untuk mendeteksi penyakit lambung dan memperoleh tingkat akurasi yang cukup baik. Ramadhan dkk. (2025) juga menerapkan metode *Support Vector Machine* untuk klasifikasi GERD berdasarkan data gejala pasien. Meskipun memberikan hasil yang memuaskan, pendekatan-pendekatan tersebut masih bergantung pada data gejala yang bersifat reaktif.

Sementara itu, berbagai penelitian telah membuktikan bahwa algoritma Random Forest menunjukkan performa yang kompetitif dalam prediksi beragam penyakit. Edric dan Tamba (2022) mengaplikasikan Random Forest untuk memprediksi penyakit jantung, Kesuma dkk. (2023) menggunakannya pada kasus penyakit liver, sedangkan Christian dan Yani (2025) menerapkannya untuk prediksi penyakit paru-paru. Hasil-hasil penelitian tersebut mengindikasikan bahwa *Random Forest* mampu menangani data dengan kompleksitas tinggi dan hubungan antar variabel yang tidak linear dengan performa yang baik. Selain itu, Rosidah dkk. (2025) membuktikan bahwa pendekatan prediksi berbasis data pola makan dengan algoritma *Random Forest* dapat digunakan untuk memprediksi risiko obesitas. Studi tersebut memiliki kesamaan dengan penelitian ini dalam hal penggunaan data pola makan sebagai dasar prediksi, meskipun objek penyakit yang dikaji berbeda.

Namun demikian, sebagian besar penelitian terdahulu masih terfokus pada data gejala atau belum secara khusus membahas prediksi risiko GERD berbasis data pola makan. Selain itu, implementasi algoritma *Random Forest* pada kasus GERD dengan pendekatan berbasis faktor risiko masih terbatas. Hal ini membuka peluang untuk mengembangkan model prediksi yang lebih preventif dengan memanfaatkan data pola makan serta menggunakan metode yang mampu mengolah kompleksitas data secara optimal.

Berdasarkan uraian permasalahan tersebut, penelitian ini bertujuan untuk membangun model prediksi risiko GERD menggunakan algoritma *Random Forest* berbasis data pola makan. Melalui pendekatan ini, diharapkan model yang

dihasilkan dapat memberikan prediksi yang lebih baik serta berfungsi sebagai alat bantu dalam deteksi dini risiko GERD.

## 1.2 Pernyataan Masalah

Bagaimana performa model *Random Forest* dalam memprediksi risiko *Gastroesophageal Reflux Disease* (GERD) berdasarkan data pola makan?

## 1.3 Batasan Masalah

Penelitian ini memiliki batasan sebagai berikut:

1. Data yang digunakan adalah *dataset* publik *Acid Reflux Dataset for Classification* yang diperoleh dari platform Kaggle.
2. Variabel yang dianalisis dalam penelitian ini mengikuti atribut yang tersedia pada *dataset*, yang mencakup kebiasaan makan, jenis pola diet, serta frekuensi konsumsi berbagai jenis makanan dan minuman.

## 1.4 Tujuan Penelitian

Tujuan dari penelitian ini adalah untuk mengevaluasi kinerja model dalam memprediksi risiko GERD dengan menggunakan metrik evaluasi yang sesuai.

## 1.5 Manfaat Penelitian

Penelitian ini diharapkan memberikan manfaat sebagai berikut:

1. Memberikan kontribusi dalam pengembangan penerapan *machine learning* di bidang kesehatan, khususnya dalam konteks prediksi risiko GERD.

2. Menyediakan informasi awal mengenai risiko GERD berdasarkan hasil prediksi model sebagai bagian dari upaya deteksi dini.
3. Menjadi acuan dalam pengembangan model prediksi dengan mempertimbangkan pemilihan parameter yang optimal.

## **BAB II**

### **STUDI PUSTAKA**

*Gastroesophageal Reflux Disease* (GERD) tergolong sebagai kelainan fungsional saluran cerna dengan prevalensi tinggi di populasi umum. Kondisi ini terjadi ketika asam lambung mengalami refluks ke kerongkongan, memicu gejala klinis seperti sensasi terbakar di ulu hati (*heartburn*), rasa asam atau pahit di mulut (regurgitasi), serta rasa tidak nyaman di dada (Bertin dkk., 2025).

GERD termasuk ke dalam spektrum gangguan gastrointestinal fungsional, yang dicirikan oleh tidak adanya kelainan struktural namun tetap terjadi disfungsi motilitas (Lee & Kim, 2025). Berbagai faktor berkontribusi terhadap munculnya GID, di antaranya gaya hidup tidak sehat seperti manajemen stres yang buruk, pola tidur yang tidak teratur, minimnya aktivitas olahraga, serta pola konsumsi makanan tinggi lemak dan rendah serat (Lee & Kim, 2025).

Penelitian terbaru memperkuat bukti keterkaitan antara perilaku makan dan risiko GERD. Survei cross-sectional terhadap 280 pasien GERD di Pakistan melaporkan bahwa frekuensi makan tinggi ( $\geq 3$  kali/hari) berhubungan dengan peningkatan risiko, dengan 88,57% responden mengonfirmasi kebiasaan tersebut. Selain itu, konsumsi makanan tinggi lemak dan gorengan terbukti memperlambat pengosongan lambung dan memperburuk gejala karena meningkatkan volume asam yang tersedia untuk refluks ke esofagus (Khan dkk., 2024).

Tren peningkatan prevalensi GERD sejalan dengan pergeseran gaya hidup kontemporer, terutama konsumsi makanan tinggi lemak dan garam, rendah serat,

serta pola makan yang tidak teratur. Data epidemiologi global mengonfirmasi bahwa GERD kini menjadi salah satu beban kesehatan utama di berbagai negara, dengan kecenderungan peningkatan yang konsisten lintas wilayah (Nirwan dkk., 2020).

Sejumlah studi telah memetakan berbagai faktor risiko GERD yang mencakup dimensi demografis, perilaku, dan antropometri. Alsahafi dkk. (2024) mengidentifikasi hiatal hernia, indeks massa tubuh yang meningkat, dan kebiasaan merokok sebagai prediktor penting terhadap munculnya *erosive esophagitis*.

Sementara itu, penelitian oleh Tang dkk. (2026) mengungkapkan bahwa pasien GERD memiliki nilai BMI yang secara signifikan lebih tinggi dibandingkan kelompok kontrol ( $p < 0,001$ ), yang mengindikasikan bahwa status gizi turut berperan dalam risiko GERD. Pola makan, khususnya konsumsi makanan berlemak, makanan asam, dan minuman berkafein, juga diketahui dapat memicu atau memperburuk gejala GERD. Dengan demikian, pendekatan prediksi risiko yang memperhitungkan faktor pola makan dan gaya hidup menjadi penting untuk dikembangkan sebagai upaya deteksi dini yang lebih komprehensif.

Seiring dengan kemajuan teknologi, pemanfaatan kecerdasan buatan (*Artificial Intelligence*) dan *machine learning* mulai banyak diadopsi dalam dunia kesehatan, termasuk untuk keperluan prediksi dan klasifikasi penyakit. Pendekatan ini memungkinkan analisis data yang lebih kompleks, sehingga dapat meningkatkan akurasi dalam proses deteksi dini penyakit, tidak terkecuali GERD.

Ramadhan dkk. (2025) mengimplementasikan algoritma Support Vector Machine (SVM) untuk mengidentifikasi pasien GERD berdasarkan data klinis yang

dikumpulkan dari RSUD Sakinah Lhokseumawe. Sebanyak dua belas indikator klinis, di antaranya nyeri dada, regurgitasi, gangguan pola tidur, dan kesulitan menelan, digunakan sebagai variabel masukan dalam proses diagnosis. Selain itu, penelitian tersebut juga menerapkan teknik normalisasi dan penanganan data tidak seimbang untuk meningkatkan performa model. Berdasarkan hasil evaluasi, model SVM memperoleh tingkat akurasi 82,5% dengan F1-Score 58,3% (Ramadhan dkk., 2025).

Penelitian oleh Dewantika dkk. (2022) menghasilkan sebuah sistem pakar berbasis web yang mengadopsi teknik forward chaining dan certainty factor untuk mendeteksi GERD. Mekanisme kerja sistem ini meniru pola penalaran seorang pakar dalam mengambil keputusan diagnosis berdasarkan gejala yang diinputkan oleh pengguna. Evaluasi terhadap sistem menunjukkan bahwa tingkat kepercayaan diagnosis mencapai 81%, dan sistem dinilai bermanfaat bagi masyarakat dalam upaya pengenalan dini gejala GERD (Dewantika dkk., 2022).

Ratnasari dkk. (2025) mengusulkan pendekatan hibrida yang menggabungkan forward chaining dan Naïve Bayes untuk mendeteksi berbagai penyakit lambung. Data yang digunakan dalam studi tersebut diperoleh dari catatan gejala dan diagnosis pasien di fasilitas kesehatan, yang selanjutnya diproses menjadi basis aturan dan model probabilistik. Dari hasil pengujian, sistem yang dikembangkan mencapai akurasi 81,25% dan berhasil mendiagnosis dengan tepat sebanyak 80 pasien pada data uji (Ratnasari dkk., 2025).

Afni dkk. (2021) menerapkan klasifikasi Naïve Bayes untuk memprediksi penyakit lambung dengan mengandalkan data gejala klinis yang bersumber dari

Kaggle, berjumlah 50 sampel. Dari proses pengujian, model yang dibangun menghasilkan akurasi 75% dan nilai AUC 0,852, yang mengindikasikan performa yang cukup baik meskipun dengan jumlah data yang terbatas (Afni dkk., 2021).

Sari dkk. (2025) merancang sebuah sistem pakar yang mengandalkan forward chaining untuk mendiagnosis dini penyakit lambung di Klinik Dr. Yolanda. Basis pengetahuan sistem ini disusun berdasarkan gejala-gejala yang dikumpulkan dari pakar dan tinjauan literatur medis. Hasil uji coba menunjukkan bahwa sistem mampu mengidentifikasi berbagai jenis gangguan lambung, termasuk GERD, gastritis, dan gastroparesis, dengan tingkat akurasi yang baik (Sari dkk., 2025).

Berdasarkan kajian terhadap beberapa penelitian terdahulu, terlihat bahwa mayoritas studi terkait deteksi penyakit lambung, termasuk GERD, masih didominasi oleh pendekatan berbasis aturan (*rule-based*) dan metode probabilistik seperti *Forward Chaining* dan *Naïve Bayes*. Pendekatan-pendekatan tersebut umumnya mengandalkan data gejala klinis sebagai dasar dalam proses diagnosis, sehingga masih memiliki keterbatasan dalam menangkap pola kompleks yang dipengaruhi oleh faktor pola makan.

Di sisi lain, perkembangan *machine learning* menunjukkan bahwa algoritma berbasis *ensemble*, seperti *Random Forest*, memiliki performa yang baik dalam berbagai kasus prediksi penyakit. Namun, implementasi metode *Random Forest* dalam konteks penyakit lambung, khususnya yang memanfaatkan data pola makan, masih relatif terbatas. Oleh karena itu, penting untuk meninjau penelitian-penelitian lain yang telah berhasil menerapkan *Random Forest* dalam prediksi penyakit sebagai dasar pertimbangan dalam penelitian ini.

Penelitian oleh Edric dan Tamba (2022) berfokus pada prediksi penyakit jantung melalui penerapan Random Forest pada dataset rekam medis yang mencakup 918 sampel dan 12 variabel, di antaranya usia, tekanan darah, kolesterol, dan temuan EKG. Model awal memperoleh akurasi 82,60%, yang kemudian meningkat menjadi 85,05% setelah dilakukan optimasi hyperparameter menggunakan K-Fold Cross Validation dan GridSearchCV (Edric & Tamba, 2022).

Kesuma dkk. (2023) mengaplikasikan Random Forest pada Liver Disease Patient Dataset untuk memprediksi penyakit liver. Proses penelitian mencakup eksplorasi data, preprocessing, hingga pembangunan model. Berdasarkan evaluasi, Random Forest memperoleh akurasi 71,33% dan F1-Score 81%, yang menunjukkan bahwa metode ini cukup efektif untuk kasus tersebut (Kesuma dkk., 2023).

Christian dan Yani (2025) menguji performa Random Forest untuk prediksi penyakit paru-paru menggunakan dataset publik Kaggle yang terdiri dari sekitar 30.000 entri dengan 11 variabel, termasuk faktor gaya hidup seperti kebiasaan merokok, tingkat aktivitas, dan kualitas tidur. Hasil eksperimen menunjukkan bahwa Random Forest mencapai akurasi 94,7%, precision 0,952, recall 0,947, dan F1-Score 0,946, yang mengungguli metode lain seperti K-Nearest Neighbors, C4.5, dan Naïve Bayes (Christian & Yani, 2025).

Rosidah dkk. (2025) mengkaji penerapan algoritma *Random Forest* untuk klasifikasi risiko obesitas berdasarkan pola makan. *Dataset* yang digunakan merupakan *dataset* publik yang terdiri dari 1.610 data dengan 15 atribut yang berkaitan dengan kebiasaan makan dan gaya hidup. Tahapan penelitian

meliputi *preprocessing data*, pembagian data latih dan data uji, implementasi model, serta optimasi *hyperparameter* menggunakan metode *Grid Search*. Hasil penelitian menunjukkan bahwa model *Random Forest* yang telah dioptimasi mencapai akurasi sebesar 85,4% dan mampu mengidentifikasi faktor-faktor penting yang memengaruhi risiko obesitas, seperti frekuensi konsumsi makanan cepat saji, jumlah makan utama per hari, kebiasaan ngemil, dan konsumsi sayuran. Studi ini membuktikan bahwa pendekatan berbasis pola makan dengan algoritma *Random Forest* dapat digunakan secara efektif untuk membangun model prediksi risiko penyakit secara preventif (Rosidah dkk., 2025).

Berdasarkan kajian terhadap penelitian-penelitian terdahulu, pendekatan dalam mendeteksi atau memprediksi GERD masih didominasi oleh metode klasifikasi sederhana yang hanya mempertimbangkan gejala klinis. Metode-metode tersebut umumnya memanfaatkan data seperti nyeri dada, regurgitasi, dan gangguan pencernaan, sehingga belum mampu menangkap faktor gaya hidup secara komprehensif. Selain itu, pemanfaatan data gaya hidup seperti pola makan dan kebiasaan sehari-hari masih relatif terbatas dalam penelitian terkait GERD.

Di sisi lain, terdapat penelitian yang membuktikan bahwa data pola makan dapat digunakan dalam membangun model prediksi penyakit menggunakan algoritma *Random Forest*, seperti pada penelitian Rosidah dkk. (2025) yang mengkaji prediksi risiko obesitas berbasis pola makan. Hal ini mengindikasikan bahwa pendekatan berbasis gaya hidup memiliki potensi untuk dikembangkan lebih lanjut dalam konteks prediksi penyakit. Oleh karena itu, diperlukan pendekatan

yang mampu mengolah berbagai faktor secara simultan, termasuk data gaya hidup, salah satunya dengan menggunakan algoritma *Random Forest*.

Random Forest adalah salah satu pendekatan ensemble learning yang mengombinasikan prediksi dari sejumlah besar pohon keputusan untuk mencapai hasil akhir yang lebih stabil dan presisi. Kelebihan utama metode ini adalah ketahanannya terhadap overfitting dan kemampuannya menghasilkan akurasi kompetitif meskipun menggunakan parameter default tanpa tuning ekstensif (Lee & Kim, 2025).

Perkembangan mutakhir dalam bidang *machine learning* menunjukkan adanya pergeseran pendekatan dari metode berbasis aturan (*rule-based*) menuju model *ensemble* yang lebih adaptif dan mampu menangani kompleksitas data yang lebih tinggi. Delfina dkk. (2025) menerapkan model *XGBoost* untuk deteksi dini GERD menggunakan data pola makan dan gaya hidup. Penelitian tersebut menggunakan teknik SMOTE untuk mengatasi ketidakseimbangan data serta seleksi fitur menggunakan *Pearson Correlation* dan *Principal Component Analysis* (PCA). Hasil penelitian menunjukkan bahwa model *XGBoost* dengan seleksi fitur mencapai nilai *F1-Score* sebesar 0,9809 dan akurasi sebesar 0,9905, yang membuktikan bahwa pendekatan *machine learning* dengan penanganan data yang tepat dapat menghasilkan performa yang sangat baik dalam prediksi GERD (Wisesty dkk., 2025).

Selain *XGBoost*, algoritma *Random Forest* juga telah diterapkan dalam berbagai penelitian terkait GERD dan penyakit pencernaan lainnya. Wang dkk. (2026) mengembangkan model *machine learning* berbasis *Random Forest* untuk

mengidentifikasi *Reflux Esophagitis* (salah satu subtype GERD) menggunakan data komposisi tubuh yang diperoleh dari CT scan abdomen. Model *Random Forest* yang dikembangkan berhasil mencapai nilai AUC sebesar 0,829, dengan analisis *Feature Importance* menggunakan SHAP (*SHapley Additive exPlanations*) menunjukkan bahwa faktor seperti lemak antar-otot (*intermuscular adipose tissue*), rasio lemak visceral/subkutan, dan BMI merupakan prediktor penting dalam identifikasi risiko *Reflux Esophagitis* (Wang dkk., 2026).

Lee dkk. (2024) menggunakan *Random Forest* untuk menganalisis data populasi yang menghubungkan GERD dengan kejadian kelahiran prematur. Hasil penelitian menunjukkan bahwa GERD memiliki hubungan yang kuat dengan kelahiran prematur dengan nilai kontribusi SHAP sebesar 0,27, yang mengindikasikan bahwa dampak GERD tidak hanya terbatas pada gangguan pencernaan tetapi juga dapat memengaruhi kondisi kesehatan lainnya (Lee dkk., 2024). Temuan ini semakin memperkuat pentingnya deteksi dini risiko GERD sebagai upaya pencegahan komplikasi kesehatan yang lebih luas.

Owess dkk. (2024) membandingkan berbagai model *supervised learning* untuk memprediksi peningkatan kadar gula darah, termasuk *Random Forest* dan *Adaboost*. Hasil penelitian menunjukkan bahwa *Random Forest* mencapai akurasi 98,4%, lebih tinggi dibandingkan *Adaboost* yang mencapai 95,2% (Owess dkk., 2024). Hal ini mengindikasikan bahwa *Random Forest* cenderung lebih unggul dalam menangani data medis dengan kompleksitas tinggi.

Berdasarkan penelitian-penelitian yang telah dikaji, terlihat bahwa algoritma *ensemble* seperti *Random Forest* dan *XGBoost* memiliki performa yang baik dalam menangani data dengan kompleksitas tinggi serta hubungan antar variabel yang tidak linear. Hal ini menunjukkan bahwa pendekatan berbasis *machine learning*, khususnya dengan memanfaatkan data pola makan dan gaya hidup, memiliki potensi besar untuk dikembangkan dalam prediksi risiko GERD. Penelitian ini bertujuan untuk mengisi celah tersebut dengan mengimplementasikan algoritma *Random Forest* untuk memprediksi risiko GERD berdasarkan data pola makan, serta mengevaluasi pengaruh variasi parameter terhadap performa model.

## BAB III

### DESAIN DAN IMPLEMENTASI SISTEM

#### 3.1 Persiapan Data

Tahap awal dalam penelitian ini adalah persiapan data yang bertujuan untuk memahami karakteristik *dataset* serta memastikan data siap digunakan dalam proses pelatihan dan pengujian model.

Data yang digunakan merupakan *dataset* publik berjudul *Acid Reflux Dataset for Classification* yang diakses melalui platform Kaggle (<https://www.kaggle.com/datasets/ahmetsametgrkan/acid-reflux-dataset-for-classification/data>). *Dataset* ini memuat informasi terkait pola konsumsi makanan, kebiasaan hidup, serta status risiko GERD.

*Dataset* ini tersusun atas 14 variabel, yang terdiri dari 13 variabel independen (*feature*) dan 1 variabel dependen (*target*) bernama *Reflu*. Variabel *Reflu* merepresentasikan status kesehatan seseorang, di mana nilai 1 mengindikasikan mengalami GERD dan nilai 0 mengindikasikan tidak mengalami GERD. Jumlah total data dalam *dataset* ini adalah 4.757 baris.

Secara garis besar, fitur-fitur yang tersedia mencerminkan aspek pola konsumsi pangan, meliputi jenis diet (omnivora, vegetarian, atau vegan), serta frekuensi konsumsi berbagai jenis makanan dan minuman seperti buah, sayuran, daging merah, makanan berlemak, makanan asin, minuman beralkohol, dan konsumsi air harian. Seluruh variabel dalam *dataset* telah melalui proses *encoding* ordinal dan biner sehingga berbentuk numerik, sehingga tidak

memerlukan proses *encoding* tambahan sebelum pelatihan model. Penjelasan lebih rinci mengenai fitur-fitur tersebut disajikan pada Tabel 3.1.

Tabel 3.1 Rincian Fitur Dataset

| No | Fitur                                                    | Jenis Data          | Deskripsi                                                    | Kategori / Skala Nilai             |
|----|----------------------------------------------------------|---------------------|--------------------------------------------------------------|------------------------------------|
| 1  | DO ( <i>diet_Omn</i> )                                   | Kategorikal (biner) | Jenis diet: Omnivora (makan semua jenis makanan)             | 1 = Ya, 0 = Tidak                  |
| 2  | DVg ( <i>diet_Veg</i> )                                  | Kategorikal (biner) | Jenis diet: Vegetarian                                       | 1 = Ya, 0 = Tidak                  |
| 3  | DVgt ( <i>diet_Vegt</i> )                                | Kategorikal (biner) | Jenis diet: Vegan / semi-vegetarian                          | 1 = Ya, 0 = Tidak                  |
| 4  | FF ( <i>fruit_frequency_encoded</i> )                    | Ordinal             | Frekuensi konsumsi buah                                      | 0 = Tidak pernah → 4 = Setiap hari |
| 5  | HFRMF ( <i>high_fat_red_meat_frequency_encoded</i> )     | Ordinal             | Frekuensi konsumsi daging merah berlemak                     | 0 = Tidak pernah → 4 = Setiap hari |
| 6  | HMF ( <i>homecooked_meals_frequency_encoded</i> )        | Ordinal             | Frekuensi makan masakan rumahan                              | 0 = Tidak pernah → 4 = Setiap hari |
| 7  | VF ( <i>vegetable_frequency_encoded</i> )                | Ordinal             | Frekuensi konsumsi sayuran                                   | 0 = Tidak pernah → 4 = Setiap hari |
| 8  | AF ( <i>alcohol_frequency_encoded</i> )                  | Ordinal             | Frekuensi konsumsi alkohol                                   | 0 = Tidak pernah → 4 = Setiap hari |
| 9  | FDF ( <i>frozen_dessert_frequency_encoded</i> )          | Ordinal             | Frekuensi konsumsi makanan pencuci mulut beku (es krim, dll) | 0 = Tidak pernah → 4 = Setiap hari |
| 10 | MCF ( <i>milk_cheese_frequency_encoded</i> )             | Ordinal             | Frekuensi konsumsi susu dan keju                             | 0 = Tidak pernah → 4 = Setiap hari |
| 11 | WF ( <i>one_liter_of_water_a_day_frequency_encoded</i> ) | Ordinal             | Kebiasaan minum $\geq 1$ liter air per hari                  | 0 = Tidak pernah → 4 = Setiap hari |
| 12 | SF ( <i>salted_snacks_frequency_encoded</i> )            | Ordinal             | Frekuensi konsumsi makanan asin / tinggi garam               | 0 = Tidak pernah → 4 = Setiap hari |

|    |                                              |                     |                                  |                                    |
|----|----------------------------------------------|---------------------|----------------------------------|------------------------------------|
| 13 | RMF<br>( <i>red_meat_frequency_encoded</i> ) | Ordinal             | Frekuensi konsumsi daging merah  | 0 = Tidak pernah → 4 = Setiap hari |
| 14 | <i>Reflu</i>                                 | Kategorikal (biner) | Status GERD / <i>Acid Reflux</i> | 1 = Mengalami GERD, 0 = Tidak      |

Berdasarkan hasil analisis terhadap karakteristik *dataset*, ditemukan adanya perbedaan rentang nilai antar atribut. Variabel jenis diet menggunakan skala biner (0 dan 1), sedangkan variabel frekuensi konsumsi makanan dan minuman menggunakan rentang nilai 0 hingga 4. Perbedaan skala ini berpotensi menyebabkan fitur dengan rentang nilai yang lebih besar mendominasi proses pembentukan model.

Untuk mengatasi perbedaan skala tersebut, dilakukan normalisasi data menggunakan metode *MinMaxScaler* yang bertujuan menyeragamkan rentang nilai seluruh fitur ke dalam skala 0 hingga 1. Proses normalisasi ini penting agar setiap fitur memberikan kontribusi yang lebih seimbang dalam proses klasifikasi menggunakan algoritma *Random Forest*.

Perbandingan kondisi *dataset* sebelum dan sesudah normalisasi disajikan pada Tabel 3.2 dan Tabel 3.3.

Tabel 3.2 *Dataset* Sebelum Normalisasi

| DO | DVg | DVgt | FF | HFRMF | HMF | VF | AF | FDF | MCF | WF | SF | RMF | Reflu |
|----|-----|------|----|-------|-----|----|----|-----|-----|----|----|-----|-------|
| 0  | 0   | 1    | 2  | 0     | 3   | 3  | 2  | 0   | 3   | 2  | 2  | 0   | 1     |
| 0  | 1   | 0    | 3  | 0     | 3   | 3  | 1  | 0   | 0   | 3  | 0  | 0   | 1     |
| 0  | 0   | 1    | 2  | 0     | 2   | 2  | 0  | 0   | 1   | 4  | 1  | 0   | 0     |
| 1  | 0   | 0    | 0  | 0     | 3   | 0  | 1  | 1   | 4   | 0  | 2  | 3   | 0     |
| 0  | 0   | 1    | 3  | 0     | 0   | 3  | 1  | 2   | 3   | 4  | 1  | 0   | 0     |

Tabel 3.3 *Dataset* Setelah Normalisasi

| DO  | DVg | DVgt | FF   | HFRMF | HMF  | VF   | AF   | FDF  | MCF  | WF   | SF   | RMF  | Reflu |
|-----|-----|------|------|-------|------|------|------|------|------|------|------|------|-------|
| 1.0 | 0.0 | 0.0  | 1.00 | 0.0   | 1.00 | 1.00 | 0.25 | 0.00 | 1.00 | 0.50 | 0.25 | 0.25 | 1     |
| 1.0 | 0.0 | 0.0  | 0.50 | 1.0   | 1.00 | 1.00 | 1.00 | 0.75 | 0.25 | 0.75 | 0.00 | 0.75 | 1     |
| 0.0 | 0.0 | 1.0  | 0.75 | 0.0   | 1.00 | 1.00 | 0.50 | 0.25 | 0.00 | 0.75 | 0.25 | 0.00 | 0     |
| 1.0 | 0.0 | 0.0  | 0.75 | 0.0   | 0.75 | 0.75 | 0.00 | 0.50 | 0.50 | 1.00 | 0.25 | 0.25 | 0     |
| 0.0 | 1.0 | 0.0  | 1.00 | 0.0   | 1.00 | 1.00 | 0.00 | 0.00 | 0.00 | 1.00 | 0.50 | 0.00 | 0     |

Berdasarkan analisis distribusi kelas pada variabel target *Reflu*, diketahui bahwa *dataset* yang digunakan memiliki kondisi ketidakseimbangan (*imbalanced data*), di mana jumlah data tidak GERD jauh lebih banyak dibandingkan data GERD. Dari total 4.757 data, sebanyak 4.013 data (84,36%) termasuk dalam kategori tidak mengalami GERD (label 0), sedangkan 744 data (15,64%) termasuk dalam kategori mengalami GERD (label 1). Kondisi ini berpotensi menyebabkan model menjadi bias terhadap kelas mayoritas, sehingga kemampuan model dalam mendeteksi kelas minoritas menjadi kurang optimal.

Untuk mengatasi permasalahan tersebut, diterapkan metode *Synthetic Minority Over-sampling Technique* (SMOTE) pada data latih. SMOTE bekerja

dengan menghasilkan data sintetis baru pada kelas minoritas berdasarkan kedekatan antar data dalam ruang fitur.

Setelah penerapan SMOTE pada data latih, jumlah data pada kedua kelas menjadi seimbang, masing-masing sebanyak 3.210 data, dengan total keseluruhan 6.420 data. Perbandingan distribusi data sebelum dan sesudah penerapan SMOTE disajikan pada Tabel 3.4.

Tabel 3.4 Distribusi Data Sebelum dan Sesudah SMOTE

| <b>Kelas<br/>(Reflu)</b> | <b>Keterangan</b> | <b>Jumlah<br/>Data</b> | <b>Persentase<br/>Sebelum<br/>SMOTE</b> | <b>Jumlah<br/>Data<br/>Setelah<br/>SMOTE</b> | <b>Persentase<br/>Setelah<br/>SMOTE</b> |
|--------------------------|-------------------|------------------------|-----------------------------------------|----------------------------------------------|-----------------------------------------|
| 0                        | Tidak GERD        | 4.013                  | 84,36%                                  | 3.210                                        | 50,00%                                  |
| 1                        | Mengalami GERD    | 744                    | 15,64%                                  | 3.210                                        | 50,00%                                  |
| <b>Total</b>             |                   | <b>4.757</b>           | <b>100%</b>                             | <b>6.420</b>                                 | <b>100%</b>                             |

Penggunaan SMOTE sebagai teknik penanganan data tidak seimbang dipilih berdasarkan sejumlah penelitian yang telah membuktikan kemampuannya dalam meningkatkan kualitas model klasifikasi. Wisesty dkk. (2025) menerapkan SMOTE pada dataset GERD dengan model XGBoost untuk deteksi dini, di mana komposisi data latih awal terdiri dari 591 sampel non-probable GERD dan 369 sampel probable GERD. Setelah proses SMOTE, proporsi kedua kelas menjadi 591 sampel masing-masing (Wisesty dkk., 2025).

Hal ini mengindikasikan bahwa penanganan ketidakseimbangan data merupakan langkah krusial dalam pengembangan model prediksi yang akurat, terutama pada kasus medis di mana kelas minoritas (pasien yang benar-benar mengalami penyakit) seringkali memiliki jumlah data yang lebih sedikit

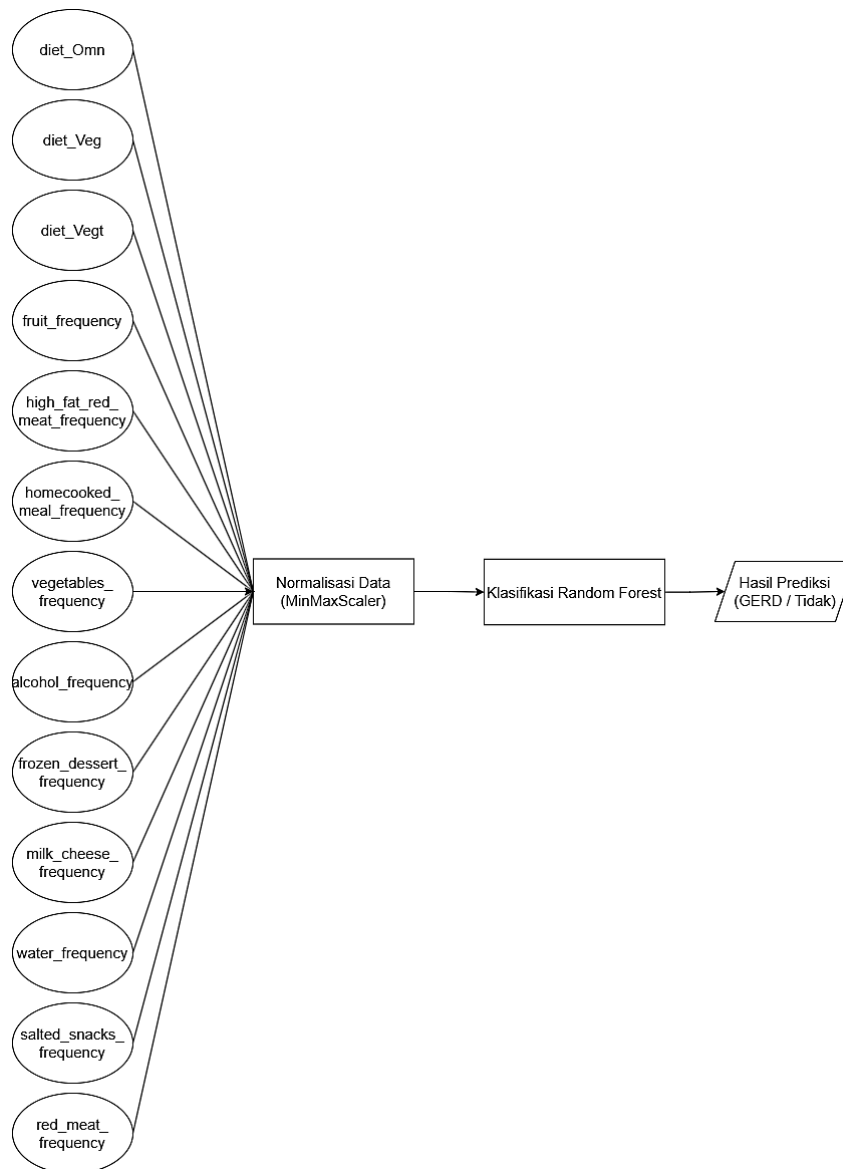
dibandingkan kelas mayoritas (pasien yang tidak mengalami penyakit). Temuan ini juga mendukung penggunaan SMOTE pada penelitian ini untuk mengatasi ketidakseimbangan data pada dataset *Acid Reflux*.

Metode SMOTE mengakomodasi ketidakseimbangan kelas dengan mensintesis sampel buatan untuk kelompok minoritas menggunakan prinsip interpolasi antar tetangga terdekat dalam ruang fitur. Strategi ini memungkinkan model mempelajari distribusi kedua kelas secara lebih proporsional.

### 3.2 Desain Sistem

Desain sistem dalam penelitian ini bertujuan untuk menggambarkan alur kerja sistem dalam memprediksi risiko GERD berdasarkan data pola makan menggunakan algoritma *Random Forest*. Sistem yang dibangun menerima masukan berupa variabel pola makan, kemudian melakukan proses *preprocessing data* sebelum dilakukan klasifikasi menggunakan model *Random Forest* untuk menghasilkan prediksi risiko GERD.

Secara umum, alur kerja sistem terdiri dari beberapa tahap, yaitu input data, preprocessing data menggunakan normalisasi *MinMaxScaler*, proses klasifikasi menggunakan *Random Forest*, dan output hasil prediksi risiko GERD. Gambaran keseluruhan desain sistem ditunjukkan pada Gambar 3.1.



Gambar 3.1 Desain Sistem Prediksi Risiko GERD

### 3.2.1 *Input Data*

Tahap *input data* merupakan langkah awal dalam sistem prediksi risiko GERD. Pada tahap ini, sistem menerima data berupa variabel-variabel yang berkaitan dengan pola makan. Variabel yang digunakan meliputi 13 fitur, seperti jenis diet, frekuensi konsumsi buah, sayur, daging merah, makanan berlemak, makanan asin, produk susu, alkohol, konsumsi air, serta beberapa kebiasaan makan lainnya yang terkait dengan risiko GERD. Seluruh atribut pada *dataset* telah tersedia dalam bentuk numerik sehingga dapat langsung digunakan dalam proses klasifikasi menggunakan algoritma *Random Forest*.

### 3.2.2 *Normalisasi Data*

Tahap normalisasi data bertujuan untuk menyeragamkan rentang nilai antar fitur sebelum data digunakan dalam proses klasifikasi. Normalisasi diperlukan karena terdapat perbedaan skala nilai pada beberapa atribut dalam *dataset*.

Dalam penelitian ini, variabel jenis diet menggunakan skala biner (0 dan 1), sedangkan variabel frekuensi konsumsi makanan dan minuman menggunakan rentang nilai 0 hingga 4. Perbedaan rentang nilai ini dapat memengaruhi proses pembentukan model karena fitur dengan rentang nilai lebih besar dapat memberikan pengaruh yang lebih dominan dibandingkan fitur lainnya.

Untuk mengatasi perbedaan skala tersebut, digunakan metode *MinMaxScaler* yang mentransformasikan nilai setiap fitur ke dalam rentang 0 hingga 1 berdasarkan nilai minimum dan maksimum dari masing-masing fitur. Proses normalisasi ini menghasilkan skala yang seragam pada seluruh fitur, sehingga dapat membantu meningkatkan stabilitas model dalam proses klasifikasi.

Adapun rumus normalisasi *MinMaxScaler* ditunjukkan pada persamaan (3.1):

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

Keterangan:

$X'$  = nilai hasil normalisasi  
 $X$  = nilai data awal  
 $X_{\min}$  = nilai minimum pada fitur  
 $X_{\max}$  = nilai maksimum pada fitur

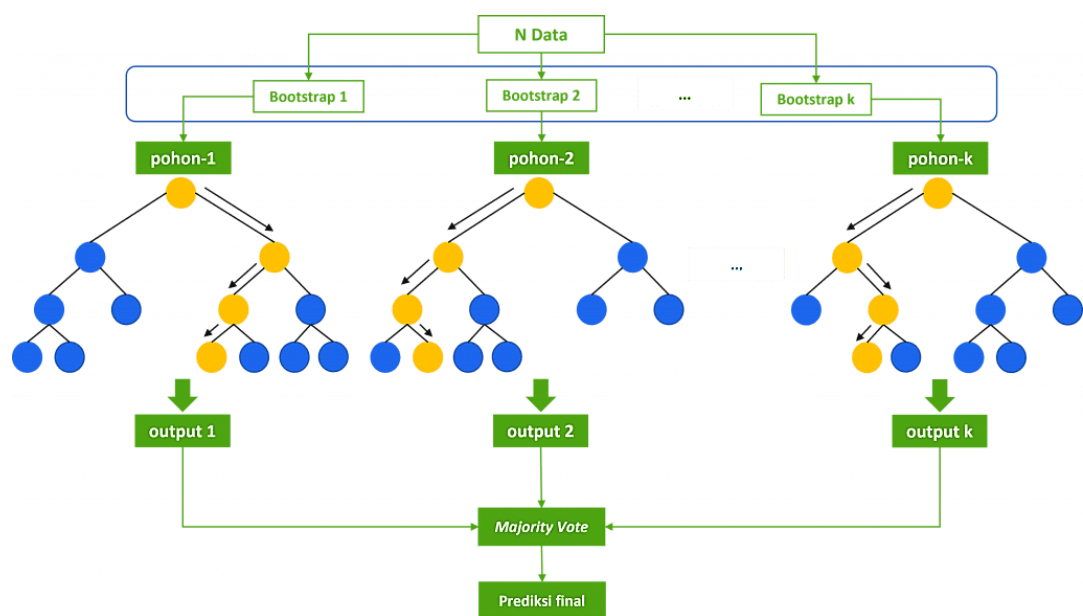
Pemilihan metode normalisasi berpengaruh signifikan terhadap performa model. *Min-Max Scaling* tidak hanya menyamakan skala fitur, tetapi terbukti mampu meningkatkan performa *Random Forest* secara keseluruhan. Penelitian komparatif tentang teknik preprocessing untuk Random Forest melaporkan bahwa penggunaan Min-Max Scaling menghasilkan performa terbaik dengan tingkat akurasi mencapai 94,72% (Haseeb dkk., 2024). Studi lain yang mengkaji implementasi Random Forest untuk keperluan klasifikasi menegaskan bahwa normalisasi Min-Max memberikan peningkatan performa paling signifikan dibandingkan alternatif normalisasi lainnya (Priyambudi & Nugroho, 2024).

### 3.2.3 Klasifikasi dengan *Random Forest*

Tahap klasifikasi dalam penelitian ini dilakukan menggunakan algoritma *Random Forest* untuk memprediksi risiko GERD berdasarkan pola makan pengguna. *Random Forest* merupakan algoritma *ensemble learning* yang bekerja dengan membangun banyak pohon keputusan (*decision tree*) untuk menghasilkan prediksi yang lebih stabil dan akurat dibandingkan penggunaan satu pohon keputusan tunggal.

Dalam algoritma *Random Forest*, proses klasifikasi dilakukan melalui beberapa tahapan utama, yaitu *bootstrap sampling*, pembentukan *decision tree*, perhitungan *Gini Index*, pembentukan *node* dan *branch*, serta *majority voting* untuk menentukan hasil klasifikasi akhir.

Desain model *Random Forest* pada penelitian ini ditunjukkan pada Gambar 3.2.



Gambar 3.2 Desain Model *Random Forest*

Sumber: <https://sainsdata.id/machine-learning/893/random-forest-untuk-model-klasifikasi-menggunakan-scikitlearn-python/>

### 1. *Bootstrap Sampling*

Tahap pertama dalam algoritma *Random Forest* adalah *bootstrap sampling*, yaitu proses pengambilan subset data secara acak dari data latih dengan metode *random with replacement*, di mana satu data dapat terpilih lebih dari satu kali dalam satu subset.

Setiap *decision tree* dibangun menggunakan subset data yang berbeda, sehingga menghasilkan variasi model pada masing-masing pohon

keputusan. Variasi ini bertujuan untuk meningkatkan stabilitas model dan mengurangi risiko *overfitting*.

Proses *bootstrap sampling* pada *Random Forest* menggunakan metode *random with replacement*. Proses ini memungkinkan setiap pohon keputusan dilatih pada subset data yang berbeda, menciptakan variasi antar pohon yang pada akhirnya meningkatkan stabilitas dan akurasi model. Data yang tidak terpilih dalam proses ini, yang disebut *Out-of-Bag* (OOB), kemudian dapat digunakan sebagai validasi internal tanpa memerlukan data uji terpisah (Malley dkk., 2011).

## 2. Pembentukan *Decision Tree*

Setelah subset data diperoleh, sistem membangun *decision tree* pada masing-masing subset. *Decision tree* bekerja dengan mempartisi data ke dalam beberapa *node* berdasarkan atribut tertentu hingga menghasilkan klasifikasi akhir.

Dalam penelitian ini, atribut yang digunakan terdiri dari variabel pola makan seperti jenis diet, frekuensi konsumsi buah, sayur, alkohol, makanan berlemak, makanan asin, produk susu, konsumsi air, serta beberapa variabel lainnya yang berkaitan dengan risiko GERD.

Pada setiap *node*, sistem mencari atribut terbaik yang mampu memisahkan data secara optimal menggunakan perhitungan *Gini Index*.

## 3. Perhitungan *Gini Index*

*Gini Index* digunakan untuk mengukur tingkat *impurity* atau ketidakhomogenan data pada suatu *node*. Semakin kecil nilai *Gini Index*,

maka data pada *node* tersebut semakin homogen, sehingga atribut tersebut dianggap semakin baik dalam memisahkan data.

Adapun rumus *Gini Index* ditunjukkan sebagai persamaan (3.2):

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

Keterangan:

$D$  = himpunan data pada *node*

$m$  = jumlah kelas target

$p_i$  = proporsi data pada kelas ke- $i$

Sebagai contoh, apabila suatu *node* terdiri dari dua kelas yaitu GERD dan Tidak GERD, maka proporsi masing-masing kelas dihitung untuk memperoleh nilai *Gini Index*. Atribut dengan nilai *impurity* terkecil dipilih sebagai atribut pemisah (*split*) terbaik.

Gini Index berfungsi sebagai metrik untuk mengevaluasi tingkat kemurnian (*impurity*) dari suatu *node* keputusan. Nilai Gini yang semakin rendah mencerminkan homogenitas data yang lebih tinggi, yang mengindikasikan bahwa atribut tersebut efektif sebagai pemisah kelas (Gwetu dkk., 2020).

#### 4. Pembentukan *Node* dan *Branch*

Berdasarkan atribut terbaik yang diperoleh dari perhitungan *Gini Index*, sistem kemudian membentuk *node* dan *branch* (cabang) pada *decision tree*. *Node* pertama disebut sebagai *root node* yang menjadi titik awal proses klasifikasi. Selanjutnya, sistem membentuk *branch* berdasarkan kondisi tertentu dari atribut yang dipilih hingga terbentuk *leaf node* yang merepresentasikan hasil klasifikasi akhir.

Proses pembentukan cabang dilakukan secara berulang hingga mencapai kondisi tertentu, seperti:

- seluruh data pada *node* berada pada kelas yang sama,
- jumlah data pada *node* terlalu sedikit,
- atau telah mencapai batas kedalaman pohon (*max\_depth*).

#### 5. *Majority Voting*

Dalam algoritma *Random Forest*, proses pembentukan *decision tree* dilakukan pada banyak pohon keputusan sekaligus. Setiap pohon akan menghasilkan prediksi terhadap data masukan.

Tahap akhir dilakukan menggunakan metode *majority voting*, yaitu menentukan hasil klasifikasi akhir berdasarkan kelas yang paling banyak diprediksi oleh seluruh pohon keputusan.

Sebagai contoh:

- Tree 1 → GERD
- Tree 2 → Tidak GERD
- Tree 3 → GERD

Berdasarkan hasil voting tersebut, prediksi akhir adalah GERD karena memperoleh jumlah voting terbanyak.

Penggunaan banyak *decision tree* pada *Random Forest* bertujuan untuk meningkatkan stabilitas model dan mengurangi risiko *overfitting* dibandingkan penggunaan satu *decision tree* saja.

### 3.2.4 Hasil Prediksi Risiko GERD

Tahap akhir dari sistem adalah menghasilkan *output* berupa prediksi risiko GERD berdasarkan data yang telah dimasukkan oleh pengguna.

*Output* yang dihasilkan berupa hasil klasifikasi risiko, yaitu:

- 0 → Tidak berisiko GERD
- 1 → Berisiko GERD

Hasil prediksi tersebut dapat digunakan sebagai informasi awal dalam memahami risiko GERD berdasarkan pola makan yang dimiliki.

### 3.3 Implementasi Sistem

Tahap implementasi sistem merupakan proses penerapan rancangan sistem ke dalam bentuk program yang dapat dijalankan untuk membangun model prediksi risiko GERD. Implementasi dilakukan menggunakan bahasa pemrograman *Python* dengan bantuan platform *Google Colab* sebagai lingkungan pengembangan berbasis *cloud computing*.

Beberapa pustaka (*library*) yang digunakan dalam penelitian ini antara lain:

- *pandas* dan *numpy* untuk pengolahan dan analisis data
- *scikit-learn* untuk implementasi algoritma *Random Forest*, pembagian data, serta evaluasi model
- *imbalanced-learn (imblearn)* untuk penerapan metode SMOTE
- *matplotlib* dan *seaborn* untuk visualisasi hasil analisis dan evaluasi model
- *joblib* untuk menyimpan model *machine learning*

Adapun langkah-langkah implementasi sistem yang dilakukan adalah sebagai berikut:

#### 3.3.1 Import Library dan Load Dataset

Pada tahap awal, seluruh pustaka yang dibutuhkan diimpor ke dalam lingkungan kerja *Python*. *Dataset* kemudian dimuat ke dalam sistem menggunakan fungsi *read\_csv()* dari pustaka *pandas*. Selanjutnya dilakukan pengecekan struktur data menggunakan fungsi seperti *shape()* dan *head()* untuk memastikan *dataset* telah terbaca dengan baik.

### 3.3.2 Pemisahan Fitur dan Target

*Dataset* dipisahkan menjadi dua bagian, yaitu fitur ( $X$ ) dan target ( $*y*$ ). Fitur mencakup seluruh variabel independen yang merepresentasikan pola makan, sedangkan target adalah variabel *Reflu* yang menunjukkan status GERD (1 = GERD, 0 = tidak GERD).

### 3.3.3 Pembagian Data dan Normalisasi Data

*Dataset* kemudian dibagi menjadi data latih (*training data*) dan data uji (*testing data*) dengan rasio 80:20. Sebanyak 80% data digunakan sebagai data latih untuk proses pelatihan model, sedangkan 20% sisanya digunakan sebagai data uji untuk mengevaluasi performa model.

Setelah proses pembagian data dilakukan, tahap selanjutnya adalah normalisasi data menggunakan metode *MinMaxScaler*. Normalisasi dilakukan untuk menyeragamkan rentang nilai antar fitur sebelum digunakan dalam proses klasifikasi menggunakan algoritma *Random Forest*.

Pada penelitian ini, beberapa atribut memiliki rentang nilai yang berbeda, seperti variabel jenis diet yang menggunakan nilai biner 0 dan 1, sedangkan variabel frekuensi konsumsi makanan dan minuman menggunakan rentang nilai 0 hingga 4.

Oleh karena itu, normalisasi diperlukan agar seluruh fitur memiliki skala yang konsisten.

Metode *MinMaxScaler* bekerja dengan mentransformasikan nilai data ke dalam rentang 0 hingga 1 berdasarkan nilai minimum dan maksimum pada masing-masing fitur.

Adapun penanganan ketidakseimbangan data menggunakan metode SMOTE (*Synthetic Minority Over-sampling Technique*) dilakukan pada tahap pengujian model menggunakan data latih.

### 3.3.4 Model *Random Forest*

Model dibangun menggunakan algoritma *Random Forest* dari pustaka *scikit-learn*. Algoritma *Random Forest* digunakan untuk melakukan klasifikasi risiko GERD berdasarkan data pola makan pengguna.

Proses pelatihan dilakukan menggunakan data latih yang telah melalui tahap normalisasi data dan penyeimbangan kelas menggunakan metode SMOTE. Pada penelitian ini, *Random Forest* bekerja dengan membangun banyak pohon keputusan (*decision tree*) menggunakan subset data yang berbeda melalui proses *bootstrap sampling*.

Setiap *decision tree* melakukan proses klasifikasi berdasarkan atribut tertentu menggunakan perhitungan *Gini Index* untuk menentukan atribut pemisah terbaik pada setiap *node*. Selanjutnya, hasil prediksi dari seluruh pohon keputusan dikombinasikan menggunakan metode *majority voting* untuk menghasilkan prediksi akhir.

Parameter yang digunakan dalam penelitian ini meliputi jumlah pohon keputusan (*n\_estimators*) dan kedalaman maksimum pohon (*max\_depth*). Variasi parameter tersebut digunakan dalam beberapa skenario pengujian untuk mengetahui pengaruhnya terhadap performa model dalam memprediksi risiko GERD.

Selain *n\_estimators* dan *max\_depth*, parameter lain seperti *bootstrap* *sampling rate* juga terbukti memengaruhi performa model. Penelitian menunjukkan bahwa jumlah sampel yang diambil dalam proses *bootstrap* (*bootstrap rate*) dapat disesuaikan untuk mengoptimalkan keseimbangan antara bias dan varians, terutama pada dataset dengan karakteristik tertentu (Iwaniuk dkk., 2025). Namun, dalam praktiknya, penggunaan nilai default sering kali sudah cukup untuk menghasilkan performa yang baik (Haseeb dkk., 2024).

Selain itu, implementasi model juga menggunakan parameter *class\_weight='balanced'* untuk membantu menangani distribusi kelas yang tidak seimbang pada data.

Model yang telah dibangun kemudian digunakan untuk melakukan prediksi terhadap data baru. Proses prediksi dilakukan dengan memasukkan data *input* ke dalam model yang telah dilatih, kemudian menghasilkan *output* berupa nilai klasifikasi risiko GERD. *Output* prediksi tersebut dapat digunakan sebagai informasi awal untuk mengetahui risiko GERD berdasarkan pola makan yang dimiliki.

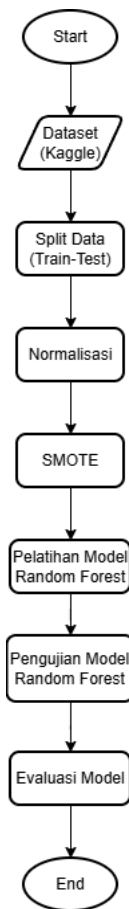
## BAB IV

### UJI COBA DAN PEMBAHASAN

#### 4.1 Skenario Pengujian

Rangkaian skenario uji coba pada penelitian ini disusun dengan tujuan untuk mengukur performa model Random Forest dalam memprediksi risiko GERD melalui variasi nilai parameter yang diterapkan.

Proses pengujian diawali dengan pembagian *dataset* menjadi data latih dan data uji, dilanjutkan dengan normalisasi menggunakan *MinMaxScaler*. Selanjutnya, dilakukan penyeimbangan data menggunakan SMOTE pada data latih sebelum pelatihan model *Random Forest*. Model yang telah dilatih kemudian diuji menggunakan data uji dan dievaluasi menggunakan *confusion matrix* dengan metrik *accuracy*, *precision*, *recall*, dan *F1-Score*. Tahapan ini divisualisasikan dalam bentuk *flowchart* pada Gambar 4.1.



Gambar 4.1 Diagram Alir Proses Pengujian

Penelitian ini memfokuskan skenario pengujian pada variasi dua parameter utama algoritma *Random Forest*, yaitu jumlah pohon ( $n\_estimators$ ) dan kedalaman pohon ( $max\_depth$ ). Pemilihan kedua parameter ini merujuk pada penelitian terdahulu yang menunjukkan pengaruh signifikan keduanya terhadap performa model, khususnya dalam prediksi berbasis data pola makan (Rosidah dkk., 2025).

Parameter  $n\_estimators$  merepresentasikan jumlah pohon keputusan yang dibangun dalam model, di mana jumlah pohon yang lebih banyak umumnya dapat meningkatkan stabilitas dan konsistensi prediksi. Sementara

itu, *max\_depth* mengatur batas kedalaman maksimum setiap pohon keputusan yang memengaruhi kemampuan model dalam menangkap pola data. Penentuan nilai *max\_depth* yang terlalu rendah berpotensi memicu kondisi *underfitting*, di mana model gagal menangkap pola data secara optimal. Sebaliknya, nilai yang terlalu tinggi atau tanpa batasan (*None*) dapat meningkatkan kompleksitas model dan memperbesar kemungkinan terjadinya *overfitting*. Oleh karena itu, pemilihan kombinasi nilai pada kedua parameter ini menjadi krusial dalam menghasilkan performa model yang optimal.

Seluruh skenario pengujian menerapkan proporsi pembagian *data training* dan *data testing* yang konsisten, yaitu 80:20, untuk memastikan validitas perbandingan antar konfigurasi parameter. Hal ini bertujuan agar perbedaan performa yang dihasilkan murni dipengaruhi oleh variasi parameter model, bukan oleh perubahan distribusi data.

Masing-masing parameter memiliki tiga variasi nilai, yaitu *n\_estimators* (50, 100, 150) dan *max\_depth* (10, 15, dan *None*). Kombinasi dari kedua parameter tersebut menghasilkan sembilan skenario pengujian. Setiap skenario merepresentasikan satu konfigurasi parameter yang berbeda, sehingga memungkinkan analisis pengaruh perubahan parameter terhadap performa model.

Rincian skenario pengujian disajikan pada Tabel 4.1.

Tabel 4.1 Variasi Parameter Pengujian

| Kode | <i>n_estimators</i> | <i>max_depth</i> |
|------|---------------------|------------------|
| A1   | 50                  | 10               |
| A2   | 50                  | 15               |
| A3   | 50                  | <i>None</i>      |
| A4   | 100                 | 10               |
| A5   | 100                 | 15               |
| A6   | 100                 | <i>None</i>      |

|    |     |      |
|----|-----|------|
| A7 | 150 | 10   |
| A8 | 150 | 15   |
| A9 | 150 | None |

Evaluasi model dilakukan dengan menggunakan *confusion matrix* yang merepresentasikan empat kategori hasil klasifikasi, yakni *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN). Nilai-nilai tersebut digunakan untuk menghitung metrik evaluasi berupa *accuracy*, *precision*, *recall*, dan *F1-Score*.

*Confusion matrix* terdiri dari empat komponen utama:

- *True Positive (TP)*: jumlah data GERD yang diprediksi benar sebagai GERD
- *True Negative (TN)*: jumlah data tidak GERD yang diprediksi benar sebagai tidak GERD
- *False Positive (FP)*: jumlah data tidak GERD yang diprediksi sebagai GERD
- *False Negative (FN)*: jumlah data GERD yang diprediksi sebagai tidak GERD

*Accuracy* mengukur tingkat ketepatan model dalam memprediksi seluruh data uji, dihitung menggunakan persamaan (4.1):

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (4.1)$$

*Precision* mengukur ketepatan model dalam memprediksi kelas positif (GERD), dihitung menggunakan persamaan (4.2):

$$Precision = \frac{TP}{TP + FP} \quad (4.2)$$

*Recall* mengukur kemampuan model dalam mendeteksi seluruh data yang benar-benar termasuk kelas GERD, dihitung menggunakan persamaan (4.3):

$$Recall = \frac{TP}{TP + FN} \quad (4.3)$$

*F1-Score* digunakan untuk melihat keseimbangan performa model antara *precision* dan *recall*, dihitung menggunakan persamaan (4.4):

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (4.4)$$

Berdasarkan hasil perhitungan tersebut, dilakukan analisis terhadap perubahan performa model pada setiap kombinasi parameter yang digunakan.

## 4.2 Hasil Pengujian

### 4.2.1 Hasil Pengujian Skenario A1

Skenario A1 menggunakan parameter  $n\_estimators = 50$  dan  $max\_depth = 10$ . Hasil pengujian disajikan pada Tabel 4.2.

Tabel 4.2 Hasil Pengujian Skenario A1

| <b>TN</b> | <b>FP</b> | <b>FN</b> | <b>TP</b> |
|-----------|-----------|-----------|-----------|
| 733       | 70        | 81        | 68        |

Model menghasilkan *accuracy* sebesar 0,8414, *precision* sebesar 0,4928, *recall* sebesar 0,4564, dan *F1-Score* sebesar 0,4739. Nilai *precision* dan *recall* yang masih rendah mengindikasikan bahwa model

belum optimal dalam mengenali pola data GERD. Hal ini disebabkan oleh keterbatasan kedalaman pohon yang membuat model belum mampu menangkap hubungan antar fitur secara kompleks.

#### 4.2.2 Hasil Pengujian Skenario A2

Skenario A2 menggunakan parameter  $n\_estimators = 50$  dan  $max\_depth = 15$ . Hasil pengujian disajikan pada Tabel 4.3.

Tabel 4.3 Hasil Pengujian Skenario A2

| <b>TN</b> | <b>FP</b> | <b>FN</b> | <b>TP</b> |
|-----------|-----------|-----------|-----------|
| 769       | 34        | 72        | 77        |

Model menghasilkan *accuracy* sebesar 0,8887, *precision* sebesar 0,6937, *recall* sebesar 0,5168, dan *F1-Score* sebesar 0,5923. Terjadi peningkatan performa dibandingkan skenario A1 setelah  $max\_depth$  ditingkatkan dari 10 menjadi 15. Kedalaman pohon yang lebih besar memungkinkan model mempelajari pola data dengan lebih baik.

#### 4.2.3 Hasil Pengujian Skenario A3

Skenario A3 menggunakan parameter  $n\_estimators = 50$  dan  $max\_depth = None$ . Hasil pengujian disajikan pada Tabel 4.4.

Tabel 4.4 Hasil Pengujian Skenario A3

| <b>TN</b> | <b>FP</b> | <b>FN</b> | <b>TP</b> |
|-----------|-----------|-----------|-----------|
| 776       | 27        | 72        | 77        |

Konfigurasi pada skenario A3 menghasilkan unjuk kerja tertinggi dengan nilai *accuracy* 0,8960, *precision* 0,7404, *recall* 0,5168, dan *F1-Score* 0,6087. Penerapan  $max\_depth = None$  memberikan keleluasaan yang lebih besar dalam

pembentukan pohon keputusan, sehingga model dapat menangkap kompleksitas pola data dengan lebih baik.

#### 4.2.4 Hasil Pengujian Skenario A4

Skenario A4 menggunakan parameter  $n\_estimators = 100$  dan  $max\_depth = 10$ . Hasil pengujian disajikan pada Tabel 4.5.

Tabel 4.5 Hasil Pengujian Skenario A4

| <b>TN</b> | <b>FP</b> | <b>FN</b> | <b>TP</b> |
|-----------|-----------|-----------|-----------|
| 741       | 62        | 79        | 70        |

Model menghasilkan *accuracy* sebesar 0,8519, *precision* sebesar 0,5303, *recall* sebesar 0,4698, dan *F1-Score* sebesar 0,4982. Peningkatan jumlah pohon dari 50 menjadi 100 pada  $max\_depth = 10$  belum memberikan peningkatan performa yang signifikan, menunjukkan bahwa pembatasan kedalaman pohon masih menjadi kendala utama.

#### 4.2.5 Hasil Pengujian Skenario A5

Skenario A5 menggunakan parameter  $n\_estimators = 100$  dan  $max\_depth = 15$ . Hasil pengujian disajikan pada Tabel 4.6.

Tabel 4.6 Hasil Pengujian Skenario A5

| <b>TN</b> | <b>FP</b> | <b>FN</b> | <b>TP</b> |
|-----------|-----------|-----------|-----------|
| 771       | 32        | 74        | 75        |

Model menghasilkan *accuracy* sebesar 0,8887, *precision* sebesar 0,7009, *recall* sebesar 0,5034, dan *F1-Score* sebesar 0,5859. Peningkatan performa terjadi setelah  $max\_depth$  ditingkatkan menjadi 15, yang memungkinkan model melakukan klasifikasi lebih baik pada kedua kelas.



#### 4.2.6 Hasil Pengujian Skenario A6

Skenario A6 menggunakan parameter  $n\_estimators = 100$  dan  $max\_depth = None$ . Hasil pengujian disajikan pada Tabel 4.7.

Tabel 4.7 Hasil Pengujian Skenario A6

| <b>TN</b> | <b>FP</b> | <b>FN</b> | <b>TP</b> |
|-----------|-----------|-----------|-----------|
| 773       | 30        | 72        | 77        |

Model menghasilkan *accuracy* sebesar 0,8929, *precision* sebesar 0,7196, *recall* sebesar 0,5168, dan *F1-Score* sebesar 0,6016. Penggunaan  $max\_depth = None$  kembali memberikan peningkatan performa dengan keseimbangan yang lebih baik antara *precision* dan *recall*.

#### 4.2.7 Hasil Pengujian Skenario A7

Skenario A7 menggunakan parameter  $n\_estimators = 150$  dan  $max\_depth = 10$ . Hasil pengujian disajikan pada Tabel 4.8.

Tabel 4.8 Hasil Pengujian Skenario A7

| <b>TN</b> | <b>FP</b> | <b>FN</b> | <b>TP</b> |
|-----------|-----------|-----------|-----------|
| 741       | 62        | 79        | 70        |

Model menghasilkan *accuracy* sebesar 0,8519, *precision* sebesar 0,5303, *recall* sebesar 0,4698, dan *F1-Score* sebesar 0,4982. Penambahan jumlah pohon menjadi 150 pada  $max\_depth = 10$  belum memberikan peningkatan dibandingkan skenario A4, menegaskan bahwa kedalaman pohon lebih berpengaruh dibandingkan jumlah pohon.

#### 4.2.8 Hasil Pengujian Skenario A8

Skenario A8 menggunakan parameter  $n\_estimators = 150$  dan  $max\_depth = 15$ . Hasil pengujian disajikan pada Tabel 4.9.

Tabel 4.9 Hasil Pengujian Skenario A8

| <b>TN</b> | <b>FP</b> | <b>FN</b> | <b>TP</b> |
|-----------|-----------|-----------|-----------|
| 771       | 32        | 74        | 75        |

Model menghasilkan *accuracy* sebesar 0,8887, *precision* sebesar 0,7009, *recall* sebesar 0,5034, dan *F1-Score* sebesar 0,5859. Penambahan jumlah pohon dari 100 menjadi 150 tidak memberikan perubahan performa yang berarti, menunjukkan bahwa model telah mencapai titik optimal pada konfigurasi sebelumnya.

#### 4.2.9 Hasil Pengujian Skenario A9

Skenario A9 menggunakan parameter  $n\_estimators = 150$  dan  $max\_depth = None$ . Hasil pengujian disajikan pada Tabel 4.10.

Tabel 4.10 Hasil Pengujian Skenario A9

| <b>TN</b> | <b>FP</b> | <b>FN</b> | <b>TP</b> |
|-----------|-----------|-----------|-----------|
| 775       | 28        | 73        | 76        |

Model menghasilkan *accuracy* sebesar 0,8939, *precision* sebesar 0,7308, *recall* sebesar 0,5101, dan *F1-Score* sebesar 0,6008. Performa mendekati skenario terbaik, namun peningkatan  $n\_estimators$  tidak memberikan peningkatan besar dibandingkan skenario A3, menunjukkan bahwa penambahan pohon yang terlalu besar tidak selalu optimal.

### 4.3 Analisis dan Pembahasan

Pada bagian ini dilakukan analisis dan pembahasan terhadap hasil pengujian model *Random Forest* yang telah diperoleh. Analisis bertujuan untuk mengetahui pengaruh variasi parameter terhadap performa model serta menentukan konfigurasi parameter terbaik.

Pembahasan dilakukan berdasarkan metrik *accuracy*, *precision*, *recall*, dan *F1-Score*, serta hasil klasifikasi pada *confusion matrix*. Selain itu, dilakukan pula analisis pengaruh parameter *Random Forest* dan *Feature Importance*.

#### 4.3.1 Analisis Performa Model pada Berbagai Skenario Pengujian

Evaluasi model dilakukan menggunakan *confusion matrix* berdasarkan hasil prediksi pada data uji. Melalui *confusion matrix*, dapat diketahui jumlah prediksi benar dan salah pada masing-masing kelas.

Berdasarkan hasil *confusion matrix*, dihitung beberapa metrik performa:

- *Accuracy*: mengukur ketepatan model dalam memprediksi keseluruhan data
- *Precision*: mengukur ketepatan model dalam memprediksi kelas positif (GERD)
- *Recall*: mengukur kemampuan model mendeteksi seluruh data kelas GERD
- *F1-Score*: rata-rata harmonis antara *precision* dan *recall*

Rekapitulasi hasil evaluasi performa seluruh skenario ditampilkan pada

Tabel 4.11 dan Tabel 4.12.

Tabel 4.11 Rekapitulasi Performa Seluruh Skenario

| <b>Kode</b> | <b><i>n_estimators</i></b> | <b><i>max_depth</i></b> | <b><i>Accuracy</i></b> | <b><i>Precision</i></b> | <b><i>Recall</i></b> | <b><i>F1-Score</i></b> |
|-------------|----------------------------|-------------------------|------------------------|-------------------------|----------------------|------------------------|
| A1          | 50                         | 10                      | 0.8414                 | 0.4928                  | 0.4564               | 0.4739                 |
| A2          | 50                         | 15                      | 0.8887                 | 0.6937                  | 0.5168               | 0.5923                 |
| A3          | 50                         | <i>None</i>             | 0.8960                 | 0.7404                  | 0.5168               | 0.6087                 |
| A4          | 100                        | 10                      | 0.8519                 | 0.5303                  | 0.4698               | 0.4982                 |
| A5          | 100                        | 15                      | 0.8887                 | 0.7009                  | 0.5034               | 0.5859                 |
| A6          | 100                        | <i>None</i>             | 0.8929                 | 0.7196                  | 0.5168               | 0.6016                 |
| A7          | 150                        | 10                      | 0.8519                 | 0.5303                  | 0.4698               | 0.4982                 |
| A8          | 150                        | 15                      | 0.8887                 | 0.7009                  | 0.5034               | 0.5859                 |
| A9          | 150                        | <i>None</i>             | 0.8939                 | 0.7308                  | 0.5101               | 0.6008                 |

Tabel 4.12 Perbandingan Performa Model Berdasarkan *F1-Score*

| <b>Kode</b> | <b><i>n_estimators</i></b> | <b><i>max_depth</i></b> | <b><i>Accuracy</i></b> | <b><i>Precision</i></b> | <b><i>Recall</i></b> | <b><i>F1-Score</i></b> |
|-------------|----------------------------|-------------------------|------------------------|-------------------------|----------------------|------------------------|
| A3          | 50                         | <i>None</i>             | 0.8960                 | 0.7404                  | 0.5168               | 0.6087                 |
| A6          | 100                        | <i>None</i>             | 0.8929                 | 0.7196                  | 0.5168               | 0.6016                 |
| A9          | 150                        | <i>None</i>             | 0.8939                 | 0.7308                  | 0.5101               | 0.6008                 |
| A2          | 50                         | 15                      | 0.8887                 | 0.6937                  | 0.5168               | 0.5923                 |
| A8          | 150                        | 15                      | 0.8887                 | 0.7009                  | 0.5034               | 0.5859                 |
| A5          | 100                        | 15                      | 0.8887                 | 0.7009                  | 0.5034               | 0.5859                 |
| A4          | 100                        | 10                      | 0.8519                 | 0.5303                  | 0.4698               | 0.4982                 |
| A7          | 150                        | 10                      | 0.8519                 | 0.5303                  | 0.4698               | 0.4982                 |
| A1          | 50                         | 10                      | 0.8414                 | 0.4928                  | 0.4564               | 0.4739                 |

Berdasarkan hasil pengujian pada sembilan skenario, setiap kombinasi parameter *n\_estimators* dan *max\_depth* menghasilkan performa yang berbeda. Perbedaan tersebut tercermin dari nilai *accuracy*, *precision*, *recall*, dan *F1-Score* yang diperoleh.

Secara umum, seluruh skenario menghasilkan *accuracy* yang cukup tinggi pada rentang 84% hingga 89%. Namun, *accuracy* yang tinggi tidak selalu mencerminkan performa terbaik karena kondisi *dataset* yang tidak seimbang. Oleh

karena itu, evaluasi model tidak hanya berfokus pada *accuracy*, tetapi juga mempertimbangkan *precision*, *recall*, dan *F1-Score*.

Nilai *precision* cenderung meningkat seiring bertambahnya *max\_depth*, menunjukkan model semakin baik dalam mengidentifikasi data yang benar-benar termasuk kelas GERD. Sementara itu, *recall* pada sebagian besar skenario masih relatif rendah pada kisaran 0,45 hingga 0,52, mengindikasikan masih terdapat data GERD yang tidak terdeteksi (*false negative*).

Berdasarkan *F1-Score*, performa terbaik diperoleh pada skenario A3 dengan parameter *n\_estimators* = 50 dan *max\_depth* = *None*, yaitu sebesar 0,6087. Nilai ini merupakan yang tertinggi karena mampu memberikan keseimbangan lebih baik antara *precision* dan *recall*. Oleh karena itu, skenario A3 dipilih sebagai model terbaik.

Secara keseluruhan, hasil pengujian menunjukkan bahwa algoritma *Random Forest* mampu digunakan untuk memprediksi risiko GERD berdasarkan data pola makan. Variasi parameter memberikan pengaruh terhadap performa, terutama *max\_depth* yang terbukti lebih berpengaruh dibandingkan *n\_estimators*.

Studi oleh Alsabhan dan Alfadhly (2025) mengimplementasikan *Min-Max Scaling* pada *Random Forest* dan berhasil memperoleh akurasi sebesar 0,8873 ketika diuji pada *dataset* medis. Penelitian yang dilakukan oleh Suryanegara dkk. (2021) juga memperkuat bukti bahwa penerapan normalisasi *Min-Max* pada *Random Forest* mampu menghasilkan tingkat akurasi yang lebih unggul dibandingkan model tanpa normalisasi. Temuan ini mendukung bahwa kombinasi

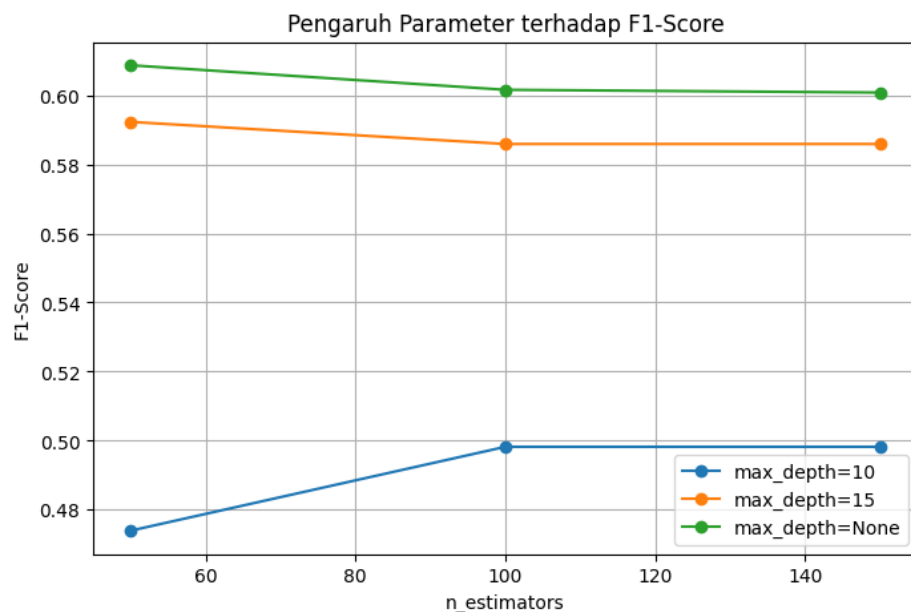
normalisasi data dan algoritma *ensemble* seperti *Random Forest* dapat memberikan performa yang konsisten (Priyambudi & Nugroho, 2024).

### 4.3.2 Analisis Pengaruh Parameter *Random Forest* terhadap Performa Model

#### Model

Berdasarkan hasil pengujian pada sembilan skenario, performa model *Random Forest* menunjukkan perbedaan pada setiap kombinasi parameter yang digunakan. Evaluasi dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-Score*.

Perbandingan pengaruh parameter terhadap *F1-Score* ditampilkan pada Gambar 4.2.



Gambar 4.2 Perbandingan Pengaruh Parameter Terhadap *F1-Score*

Dari visualisasi pada Gambar 4.2, tampak bahwa parameter *max\_depth* memiliki kontribusi yang lebih dominan terhadap performa model dibandingkan *n\_estimators*. Indikasi ini terlihat dari fluktuasi nilai *F1-Score*

*Score* yang cukup signifikan ketika *max\_depth* diubah, sedangkan variasi pada *n\_estimators* hanya menghasilkan perbedaan performa yang tergolong kecil.

Performa terbaik diperoleh ketika *max\_depth* = *None*, menunjukkan bahwa model memiliki kemampuan lebih baik dalam mempelajari pola data ketika kedalaman pohon tidak dibatasi. Sebaliknya, pembatasan kedalaman pohon menyebabkan model kurang optimal dalam mengenali pola data GERD.

Hasil pengujian juga menunjukkan bahwa *n\_estimators* = 50 sudah cukup untuk menghasilkan performa optimal. Penambahan jumlah pohon tidak memberikan peningkatan performa signifikan, namun meningkatkan kompleksitas dan waktu komputasi.

Temuan ini sejalan dengan penelitian sebelumnya yang menunjukkan bahwa kedalaman pohon yang terlalu dangkal menyebabkan *underfitting*, sementara kedalaman yang terlalu besar atau tidak dibatasi dapat meningkatkan risiko *overfitting*.

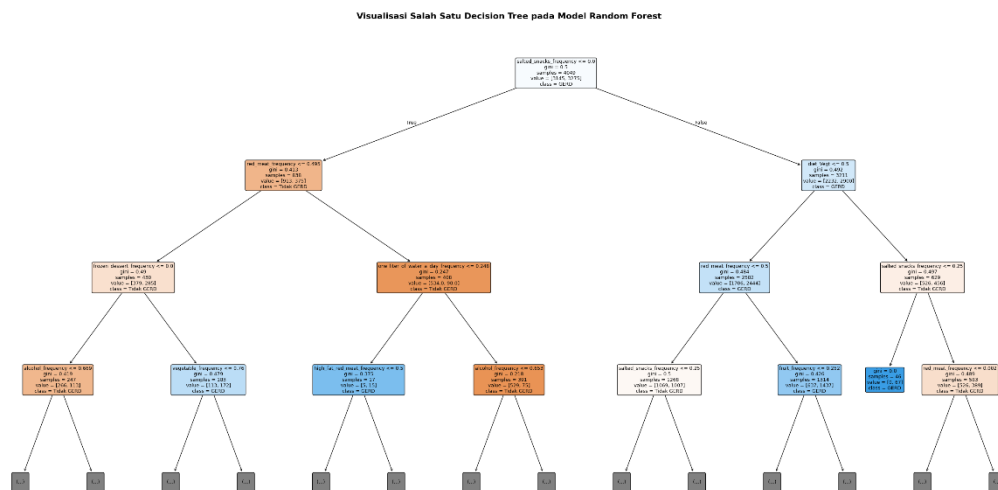
Optimasi parameter seperti *max\_depth* dan *min\_samples\_split* memiliki peran penting dalam meningkatkan akurasi model, di mana kombinasi hyperparameter yang tepat dapat meningkatkan akurasi secara signifikan dibandingkan parameter default.

#### 4.3.3 Visualisasi *Decision Tree Random Forest*

Untuk memberikan gambaran proses pengambilan keputusan pada model *Random Forest*, ditampilkan salah satu *decision tree* yang terbentuk pada model terbaik (skenario A3). Visualisasi ini bertujuan menunjukkan bagaimana model melakukan klasifikasi berdasarkan atribut pola makan hingga menghasilkan

prediksi risiko GERD. Perlu diketahui bahwa *Random Forest* terdiri atas banyak *decision tree*, sehingga gambar yang ditampilkan hanya merupakan salah satu pohon keputusan sebagai representasi.

Visualisasi pada Gambar 4.3 merupakan salah satu *decision tree* dari model terbaik dengan parameter  $n\_estimators = 50$  dan  $max\_depth = None$ . Meskipun model dibangun menggunakan banyak *decision tree*, hanya ditampilkan satu pohon sebagai representasi. Agar struktur pohon lebih mudah dibaca, visualisasi dibatasi hingga kedalaman 3 tanpa memengaruhi proses pelatihan maupun hasil prediksi.



Gambar 4.3 Contoh Decision Tree dari Model *Random Forest*

Berdasarkan Gambar 4.3, proses klasifikasi dimulai dari atribut yang memiliki kemampuan pemisahan data paling baik. Data kemudian dipisahkan berdasarkan nilai ambang (*threshold*) pada setiap atribut hingga mencapai *leaf node* yang menghasilkan klasifikasi akhir. Meskipun visualisasi hanya

menampilkan satu *decision tree*, prediksi akhir model diperoleh melalui mekanisme *majority voting* dari seluruh *decision tree* yang dibangun.

#### 4.3.4 Analisis *Feature Importance*

Setelah diperoleh model terbaik, dilakukan analisis *Feature Importance* untuk mengetahui kontribusi masing-masing fitur terhadap prediksi risiko GERD menggunakan atribut *feature\_importances\_* pada algoritma *Random Forest*.

Nilai *Feature Importance* menunjukkan seberapa besar pengaruh suatu fitur dalam proses pembentukan pohon keputusan. Semakin tinggi nilai *importance*, semakin besar kontribusinya terhadap hasil klasifikasi. Hasil perhitungan *Feature Importance* ditampilkan pada Tabel 4.13.

Tabel 4.13 Nilai *Feature Importance Random Forest*

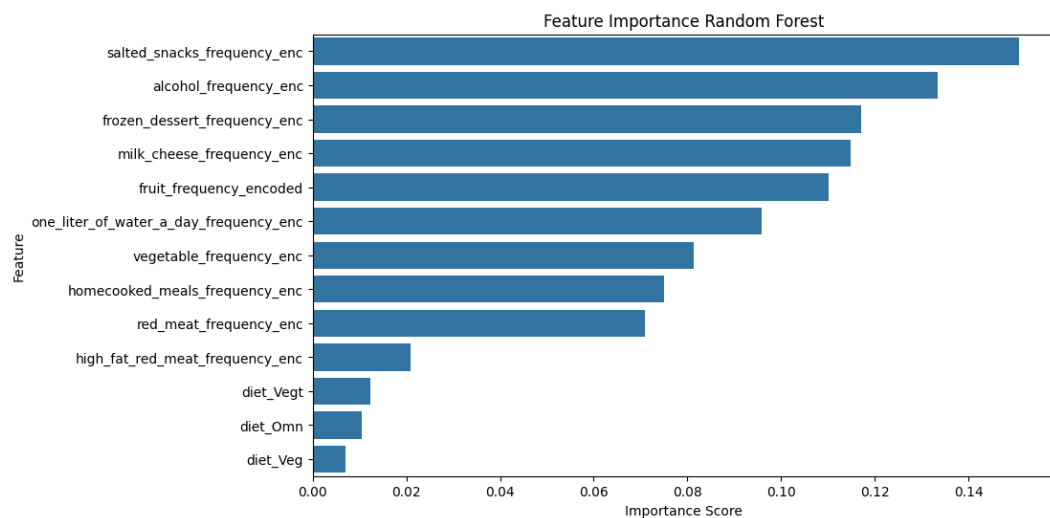
|    | <i>Feature</i>                                | <i>Importance</i> |
|----|-----------------------------------------------|-------------------|
| 11 | <i>salted snacks frequency enc</i>            | 0.150794          |
| 7  | <i>alcohol frequency enc</i>                  | 0.133444          |
| 8  | <i>frozen dessert frequency enc</i>           | 0.117021          |
| 9  | <i>milk cheese frequency enc</i>              | 0.114960          |
| 3  | <i>fruit frequency encoded</i>                | 0.110202          |
| 10 | <i>one liter of water a day frequency enc</i> | 0.095830          |
| 6  | <i>vegetable frequency enc</i>                | 0.081360          |
| 5  | <i>homecooked meals frequency enc</i>         | 0.075086          |
| 12 | <i>red meat frequency enc</i>                 | 0.070840          |
| 4  | <i>high fat red meat frequency enc</i>        | 0.020818          |
| 2  | <i>diet Vegt</i>                              | 0.012218          |
| 0  | <i>diet Omn</i>                               | 0.010439          |
| 1  | <i>diet Veg</i>                               | 0.006990          |

Hasil analisis menunjukkan bahwa frekuensi konsumsi makanan asin dan alkohol merupakan fitur paling berpengaruh dalam prediksi risiko GERD. Temuan ini sejalan dengan penelitian sebelumnya yang menunjukkan hubungan erat antara

pola makan dan kejadian GERD. Alaya dkk. (2025) menemukan bahwa semakin buruk pola makan seseorang, semakin tinggi risiko GERD. Konsumsi makanan asin dan alkohol diketahui memicu peningkatan produksi asam lambung dan iritasi saluran pencernaan.

Dalam konteks yang lebih luas, Lee dkk. (2024) menggunakan *Random Forest* pada data populasi dan menemukan hubungan kuat antara GERD dan kejadian kelahiran prematur dengan nilai kontribusi SHAP sebesar 0,27. Hal ini menunjukkan bahwa dampak GERD tidak hanya terbatas pada gangguan pencernaan, sehingga deteksi dini menjadi sangat penting.

Visualisasi *Feature Importance* ditampilkan pada Gambar 4.4 untuk mempermudah interpretasi kontribusi masing-masing fitur.



Gambar 4.4 Diagram *Feature Importance Random Forest*

Berdasarkan analisis, fitur paling berpengaruh adalah frekuensi konsumsi makanan asin (*salted\_snacks\_frequency\_encoded*) dengan nilai importance 0,1508.

Fitur lain dengan pengaruh cukup tinggi adalah konsumsi alkohol, makanan pencuci mulut beku, produk susu dan keju, serta konsumsi buah.

Hasil ini menunjukkan bahwa pola konsumsi makanan tertentu memiliki kontribusi besar terhadap prediksi risiko GERD. Konsumsi makanan asin dan alkohol yang tinggi memicu peningkatan produksi asam lambung dan iritasi saluran pencernaan. Sementara itu, fitur jenis diet memiliki nilai *importance* relatif rendah, menunjukkan bahwa frekuensi konsumsi makanan lebih berpengaruh dibandingkan kategori jenis diet.

#### **4.3.5 Pengujian Tambahan Menggunakan *Random UnderSampling* dari Model Terbaik**

Dari hasil evaluasi, model terbaik yang diperoleh pada skenario A3 menunjukkan nilai *F1-Score* 0,6087 dengan tingkat *recall* yang masih tergolong rendah, yaitu 0,5168. Angka ini mengindikasikan bahwa model masih belum mampu mendeteksi seluruh kasus GERD secara sempurna, ditandai dengan masih tingginya jumlah *false negative*. Selain itu, *confusion matrix* menunjukkan *true negative* jauh lebih besar dibandingkan *true positive*, yang dipengaruhi oleh ketidakseimbangan data. Oleh karena itu, dilakukan pengujian tambahan menggunakan *Random UnderSampling* untuk meningkatkan kemampuan model dalam mendeteksi kelas GERD.

Perbandingan distribusi data sebelum dan sesudah *Random UnderSampling* disajikan pada Tabel 4.14.

Tabel 4.14 Distribusi Data Sebelum dan Sesudah *Random UnderSampling*

| Kelas<br>( <i>Reflu</i> ) | Keterangan     | Jumlah Data<br>Sebelum<br><i>UnderSampling</i> | Jumlah Data Setelah<br><i>UnderSampling</i> |
|---------------------------|----------------|------------------------------------------------|---------------------------------------------|
| 0                         | Tidak GERD     | 3.210                                          | 744                                         |
| 1                         | Mengalami GERD | 3.210                                          | 744                                         |
| <b>Total</b>              |                | <b>6.420</b>                                   | <b>1.488</b>                                |

Berdasarkan hasil *Random UnderSampling*, data pada masing-masing kelas berhasil diseimbangkan menjadi 744 data untuk kelas tidak GERD dan 744 data untuk kelas GERD. Proses ini bertujuan agar model dapat mempelajari pola data dari kedua kelas secara lebih seimbang.

Hasil pengujian sebelum dan sesudah *Random UnderSampling* disajikan pada Tabel 4.15.

Tabel 4.15 Perbandingan Performa Model Sebelum dan Sesudah *Random UnderSampling*

| <i>Metric</i>    | Sebelum | Sesudah |
|------------------|---------|---------|
| <i>Accuracy</i>  | 0.8960  | 0.7583  |
| <i>Precision</i> | 0.7404  | 0.7984  |
| <i>Recall</i>    | 0.5168  | 0.6912  |
| <i>F1-Score</i>  | 0.6087  | 0.7410  |

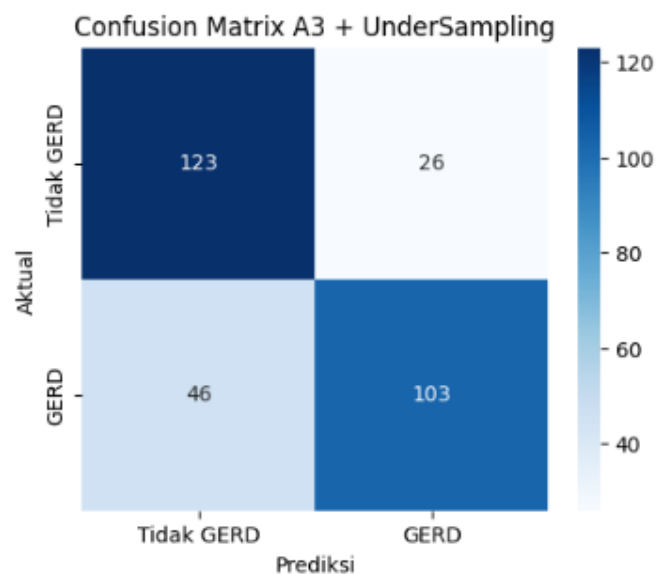
Berdasarkan hasil uji coba tambahan, terjadi peningkatan yang cukup berarti pada metrik *recall* dan *F1-Score*, walaupun nilai *accuracy* justru menunjukkan penurunan. Kenaikan *recall* dari 0,5168 menjadi 0,6912 mencerminkan bahwa model kini lebih sensitif dalam mengidentifikasi kasus GERD, yang berarti berkurangnya jumlah *false negative*. Hal ini sejalan dengan temuan Wisesty dkk. (2025) bahwa penanganan ketidakseimbangan data menggunakan SMOTE mampu meningkatkan *F1-Score* secara signifikan. Pada kasus medis seperti prediksi GERD, peningkatan deteksi kelas positif sangat

penting karena *false negative* dapat menyebabkan pasien tidak mendapat peringatan dini.

Meskipun *accuracy* menurun dari 0,8960 menjadi 0,7583, penurunan ini dapat diterima karena model menjadi lebih seimbang. Pada data tidak seimbang, *accuracy* seringkali menyesatkan. *F1-Score* yang meningkat dari 0,6087 menjadi 0,7410 menunjukkan keseimbangan lebih baik antara *precision* dan *recall*.

Hasil ini didukung oleh Wang dkk. (2026) yang menunjukkan bahwa *Random Forest* dengan penanganan data tepat mampu mencapai performa baik (AUC 0,829) dalam mengidentifikasi *Reflux Esophagitis*. Hal ini mengindikasikan *Random Forest* robust dan andal untuk prediksi risiko GERD dengan teknik penanganan data yang sesuai.

Hasil pengujian divisualisasikan menggunakan *confusion matrix* pada Gambar 4.5.



Gambar 4.5 *Confusion Matrix* Skenario A3 dengan *UnderSampling*

Berdasarkan *confusion matrix*, jumlah data GERD yang terdeteksi meningkat dibandingkan sebelumnya, dan *false negative* menurun. Meskipun *accuracy* menurun, peningkatan *recall* dan *F1-Score* menunjukkan model lebih seimbang dalam klasifikasi kedua kelas. Pada kasus kesehatan, peningkatan deteksi kelas positif sangat penting untuk mengurangi risiko kesalahan prediksi.

Hasil pengujian tambahan dengan *Random UnderSampling* mampu meningkatkan performa model, terutama pada *recall* dan *F1-Score*, menunjukkan bahwa penanganan ketidakseimbangan data berpengaruh terhadap kemampuan model dalam mendeteksi risiko GERD secara lebih optimal.

Berdasarkan hasil penelitian yang telah dilakukan, terbukti bahwa faktor pola makan berpengaruh terhadap risiko GERD. Temuan dari analisis *Feature Importance* mengungkapkan bahwa konsumsi makanan asin (*salted snacks*) merupakan faktor dengan kontribusi tertinggi dalam proses prediksi, diikuti oleh konsumsi alkohol, makanan penutup beku (*frozen dessert*), produk susu dan keju, serta buah. Hal ini menegaskan bahwa kebiasaan konsumsi makanan dan minuman memiliki peran signifikan terhadap risiko GERD, sehingga pengendalian pola makan menjadi salah satu langkah penting dalam upaya pemeliharaan kesehatan.

Temuan ini sejalan dengan firman Allah SWT dalam QS. Al-A'raf ayat 31:

يَا بَنِي آدَمَ خُذُوا زِينَتَكُمْ عِنْدَ كُلِّ مَسْجِدٍ وَكُلُوا وَاشْرَبُوا وَلَا تُسْرِفُوا إِنَّهُ لَا يُحِبُّ الْمُسْرِفِينَ ٣١

"Wahai anak cucu Adam! Pakailah pakaianmu yang bagus setiap (memasuki) masjid, makan dan minumlah, tetapi jangan berlebihan! Sesungguhnya, Allah tidak menyukai orang yang berlebih-lebihan." (QS. Al-A'raf : 31)

Menurut Tafsir Ibnu Katsir, ayat ini tidak sekadar memerintahkan manusia untuk makan dan minum, tetapi juga melarang sikap berlebihan (*israf*) dalam mengonsumsi sesuatu. Larangan ini bertujuan agar manusia senantiasa menjaga keseimbangan dalam memenuhi kebutuhan tubuh serta menghindari kebiasaan yang dapat mendatangkan mudarat bagi kesehatan. Bahkan sebagian ulama yang dikutip dalam tafsir tersebut menyatakan bahwa prinsip kesehatan secara ringkas tercermin dalam perintah untuk makan dan minum secukupnya serta menjauhi perilaku berlebihan. Prinsip ini selaras dengan hasil penelitian yang menunjukkan bahwa beberapa jenis konsumsi makanan dan minuman memiliki kontribusi besar terhadap risiko GERD apabila dikonsumsi secara tidak terkendali.

Penelitian ini juga mengungkap bahwa konsumsi alkohol termasuk salah satu faktor dengan kontribusi tinggi dalam prediksi risiko GERD. Temuan ini mengindikasikan bahwa pola konsumsi minuman tertentu dapat memengaruhi kondisi kesehatan seseorang. Dalam ajaran Islam, konsumsi minuman yang memabukkan (*khamr*) telah dilarang sebagaimana firman Allah SWT dalam QS. Al-Ma'idah ayat 90:

يَا أَيُّهَا الَّذِينَ ءَامَنُوا إِنَّمَا الْخَمْرُ وَالْمَيْسِرُ وَالْأَنْصَابُ وَالْأَزْلَمُ رَجْسٌ مِّنْ عَمَلِ الشَّيْطَانِ فَاجْتَنِبُوهُ لَعَلَّكُمْ تُفْلِحُونَ ٩٠

*"Wahai orang-orang yang beriman! Sesungguhnya minuman keras, berjudi, (berkurban untuk) berhala, dan mengundi nasib dengan anak panah, adalah perbuatan keji dan termasuk perbuatan setan. Maka jauhilah (perbuatan-perbuatan) itu agar kamu beruntung." (QS. Al-Maidah : 90)*

Ayat ini memerintahkan umat Islam untuk menjauhi *khamr* karena tergolong perbuatan keji yang bersumber dari setan. Larangan ini tidak hanya

memiliki hikmah spiritual, tetapi juga mengandung nilai perlindungan terhadap kesehatan manusia. Hasil penelitian yang menunjukkan alkohol sebagai faktor risiko tinggi GERD memperkuat relevansi ajaran Islam dalam mengarahkan manusia untuk menghindari konsumsi minuman yang berpotensi merusak kesehatan dan kehidupan sosial. Dengan demikian, penelitian ini dapat menjadi media edukasi dan dakwah kepada masyarakat tentang pentingnya menerapkan pola makan dan konsumsi yang selaras dengan ajaran Islam sebagai upaya menjaga kesehatan.

Dari perspektif hubungan manusia dengan Allah SWT (*hablum minallah*), menjaga kesehatan merupakan bentuk ketaatan dan rasa syukur atas nikmat tubuh yang telah diberikan. Dengan menerapkan pola makan yang baik, seseorang berupaya melaksanakan amanah yang dititipkan Allah SWT serta mempersiapkan diri agar mampu menjalankan ibadah secara optimal. Hal ini sejalan dengan sabda Rasulullah SAW yang diriwayatkan oleh Imam Bukhari, "*Sesungguhnya tubuhmu mempunyai hak atasmu.*" Hadis ini menegaskan bahwa menjaga kesehatan adalah bagian dari tanggung jawab seorang muslim terhadap amanah Allah SWT. Oleh karena itu, pengaturan pola makan merupakan salah satu bentuk ikhtiar dalam memelihara kesehatan agar manusia dapat menjalankan ibadah kepada Allah secara maksimal.

Dari aspek hubungan manusia dengan sesama (*hablum minannas*), kesehatan yang baik memungkinkan seseorang menjalankan tanggung jawab sosial, pekerjaan, pendidikan, maupun aktivitas kemasyarakatan secara optimal. Upaya deteksi dini risiko GERD melalui pemanfaatan teknologi *machine learning* dapat

membantu masyarakat meningkatkan kesadaran akan pentingnya menjaga kesehatan, sehingga memberikan manfaat bagi diri sendiri maupun lingkungan sosial. Hal ini sejalan dengan firman Allah SWT dalam QS. Al-Ma'idah ayat 2 yang mendorong umat Islam untuk saling tolong-menolong dalam kebajikan dan takwa. Melalui pemanfaatan teknologi, hasil penelitian ini diharapkan menjadi sarana edukasi yang membantu masyarakat mengenali risiko GERD lebih dini, sehingga dapat mengadopsi pola makan yang lebih sehat dan meningkatkan kualitas hidup.

Sementara itu, dari perspektif hubungan manusia dengan lingkungan (*hablum minal alam*), pola hidup sehat mencerminkan sikap bijak dalam memanfaatkan sumber daya yang telah disediakan Allah SWT. Konsumsi makanan yang seimbang, tidak berlebihan, dan sesuai kebutuhan merupakan bentuk pemanfaatan nikmat Allah secara bertanggung jawab. Dengan demikian, penelitian ini tidak hanya memberikan kontribusi di bidang teknologi dan kesehatan, tetapi juga selaras dengan nilai-nilai Islam yang mendorong manusia untuk menjaga amanah, menerapkan pola hidup yang baik, serta mempersiapkan kehidupan yang lebih sehat di masa mendatang. Hal ini sejalan dengan firman Allah SWT dalam QS. Al-A'raf ayat 56 yang melarang manusia berbuat kerusakan di muka bumi, di mana salah satu bentuk penerapannya adalah memanfaatkan sumber daya, termasuk makanan, secara bijaksana dan tidak berlebihan.

## BAB V

### KESIMPULAN DAN SARAN

#### 5.1 Kesimpulan

Dari seluruh rangkaian eksperimen yang telah dijalankan, dapat disimpulkan bahwa algoritma *Random Forest* menunjukkan performa yang memadai untuk memprediksi risiko GERD berbasis data pola makan. Di antara sembilan skenario yang diuji, konfigurasi optimal ditemukan pada skenario A3 dengan nilai  $n\_estimators = 50$  dan  $max\_depth = None$ , yang menghasilkan  $accuracy$  89,60%,  $precision$  74,04%,  $recall$  51,68%, dan  $F1-Score$  60,87%. Pengujian tambahan menggunakan *Random UnderSampling* berhasil meningkatkan kemampuan model dalam mendeteksi kelas minoritas, dengan peningkatan  $recall$  dari 51,68% menjadi 69,12% dan  $F1-Score$  dari 60,87% menjadi 74,10%, meskipun  $accuracy$  mengalami penurunan. Dengan demikian, algoritma *Random Forest* dapat direkomendasikan sebagai pendekatan yang efektif untuk prediksi risiko GERD berbasis data pola makan.

#### 5.2 Saran

Berdasarkan hasil penelitian yang telah dilakukan, terdapat beberapa saran yang dapat dipertimbangkan untuk pengembangan penelitian selanjutnya, yaitu:

1. Mengkaji lebih lanjut efektivitas berbagai strategi penyeimbangan data, di antaranya *SMOTE*, *Random UnderSampling*, dan *hybrid sampling*, guna menemukan pendekatan yang paling optimal untuk mengatasi ketidakseimbangan kelas dalam prediksi risiko GERD.

2. Memperkaya *dataset* dengan variabel-variabel lain yang relevan terhadap risiko GERD, seperti kualitas tidur, tingkat stres, aktivitas fisik, riwayat pengobatan, serta parameter antropometri seperti indeks massa tubuh (*BMI*) yang diketahui memiliki asosiasi kuat dengan GERD.
3. Mengeksplorasi perbandingan performa *Random Forest* dengan pendekatan klasifikasi alternatif seperti *XGBoost*, *Support Vector Machine* (SVM), maupun *Neural Network* guna mengidentifikasi algoritma yang paling efektif untuk prediksi risiko GERD berbasis data diet.

## DAFTAR PUSTAKA

- Afni, N., Salim, A., & Maulana, Y. I. (2021). Penerapan Algoritma Naive Bayes Dalam Memprediksi Penyakit Lambung.
- Alaya, V. N., Permata, A., & Rahmawati, S. (2025). Hubungan Pola Makan dengan Kejadian Gastroesophageal Reflux Disease pada Mahasiswa di ITSK RS dr.Soepraoen Malang. *Journal Syifa Sciences and Clinical Research*, 7(1). <https://doi.org/10.37311/jsscr.v7i1.30681>
- Alsabhan, W., & Alfadhly, A. (2025). Effectiveness of machine learning models in diagnosis of heart disease: A comparative study. *Scientific Reports*, 15(1), 24568. <https://doi.org/10.1038/s41598-025-09423-y>
- Alsahafi, M., Salah, F., Mimish, H., Hejazi, M., Alkhiari, R., Alkhowaiter, S., & Mosli, M. (2024). The prevalence, severity, and risk factors of erosive esophagitis in a Middle Eastern population. *Saudi Journal of Gastroenterology*, 30(6), 376–380. [https://doi.org/10.4103/sjg.sjg\\_91\\_24](https://doi.org/10.4103/sjg.sjg_91_24)
- Bertin, L., Caldart, F., & Savarino, E. V. (2025). Non-pharmacological approaches in gastroesophageal reflux disease: Evidence-based dietary and lifestyle interventions. *Best Practice & Research Clinical Gastroenterology*, 79, 102083. <https://doi.org/10.1016/j.bpg.2025.102083>
- Christian, A., & Yani, A. (2025). Analisis Machine Learning Untuk Prediksi Penyakit Paru-Paru Menggunakan. 7.
- Deng, Y., Xie, Z., You, L., Wang, J., & Wu, M. (2026). Ecological co-burden patterns of GERD and asthma: A multi-method analysis of representative global regions. *Frontiers in Nutrition*, 13, 1774196. <https://doi.org/10.3389/fnut.2026.1774196>
- Dewantika, P., Lubis, A. P., & Putri, P. (2022). Penerapan Teknik Forward Chaining Dan Certainty Factor Untuk Mendeteksi Penyakit Gastroesophageal Reflux Disease (GERD). *Building of Informatics, Technology and Science (BITS)*, 3(4). <https://doi.org/10.47065/bits.v3i4.1439>
- Edric, & Tamba, S. P. (2022). Prediksi Penyakit Gagal Jantung Dengan Menggunakan Random Forest. *Jurnal Sistem Informasi dan Ilmu Komputer Prima(JUSIKOM PRIMA)*, 5(2), 176–181. <https://doi.org/10.34012/jurnalsisteminformasidanilmukomputer.v5i2.2445>
- Fatimah, L. S., Wikarta, D. B., & Lee, Y. Y. (2024). Prevalence and Determinants of Gastroesophageal. 25(1).
- Gde Agung Brahmana Suryanegara, Adiwijaya, & Mahendra Dwifabri Purbolaksono. (2021). Peningkatan Hasil Klasifikasi pada Algoritma Random Forest untuk Deteksi Pasien Penderita Diabetes Menggunakan Metode Normalisasi. *Jurnal RESTI (Rekayasa Sistem dan Teknologi*

Informasi), 5(1), 114–122. <https://doi.org/10.29207/resti.v5i1.2880>

Gwetu, M. V., Tapamo, J.-R., & Viriri, S. (2020). Random Forests with a Steepend Gini-Index Split Function and Feature Coherence Injection. Dalam S. Boumerdassi, É. Renault, & P. Mühlethaler (Ed.), *Machine Learning for Networking* (Vol. 12081, hlm. 255–272). Springer International Publishing. [https://doi.org/10.1007/978-3-030-45778-5\\_17](https://doi.org/10.1007/978-3-030-45778-5_17)

Haseeb, A., Cleland, I., Nugent, C., & McLaughlin, J. (2024). Optimizing Machine Learning for ResourceConstrained Devices: A Comparative Analysis of Preprocessing Techniques and Machine Learning Algorithms. 2024 35th Irish Signals and Systems Conference (ISSC), 1–5. <https://doi.org/10.1109/ISSC61953.2024.10603066>

Iwaniuk, M., Jarosz, M., Borycki, B., Jezierski, B., Cwalina, J., Kaźmierczak, S., & Mańdziuk, J. (2025). The Impact of Bootstrap Sampling Rate on Random Forest Performance in Regression Tasks (Versi 1). arXiv. <https://doi.org/10.48550/ARXIV.2511.13952>

Kesuma, M. (2023). Prediksi Penyakit Liver Menggunakan Algoritma Random Forest. 11(2).

Khan, M., Shah, K., Gill, S. K., Gul, N., J, J. K., Valladares, V., Khan, L. A., & Raza, M. (2024). Dietary Habits and Their Impact on Gastroesophageal Reflux Disease (GERD). *Cureus*. <https://doi.org/10.7759/cureus.65552>

Lee, K.-S., & Kim, E. S. (2025). Genetic Artificial Intelligence in Gastrointestinal Disease: A Systematic Review. *Diagnostics*, 15(17), 2227. <https://doi.org/10.3390/diagnostics15172227>

Lee, K.-S., Kim, E. S., Kim, D., Song, I.-S., & Ahn, K. H. (2021). Association of Gastroesophageal Reflux Disease with Preterm Birth: Machine Learning Analysis. *Journal of Korean Medical Science*, 36(43), e282. <https://doi.org/10.3346/jkms.2021.36.e282>

Lee, K.-S., Song, I.-S., Kim, E. S., Kim, J., Jung, S., Nam, S., & Ahn, K. H. (2024). Machine learning analysis with population data for the associations of preterm birth with temporomandibular disorder and gastrointestinal diseases. *PLOS ONE*, 19(1), e0296329. <https://doi.org/10.1371/journal.pone.0296329>

Li, S., Wang, G., Guo, H., Cheng, J., & Yu, D. (2026). Machine learning for screening laryngopharyngeal reflux symptoms in college students: A cross-sectional study. *Annals of Medicine*, 58(1), 2610063. <https://doi.org/10.1080/07853890.2025.2610063>

Malley, J. D., Malley, K. G., & Pajevic, S. (2011). *Statistical Learning for Biomedical Data* (1 ed.). Cambridge University Press. <https://doi.org/10.1017/CBO9780511975820>

Nirwan, J. S., Hasan, S. S., Babar, Z.-U.-D., Conway, B. R., & Ghori, M. U. (2020). Global Prevalence and Risk Factors of Gastro-Oesophageal Reflux Disease

- (GORD): Systematic Review with Meta-Analysis. *Scientific Reports*, 10(1), 5814. <https://doi.org/10.1038/s41598-020-62795-1>
- Owess, M. M., Owda, A. Y., Owda, M., & Massad, S. (2024). Supervised Machine Learning-Based Models for Predicting Raised Blood Sugar. *International Journal of Environmental Research and Public Health*, 21(7), 840. <https://doi.org/10.3390/ijerph21070840>
- Priyambudi, Z. S., & Nugroho, Y. S. (2024). Which algorithm is better? An implementation of normalization to predict student performance. 020110. <https://doi.org/10.1063/5.0182879>
- Ramadhan, T. F., Asrianda, & Risawandi. (2025). Penerapan Metode Algoritma SVM (Support Vector Machine) Untuk Klasifikasi Penderita Penyakit Gastroesophageal Reflux Disease. *Rabit: Jurnal Teknologi dan Sistem Informasi Univrab*, 10(2), 1212–1219. <https://doi.org/10.36341/rabit.v10i2.6466>
- Ratnasari, E. D., Farida, I. N., & Kasih, P. (2025). Implementasi Metode Forward Chaining Dan Metode Naïve Bayes Untuk Mendeteksi Penyakit Lambung. 4.
- Rosidah, N., Prathivi, R., & Susanto, S. (2025). Optimasi Metode Random Forest Untuk Klasifikasi Risiko Obesitas Berdasarkan Pola Makan. *Jurnal Informatika Teknologi dan Sains (Jinteks)*, 7(1), 56–62. <https://doi.org/10.51401/jinteks.v7i1.5065>
- Sari, H. A., Hermawan, R., & Harris, S. (2025). Sistem Pakar Diagnosa Dini Penyakit Pada Lambung Dengan Metode Forward Chaining Di Klinik Dr. Yolanda. *Semnas Ristek (Seminar Nasional Riset dan Inovasi Teknologi)*, 9(1), 496–502. <https://doi.org/10.30998/semnasristek.v9i1.7573>
- Tang, F., Yu, H., Zhao, T., Lin, L., Dong, P., Sun, K., Sun, X., & Wang, Q. (2026). Diagnostic accuracy of gastric filling ultrasound combined with BMI for gastroesophageal reflux disease. *Frontiers in Medicine*, 13, 1727623. <https://doi.org/10.3389/fmed.2026.1727623>
- Wang, T., Chen, L., Li, Y., Zhang, Q., Ou, X., & Lu, Q. (2026). From body composition to reflux esophagitis: An interpretable machine learning model based on CT-derived features. *Frontiers in Physiology*, 17, 1788389. <https://doi.org/10.3389/fphys.2026.1788389>
- Wisesty, U. N., Delfina, H. A., & Kurniawan, I. (2025). Gastroesophageal Reflux Disease Early Detection using XGBoost Method Classifier. *Jurnal Teknik Informatika (Jutif)*, 6(2), 855–870. <https://doi.org/10.52436/1.jutif.2025.6.2.4143>