



METODE ENSEMBLE UNTUK MENGATASI DATA TIDAK SEIMBANG PADA PREDIKSI KANKER SERVIKS

Yudhistira¹, Suhardjono², Imam Nawawi³, Ahmad Hafidzul Kahfi⁴, Febri Ainun Jariyah⁵

^{1,2,3,4,5}Universitas Bina Sarana Informatika

^{1,2,3,4,5}Jl. Kramat Raya No. 98, Senen, Jakarta Pusat 10450.

E- mail : yudhistira.yht@bsi.ac.id

ABSTRAK

Kanker serviks merupakan salah satu penyakit dengan tingkat mortalitas yang tinggi pada perempuan, sehingga deteksi dini menjadi sangat penting dalam upaya pencegahan dan penanganan klinis. Penerapan machine learning pada prediksi kanker serviks sering menghadapi permasalahan ketidakseimbangan kelas, di mana jumlah data pasien kanker jauh lebih sedikit dibandingkan non-kanker, yang berpotensi menurunkan kemampuan model dalam mendeteksi kasus positif. Penelitian ini bertujuan untuk mengevaluasi efektivitas pendekatan ensemble learning khusus data tidak seimbang, yaitu Balanced Random Forest, EasyEnsemble, dan RUSBoost, dalam meningkatkan performa prediksi risiko kanker serviks. Evaluasi dilakukan menggunakan Stratified K-Fold Cross Validation dengan lima fold dan metrik akurasi, precision, recall, F1-score, serta Area Under the ROC Curve (AUC-ROC). Hasil eksperimen menunjukkan bahwa Balanced Random Forest memberikan performa terbaik dengan nilai akurasi sebesar 95.46%, recall 87.27%, F1-score 71.12%, dan AUC-ROC tertinggi sebesar 0.9466. EasyEnsemble menghasilkan akurasi sebesar 95.69% dengan F1-score tertinggi sebesar 72.18% dan AUC-ROC 0.9304, sedangkan RUSBoost menunjukkan sensitivitas yang baik terhadap kelas minoritas dengan nilai recall sebesar 81.82% dan AUC-ROC 0.9229. Hasil penelitian ini membuktikan bahwa pendekatan ensemble khusus untuk data tidak seimbang mampu meningkatkan kemampuan deteksi kelas minoritas secara signifikan dan lebih efektif dibandingkan pendekatan klasifikasi konvensional pada prediksi risiko kanker serviks.

Kata Kunci : *Balanced Random Forest, Data Tidak Seimbang, EasyEnsemble, Kanker Serviks, RUSBoost.*

ABSTRACT

Cervical cancer is a disease with a high mortality rate in women, so early detection is very important in prevention efforts and clinical treatment. The application of machine learning in cervical cancer prediction often faces class reflection problems, where the amount of cancer patient data is much less than non-cancer, which has the potential to reduce the model's ability to detect positive cases. This study aims to determine the effectiveness of ensemble learning approaches specifically for imbalanced data, namely Balanced Random Forest, EasyEnsemble, and RUSBoost, in improving the performance of cervical cancer risk prediction. The evaluation was carried out using Stratified K-Fold Cross Validation with five-fold and metrics of accuracy, precision, recall, F1-score, and Area Under the ROC Curve (AUC-ROC). The experimental results showed that Balanced Random Forest provided the best performance with an accuracy value of 95.46%, recall of 87.27%, F1-score of 71.12%, and the highest AUC-ROC of 0.9466. EasyEnsemble produced an accuracy of 95.69% with the highest F1-score of 72.18% and an AUC-ROC of 0.9304, while RUSBoost showed good sensitivity to the minority class with a recall value of 81.82% and an AUC-ROC of 0.9229. The results of this study prove that the ensemble approach specifically for imbalanced data can significantly improve the detection ability of minority classes and is more effective than conventional classification approaches in predicting cervical cancer risk.

Keywords: *Balanced Random Forest, Cancer Cerviks, EasyEnsemble, Imbalance Data, RUSBoost.*



1. PENDAHULUAN

Kanker serviks merupakan ancaman kesehatan yang signifikan bagi perempuan di seluruh dunia (Vazquez et al., 2025). Dalam beberapa dekade terakhir, pendekatan *machine learning* telah banyak dimanfaatkan untuk mendukung deteksi dini penyakit ini ((Salmi, Atif, Oliva, Abraham, & Sebastian Ventura, 2024). Pendekatan ini memungkinkan pemodelan hubungan nonlinier pada data medis serta evaluasi performa model secara sistematis menggunakan teknik validasi dan metrik evaluasi yang komprehensif, sebagaimana dijelaskan dalam literatur *machine learning* modern (Geron, 2022). Namun, penerapan *machine learning* pada kasus medis sering menghadapi salah satu tantangan terbesar berupa data tidak seimbang (*imbalanced class*), di mana jumlah sampel dari kelas minoritas (penderita kanker) jauh lebih sedikit dibandingkan kelas mayoritas (bukan kanker) (Yang, Khorshidi, & Aickelin, 2024), (Altalhan, Algarni, & Monia, 2025).

Dominasi kelas mayoritas dalam data latih sering menyebabkan model prediksi menjadi bias, sehingga kemampuan model dalam mendeteksi kasus langka namun krusial secara klinis menjadi menurun. Kondisi ini umum dijumpai pada dataset kanker serviks, di mana meskipun nilai akurasi keseluruhan terlihat tinggi, performa model dalam mengenali kelas minoritas sering kali rendah (Mudawi & Alazeb, 2022). Akibatnya, metrik akurasi menjadi kurang representatif dalam menggambarkan kualitas prediksi yang sesungguhnya, terutama dari sudut pandang klinis yang menekankan pentingnya deteksi dini (Muraru, Simó, & Iantovics, 2024).

Penanganan data tidak seimbang menjadi aspek krusial dalam prediksi kanker serviks, mengingat kesalahan klasifikasi pada kelas minoritas (*false negative*) berpotensi menimbulkan dampak serius bagi pasien (Gurcan & Soylu, 2024). Berbagai pendekatan telah dikembangkan untuk mengatasi permasalahan ini, antara lain teknik *resampling* seperti *SMOTE* dan *ADASYN*, serta pendekatan algoritmik berbasis *ensemble learning* (Siregar & Arifin, 2024).

Sejumlah studi menunjukkan bahwa tanpa penerapan teknik *balancing*, model klasifikasi kanker serviks dapat menghasilkan performa yang mendekati tebakan acak akibat ketidakseimbangan data yang ekstrem (Glučina, Ariana Lorencin, Nikola Anđelić, & Ivan Lorencin, 2023). Penelitian lain yang membandingkan berbagai metode *resampling* pada dataset kanker serviks menemukan bahwa teknik *balancing* secara signifikan meningkatkan kemampuan model dalam mengenali kelas minoritas dibandingkan model tanpa

penanganan *imbalance* (Muraru et al., 2024). Bahkan, penerapan *ADASYN* pada *Random Forest* terbukti mampu meningkatkan nilai recall kelas minoritas pada dataset dengan rasio ketidakseimbangan tinggi (95:5), yang menegaskan pentingnya strategi *balancing* sebelum proses evaluasi model (Saputra et al., 2025).

Namun demikian, penggunaan *Stratified K-Fold Cross Validation (SKCV)* pada data tidak seimbang belum sepenuhnya efektif dalam menekan bias pembelajaran, meskipun proporsi kelas pada setiap *fold* tetap terjaga (Çorbacioğlu & Aksel, 2023). Hal ini menunjukkan bahwa pendekatan validasi saja belum cukup, sehingga diperlukan metode klasifikasi yang secara eksplisit mengintegrasikan mekanisme penanganan ketidakseimbangan kelas ke dalam proses pembelajaran model, khususnya pada tahap pelatihan (Huang & Dai, 2021).

Salah satu strategi yang dinilai efektif untuk mengatasi permasalahan tersebut adalah pendekatan *ensemble learning* (Ayodele, 2023). Penelitian sebelumnya menunjukkan bahwa *ensemble learning* mampu memberikan peningkatan performa yang signifikan dalam pengklasifikasian data tidak seimbang. Sebagai contoh, implementasi *Balanced Random Forest* serta variasi pendekatan *SMOTE-RF* dilaporkan memiliki performa yang kompetitif dibandingkan berbagai teknik klasifikasi lainnya pada skenario data tidak seimbang (Fulazzaky, Saefuddin, & Soleh, 2024). Selain itu, hasil eksperimen pada berbagai aplikasi medis menunjukkan bahwa kombinasi teknik *resampling* dan *ensemble learning* dapat meningkatkan kemampuan model dalam mendeteksi kelas minoritas secara substansial, terutama ketika diterapkan pada dataset dengan tingkat ketidakseimbangan yang tinggi seperti data kanker (Gurcan & Soylu, 2024).

Studi terbaru juga menunjukkan bahwa metode *imbalance-aware ensemble* seperti *Balanced Random Forest* dan *RUSBoost* mampu meningkatkan performa *balanced accuracy* dan *recall* pada berbagai dataset tidak seimbang, serta menunjukkan keunggulan dibandingkan model konvensional maupun teknik *resampling* murni seperti *SMOTE-RF* atau *SMOTEBoost* (Fulazzaky et al., 2024). Selain itu, pendekatan *sampling* berbasis *EasyEnsemble* terbukti memberikan peningkatan kinerja yang signifikan dalam skenario data tidak seimbang, terutama pada metrik yang berfokus pada kelas minoritas (Gurcan & Soylu, 2024). Penelitian lain dalam domain diagnosis kanker juga melaporkan bahwa model *ensemble* seperti *Balanced Random Forest*, *EasyEnsemble*, dan *RUSBoost* merupakan pendekatan yang menjanjikan untuk meningkatkan sensitivitas terhadap kelas minoritas



yang secara klinis penting (Fulazzaky et al., 2024).

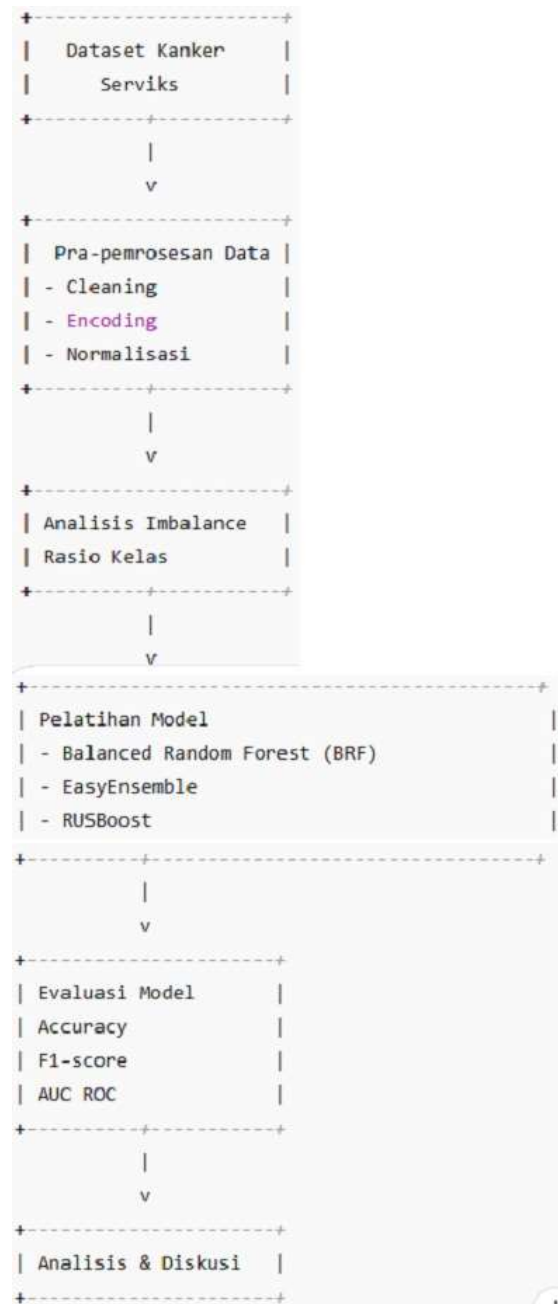
Berdasarkan latar belakang tersebut, penelitian ini mengadopsi pendekatan berbasis *ensemble* khusus untuk data tidak seimbang, yaitu *Balanced Random Forest*, *EasyEnsemble*, dan *RUSBoost*, dalam rangka meningkatkan performa prediksi risiko kanker serviks dengan fokus utama pada peningkatan kemampuan deteksi kelas minoritas.

Tujuan penelitian untuk mengevaluasi efektivitas pendekatan berbasis *ensemble* khusus untuk menangani data tidak seimbang dalam prediksi risiko kanker serviks. Secara khusus, penelitian ini mengkaji kinerja metode *Balanced Random Forest*, *EasyEnsemble*, dan *RUSBoost* dalam meningkatkan kemampuan model dalam mendeteksi kelas minoritas. Evaluasi performa dilakukan secara komprehensif menggunakan metrik akurasi, *F1-score*, dan *AUC ROC* guna memperoleh gambaran yang lebih representatif terhadap kualitas prediksi model pada data kanker serviks yang tidak seimbang.

Penelitian ini memberikan kontribusi dan kebaruan dalam pengembangan model prediksi risiko kanker serviks berbasis *machine learning*. Kontribusi utama penelitian ini terletak pada penerapan dan evaluasi teknik *ensemble* yang secara khusus dirancang untuk menangani data tidak seimbang, yaitu *Balanced Random Forest*, *EasyEnsemble*, dan *RUSBoost*, yang masih relatif terbatas dalam kajian prediksi kanker serviks (Mulugeta, Zewotir, Tegegne, Juhar, & Muleta, 2023). Evaluasi model dilakukan secara menyeluruh dengan memanfaatkan SKCV serta metrik evaluasi yang berfokus pada kelas minoritas, sehingga hasil penelitian ini diharapkan dapat menjadi acuan empiris yang dapat direplikasi dan digunakan sebagai dasar *benchmarking* pada penelitian selanjutnya di bidang prediksi penyakit berbasis data medis.

2. METODE PENELITIAN

Penelitian ini menggunakan pendekatan eksperimen kuantitatif dengan memanfaatkan algoritma *machine learning* (Ridwansyah, Iqbal, Destiana, Sugiono, & Hamid, 2024), untuk prediksi risiko kanker serviks pada data yang bersifat tidak seimbang (*imbalanced class*). Fokus utama penelitian adalah mengevaluasi efektivitas *ensemble learning* khusus data tidak seimbang, yaitu *Balanced Random Forest*, *EasyEnsemble*, dan *RUSBoost*, dibandingkan dengan metode klasifikasi konvensional yang dikombinasikan dengan teknik *resampling*. Berikut alur penelitian yang digunakan pada Gambar 1.



Gambar. 1 Alur Penelitian

Gambar 1 menunjukkan alur metodologi penelitian yang digunakan dalam pengembangan model prediksi risiko kanker serviks berbasis data tidak seimbang. Proses diawali dengan pengumpulan dataset kanker serviks yang memiliki distribusi kelas tidak seimbang antara pasien dengan dan tanpa indikasi kanker. Dataset ini kemudian melalui tahap prapemrosesan data, yang mencakup penanganan nilai hilang, transformasi data ke bentuk numerik (Sartini, Sumarna, Hamid, Kahf, & Nicodias Palasar, 2025), serta normalisasi fitur untuk memastikan kesesuaian data dengan algoritma pembelajaran mesin yang digunakan (Nurdin, Carolina, Andharsaputri, Wuryanto, & Ridwansyah, 2024).



Pada tahap selanjutnya, data yang telah dipraproses dibagi menggunakan metode *SKCV*. Pendekatan ini bertujuan untuk menjaga proporsi kelas minoritas dan mayoritas pada setiap *fold*, sehingga proses evaluasi model menjadi lebih stabil dan representatif terhadap kondisi data asli. Setiap *fold* digunakan secara bergantian sebagai data uji, sementara *fold* lainnya digunakan sebagai data latih.

Tahap pemodelan dilakukan dengan menerapkan beberapa metode *ensemble learning* yang dirancang khusus untuk menangani ketidakseimbangan kelas, yaitu *Balanced Random Forest*, *EasyEnsemble*, dan *RUSBoost*.

Seluruh model yang dibangun kemudian dievaluasi menggunakan metrik performa yang berfokus pada kemampuan deteksi kelas minoritas, yaitu *accuracy*, *precision*, *recall*, *F1-score*, dan *Area Under the ROC Curve (AUC ROC)* (Purnama, Nawawi, Rosyida, Ridwansyah, & Risandar, 2020). Hasil evaluasi ini selanjutnya dianalisis secara komparatif untuk menentukan metode yang paling optimal dalam meningkatkan performa prediksi kanker serviks pada dataset dengan ketidakseimbangan kelas yang signifikan.

A. Dataset dan Pra-pemrosesan Data

Dataset kanker serviks digunakan sebagai sumber data utama, yang terdiri dari sejumlah atribut klinis dan demografis dengan label kelas kanker dan non-kanker. Tahapan pra-pemrosesan meliputi (Priyatama & Ridwansyah, 2022):

1. *Data cleaning* untuk menangani nilai hilang dan inkonsistensi data.
2. *Encoding* untuk mengubah atribut kategorikal menjadi numerik.
3. Tahapan pra-pemrosesan meliputi data *cleaning* untuk menangani nilai hilang dan transformasi data ke bentuk numerik. Normalisasi fitur tidak diterapkan karena algoritma berbasis pohon keputusan yang digunakan tidak sensitif terhadap perbedaan skala fitur.

B. Penanganan Ketidakseimbangan Data

Dataset yang digunakan menunjukkan ketidakseimbangan kelas yang signifikan, di mana jumlah sampel kelas minoritas jauh lebih sedikit dibandingkan kelas mayoritas. Untuk mengatasi permasalahan ini, penelitian menerapkan pendekatan *ensemble* khusus data tidak seimbang, yang secara eksplisit mengintegrasikan mekanisme sampling ke dalam proses pelatihan model.

Ensemble learning menggabungkan beberapa model untuk menghasilkan prediksi yang lebih stabil dan akurat dibandingkan model tunggal. Teknik-teknik *ensemble* yang lebih maju tidak hanya meningkatkan akurasi keseluruhan tetapi juga mampu memperbaiki performa pada kelas minoritas.

C. Stratified K-Fold Cross Validation

Evaluasi model dilakukan menggunakan *SKCV* dengan $k=5$. Teknik ini memastikan proporsi kelas pada setiap *fold* tetap konsisten dengan distribusi data asli.

Secara matematis, performa model dihitung sebagai rata-rata dari seluruh *fold*:

$$performance = \frac{1}{k} \sum_{i=1}^k M_i \quad (1)$$

Dengan M_i merupakan nilai metrik evaluasi *fold* ke- i , dan k jumlah *fold*.

D. Metode Klasifikasi

Beberapa teknik klasifikasi yang digunakan dalam penelitian antara lain:

Balanced Random Forest

Random Forest merupakan metode *ensemble* berbasis pohon keputusan yang menggabungkan banyak classifier untuk meningkatkan stabilitas dan akurasi prediksi, dalam kerangka *statistical learning*. *Balanced Random Forest* merupakan pengembangan dari *Random Forest* yang menerapkan *undersampling* kelas mayoritas pada setiap pohon.

$$D_i = \text{Undersample}(D_{\text{majority}}) \cup D_{\text{minority}} \quad (2)$$

Setiap pohon dilatih menggunakan subset data seimbang, sehingga meningkatkan sensitivitas terhadap kelas minoritas. Modifikasi *random forest* yang melakukan *undersampling* otomatis pada setiap pohon untuk menjaga keseimbangan kelas secara internal

EasyEnsemble

EasyEnsemble membangun beberapa model *ensemble* berdasarkan beberapa subset *undersampling* berbeda dari kelas mayoritas. Prediksi akhir diperoleh melalui agregasi:

$$y = \text{majority_vote}(h_1(x), h_2(x), \dots, h_n(x)) \quad (3)$$

Pendekatan ini mampu mempertahankan informasi kelas mayoritas tanpa mengorbankan performa kelas minoritas.

RUSBoost

RUSBoost merupakan metode *ensemble* yang menggabungkan teknik *Random Under-Sampling (RUS)* dengan algoritma *boosting*, di mana bobot sampel diperbarui pada setiap iterasi untuk meningkatkan sensitivitas model terhadap kelas minoritas..

$$W_i^{(t+1)} = W_i^{(t)} \cdot e^{arI((y_i \neq h_t(x)))} \quad (4)$$



Metode ini efektif dalam meningkatkan fokus model pada sampel minoritas yang sulit diklasifikasikan

E. Metrik Evaluasi

Evaluasi model dilakukan menggunakan metrik berikut (Ridwansyah, Riyanto, Hamid, Rahayu, & Purnama, 2022):

1. Akurasi

Accuracy

$$= \frac{True\ Positive\ (TP) + True\ Negative\ (TN)}{TP + TN + False\ Positive\ (FP) + False\ Negative\ (FN)} \quad (5)$$

2. Precision

$$Precision = \frac{TP}{TP + FP} \quad (6)$$

3. Recall

$$Recall = \frac{TP}{TP + FN} \quad (7)$$

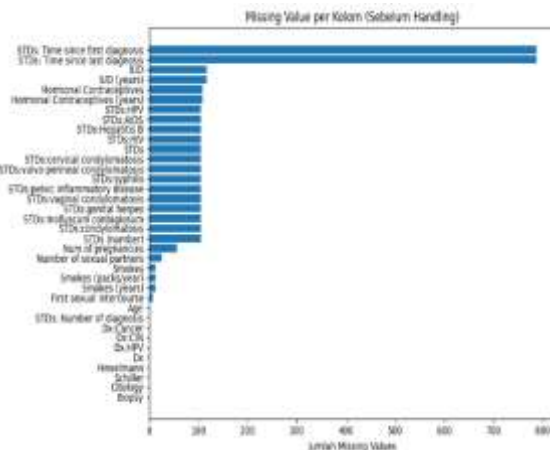
4. F1-Score

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (8)$$

5. AUC ROC

AUC mengukur kemampuan model dalam membedakan kelas positif dan negatif berdasarkan kurva ROC (Riyanto, Destiana, Prihatin, Sugiono, & Wijaya, 2025).

3. HASIL DAN PEMBAHASAN

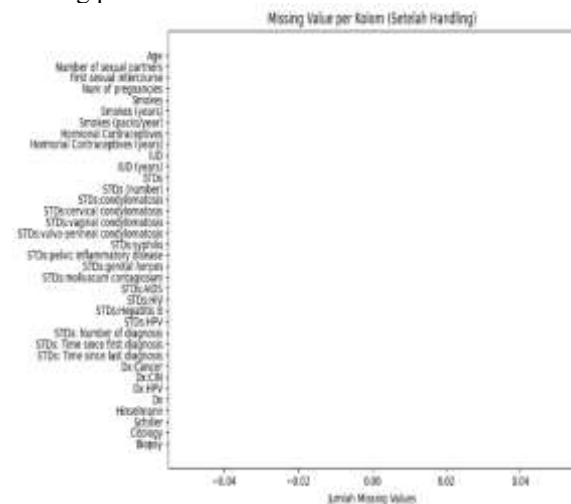


Gambar. 2 Distribusi Missing Value per Kolom (Sebelum *Handling*)

Gambar 2 menunjukkan distribusi jumlah missing value pada setiap atribut sebelum dilakukan penanganan data hilang. Terlihat bahwa beberapa atribut memiliki jumlah missing value yang cukup signifikan, yang berpotensi memengaruhi proses pembelajaran model jika tidak ditangani dengan

baik. Kondisi ini umum dijumpai pada dataset medis karena keterbatasan pemeriksaan klinis atau pencatatan data pasien yang tidak lengkap.

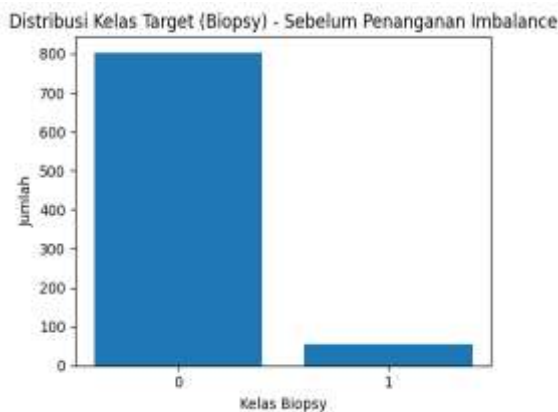
Berdasarkan hasil tersebut, dilakukan proses penanganan missing value dengan pendekatan imputasi adaptif, yaitu menggunakan nilai modus untuk atribut dengan variasi kecil (atribut biner atau kategorikal sederhana) dan nilai median untuk atribut numerik dengan variasi yang lebih besar. Pendekatan ini dipilih karena median relatif lebih robust terhadap outlier, sedangkan modus efektif untuk mempertahankan distribusi atribut diskrit. Setelah proses imputasi, seluruh missing value berhasil diatasi, sebagaimana ditunjukkan pada Gambar 3, di mana tidak ditemukan lagi nilai kosong pada dataset.



Gambar. 3 Distribusi Missing Value per Kolom (Setelah *Handling*)

Hasil pada Gambar 3 menegaskan bahwa dataset telah siap digunakan untuk tahap pemodelan tanpa risiko bias akibat data hilang.

Setelah dataset telah siap maka akan dilakukan distribusi kelas dengan memperlihatkan distribusi kelas target Biopsy sebelum dilakukan penanganan ketidakseimbangan kelas. Terlihat adanya dominasi kelas mayoritas (non-kanker) dibandingkan kelas minoritas (kanker serviks), yang mengindikasikan bahwa dataset bersifat tidak seimbang (imbalanced class) yang terlihat pada Gambar 4.



Gambar. 4 Distribusi Kelas Target Sebelum Penanganan Imbalance

Ketidakeimbangan pada Gambar 4 ini berpotensi menyebabkan model klasifikasi menjadi bias terhadap kelas mayoritas dan menurunkan kemampuan deteksi kelas minoritas yang secara klinis lebih penting. Oleh karena itu, diperlukan pendekatan klasifikasi yang secara eksplisit dirancang untuk menangani kondisi data tidak seimbang.

Untuk memastikan evaluasi model yang adil dan stabil, penelitian ini menggunakan *SKCV* dengan lima *fold*. Teknik ini mempertahankan proporsi kelas pada setiap *fold* sehingga distribusi kelas pada data latih dan data uji tetap representatif terhadap distribusi data asli.

Namun demikian, sebagaimana dibahas pada bagian pendahuluan, stratifikasi saja belum cukup untuk mengatasi bias pembelajaran pada data tidak seimbang. Oleh karena itu, *SKCV* digunakan sebagai mekanisme evaluasi, sementara penanganan imbalance dilakukan secara langsung melalui algoritma ensemble yang digunakan.

Tiga metode ensemble khusus data tidak seimbang yang diuji dalam penelitian ini adalah *Balanced Random Forest*, *EasyEnsemble*, dan *RUSBoost*. Evaluasi performa dilakukan menggunakan lima metrik, yaitu *Accuracy*, *Precision*, *Recall*, *F1-score*, dan *AUC-ROC*. Gambar 5 menampilkan perbandingan performa ketiga model berdasarkan masing-masing metrik evaluasi.

	Model	Accuracy	Precision	Recall	F1-score	AUC
0	Balanced RF	0.954563	0.601111	0.872727	0.711150	0.946581
1	EasyEnsemble	0.956888	0.617183	0.872727	0.721817	0.930420
2	RUSBoost	0.944098	0.554405	0.818182	0.658543	0.922863

Gambar. 5 Perbandingan Performa Model Ensemble pada Berbagai Metrik

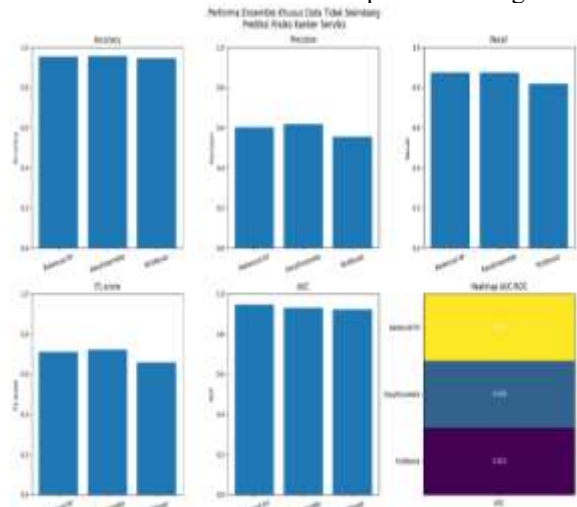
Berdasarkan Gambar 5 hasil evaluasi

menggunakan *SKCV*, *Balanced Random Forest* memperoleh nilai akurasi sebesar 95.46%, *precision* 60.11%, *recall* 87.27%, *F1-score* 71.12%, dan *AUC-ROC* sebesar 0.9466. Hasil ini menunjukkan bahwa *Balanced Random Forest* memiliki kemampuan diskriminasi kelas yang sangat baik dengan keseimbangan antara akurasi dan sensitivitas terhadap kelas minoritas.

EasyEnsemble mencatatkan nilai akurasi tertinggi sebesar 95.69%, dengan *precision* 61.72%, *recall* 87.27%, *F1-score* 72.18%, serta *AUC-ROC* sebesar 0,9304. Nilai *F1-score* tertinggi yang diperoleh *EasyEnsemble* menunjukkan kestabilan performa dalam menangani ketidakseimbangan kelas.

Sementara itu, *RUSBoost* menghasilkan nilai *recall* sebesar 81.82%, namun dengan akurasi yang lebih rendah yaitu 94.41% dan *F1-score* 65.85%. Hal ini mengindikasikan bahwa *RUSBoost* lebih menekankan sensitivitas terhadap kelas minoritas, tetapi dengan kompromi pada ketepatan prediksi keseluruhan.

Heatmap AUC-ROC pada Gambar 6 memberikan gambaran komparatif kemampuan masing-masing model dalam membedakan kelas positif dan negatif.



Gambar. 6 Heatmap AUC-ROC Model Ensemble

Gambar 6 memperlihatkan nilai *AUC-ROC* yang dihasilkan masing-masing model adalah 0.9466 untuk *Balanced Random Forest*, 0.9304 untuk *EasyEnsemble*, dan 0.9229 untuk *RUSBoost*. Nilai AUC yang tinggi pada seluruh model menunjukkan bahwa pendekatan *ensemble* khusus data tidak seimbang memiliki kemampuan yang baik dalam membedakan kelas positif dan negatif pada dataset kanker serviks.

Hasil eksperimen menunjukkan bahwa pendekatan ensemble khusus untuk data tidak seimbang secara konsisten mampu meningkatkan performa klasifikasi pada dataset kanker serviks dibandingkan pendekatan konvensional tanpa



penanganan *imbalance* secara eksplisit. Keunggulan utama metode ensemble terlihat pada metrik yang berfokus pada kelas minoritas, seperti *recall* dan *F1-score*, yang sangat penting dalam konteks medis untuk meminimalkan kesalahan *false negative*.

Secara kuantitatif, model *Balanced Random Forest* menunjukkan performa terbaik secara keseluruhan dengan nilai AUC tertinggi sebesar 0.9466, akurasi sebesar 95.46%, dan nilai *recall* sebesar 87.27%. Nilai *recall* yang tinggi ini sangat krusial dalam konteks klinis karena menunjukkan kemampuan model dalam mendeteksi sebagian besar kasus kanker serviks yang sebenarnya. *EasyEnsemble* menghasilkan performa yang stabil dengan nilai akurasi sebesar 95.69% dan *F1-score* tertinggi sebesar 72.18%, yang mengindikasikan keseimbangan antara *precision* dan *recall*. Sementara itu, *RUSBoost* menunjukkan peningkatan sensitivitas terhadap kelas minoritas dengan nilai *recall* sebesar 81.82%, meskipun menghasilkan nilai *precision* yang lebih rendah dibandingkan dua metode lainnya.

Secara keseluruhan, hasil ini mengonfirmasi bahwa integrasi mekanisme penanganan ketidakseimbangan kelas langsung ke dalam algoritma pembelajaran, sebagaimana dilakukan oleh *Balanced Random Forest*, *EasyEnsemble*, dan *RUSBoost*, merupakan pendekatan yang lebih efektif dibandingkan mengandalkan teknik validasi atau metrik evaluasi semata.

4. KESIMPULAN

Penelitian ini membuktikan bahwa pendekatan ensemble khusus untuk data tidak seimbang mampu meningkatkan performa klasifikasi pada prediksi risiko kanker serviks, terutama dalam mendeteksi kelas minoritas yang secara klinis krusial. Berdasarkan evaluasi menggunakan SKCV, *Balanced Random Forest* menunjukkan performa terbaik dengan nilai AUC-ROC sebesar 0.9466, yang menandakan kemampuan diskriminasi kelas yang sangat baik.

EasyEnsemble dan *RUSBoost* juga memberikan performa yang kompetitif, dengan karakteristik keunggulan yang berbeda. *EasyEnsemble* menunjukkan kestabilan performa yang baik pada metrik *F1-score*, sedangkan *RUSBoost* lebih menonjol dalam meningkatkan sensitivitas terhadap kelas minoritas. Hasil ini menegaskan bahwa integrasi mekanisme penanganan ketidakseimbangan kelas langsung ke dalam algoritma pembelajaran merupakan strategi yang lebih efektif dibandingkan pendekatan konvensional tanpa penanganan *imbalance* secara eksplisit.

Saran untuk kedepannya mengeksplorasi

pendekatan lain dalam penanganan data tidak seimbang, seperti *cost-sensitive learning*, *hybrid ensemble-resampling*, maupun penerapan *deep learning* berbasis *attention* yang mampu menangkap hubungan nonlinier antar fitur klinis. Selain itu, penggunaan teknik seleksi fitur dan optimasi hiperparameter, seperti *Particle Swarm Optimization* atau *Bayesian Optimization*, berpotensi meningkatkan efisiensi dan generalisasi model

5. REFERENSI

- Altalhan, M., Algarni, A., & Monia, T. H. A. (2025). Imbalanced Data Problem in Machine Learning: A Review. *Turki Hadj Alouane Monia*, 11. <https://doi.org/10.1109/ACCESS.2025.3531662>
- Ayodele, A. (2023). A comparative study of ensemble learning techniques for imbalanced classification problems. *World Journal of Advanced Research and Review*, 19(1), 1633–1643. <https://doi.org/https://doi.org/10.30574/wjarr.2023.19.1.1202>
- Çorbacioğlu, Ş. K., & Aksel, G. (2023). Receiver operating characteristic curve analysis in diagnostic accuracy studies. *Turkish Journal of Emergency Medicine*, 23(4). https://doi.org/10.4103/tjem.tjem_182_23
- Fulazzaky, T., Saefuddin, A., & Soleh, A. M. (2024). Evaluating Ensemble Learning Techniques for Class Imbalance in Machine Learning: A Comparative Analysis of Balanced Random Forest, SMOTE-RF, SMOTEBoost, and RUSBoost. *Scientific Journal of Informatics*, 11(4). <https://doi.org/https://doi.org/10.15294/sji.v11i4.15937>
- Geron, A. (2022). *Hands on Machine learning with Scikit Learn Keras and Tensor Flow Concepts, Tools, and Techniques to Build Intelligent Systems* (3rd ed.). Sebastopol, CA, USA: O'Reilly Media, Inc.
- Glučina, M., Ariana Lorencin, Nikola Andelić, & Ivan Lorencin. (2023). Cervical Cancer Diagnostics Using Machine Learning Algorithms and Class Balancing Techniques. *Applied Sciences*, 13(12). <https://doi.org/https://doi.org/10.3390/app13021061>
- Gurcan, F., & Soylu, A. (2024). Learning from Imbalanced Data: Integration of Advanced Resampling Techniques and Machine Learning Models for Enhanced Cancer



- Diagnosis and Prognosis. *Cancers*, 16(19). <https://doi.org/https://doi.org/10.3390/cancers16193417>
- Huang, C. Y., & Dai, H. L. (2021). Learning from class-imbalanced data: review of data driven methods and algorithm driven methods. *Data Science in Finance and Economics*, 1(1), 21–36. <https://doi.org/10.3934/DSFE.2021002>
- Mudawi, N. Al, & Alazeb, A. (2022). A Model for Predicting Cervical Cancer Using Machine Learning Algorithms. *Sensors*, 22(11). <https://doi.org/https://doi.org/10.3390/s22114132>
- Mulugeta, G., Zewotir, T., Tegegne, A. S., Juha, L. H., & Muleta, M. B. (2023). Classification of imbalanced data using machine learning algorithms to predict the risk of renal graft failures in Ethiopia. *BMC Medical Informatics and Decision Making*, 23(98). <https://doi.org/10.1186/s12911-023-02185-5>
- Muraru, M. M., Simó, Z., & Iantovics, L. B. (2024). Cervical Cancer Prediction Based on Imbalanced Data Using Machine Learning Algorithms with a Variety of Sampling Methods. *Applied Sciences*, 14(22). <https://doi.org/https://doi.org/10.3390/app142210085>
- Nurdin, H., Carolina, I., Andharsaputri, R. L., Wuryanto, A., & Ridwansyah. (2024). Forward Selection as a Feature Selection Method in the SVM Kernel for Student Graduation Data. *Sinkron: Jurnal Dan Penelitian Teknik Informatika*, 8(October), 2531–2537. <https://doi.org/10.33395/sinkron.v8i4.14172>
- Priyatama, I. M. D., & Ridwansyah. (2022). Klasifikasi Anak Berkebutuhan Khusus Tunagrahita Menggunakan Metode Algoritma C4.5. *Paradigma*, 24(1), 90–95. <https://doi.org/https://doi.org/10.31294/paradigma.v24i1.1087>
- Purnama, J. J., Nawawi, H. M., Rosyida, S., Ridwansyah, & Risandar. (2020). Klasifikasi Mahasiswa Her Berbasis Algoritma Svm Dan Decision Tree. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 7(6), 1253–1260. <https://doi.org/10.25126/jtiik.202073080>
- Ridwansyah, Iqbal, M., Destiana, H., Sugiono, & Hamid, A. (2024). Data Mining Berbasis Machine Learning Untuk Analitik Prediktif Dalam Kelulusan. *Semantik*, 10(2), 1–10. <https://doi.org/https://doi.org/10.55679/semantik.v10i2.67>
- Ridwansyah, R., Riyanto, V., Hamid, A., Rahayu, S., & Purnama, J. J. (2022). Grouping Data in Predicting Infant Mortality Using K-Means and Decision Tree. *Paradigma*, 24(2), 168–174. <https://doi.org/10.31294/paradigma.v24i2.1399>
- Riyanto, V., Destiana, H., Prihatin, T., Sugiono, & Wijaya, G. (2025). Mengoptimalkan Prediksi Gagal Jantung Dengan Kombinasi Svm Dan Forward Selection. *JIRE (Jurnal Informatika & Rekayasa Elektronika)*, 8(1), 103–111. <https://doi.org/https://doi.org/10.36595/jire.v8i1.1541>
- Salmi, M., Atif, D., Oliva, D., Abraham, A., & Sebastian Ventura. (2024). Handling imbalanced medical datasets: review of a decade of research. *Artificial Intelligence Review (Springer Nature)*, 57(10). <https://doi.org/10.1007/s10462-024-10884-2>
- Saputra, R. M., Alzami, F., Pramudi, Y. T. C., Erawan, L., Megantara, R. A., Ricardus Anggi Pramunendar, & Yusuf, M. (2025). Improving Cervical Cancer Classification Using ADASYN and Random Forest with GridSearchCV Optimization. *Informatics, Electrical Engineering, and Mechanical Engineering*, 16(1). <https://doi.org/https://doi.org/10.35970/infotekmesin.v16i1.2552>
- Sartini, S., Sumarna, S., Hamid, A., Kahf, A. H., & Nicodias Palasar. (2025). REVOLUSI DIAGNOSIS : OPTIMASI RANDOM TREE-PSO UNTUK PENYAKIT GINJAL KRONIS. *JIRE (Jurnal Informatika & Rekayasa Elektronika)*, 8(1), 149–158. <https://doi.org/https://doi.org/10.36595/jire.v8i1.1542>
- Siregar, A. Y., & Arifin, A. S. (2024). Enhancing XGBoost Classification with SVM-SMOTE & EasyEnsemble for Imbalanced Telemedicine Sentiment Data. *Jurnal Indonesia Sosial Teknologi*, 5(10). <https://doi.org/https://doi.org/10.59141/jist.v5i10.1160>
- Vazquez, B., Rojas-García, M., Rodríguez-Esquivel, J. I., Marquez-Acosta, J., Aranda-Flores, C. E., Cetina-Pérez, L. del C., ... Torres-Poveda, K. (2025). Machine and Deep Learning for the Diagnosis, Prognosis, and Treatment of Cervical Cancer: A Scoping Review. *Diagnostics*, 15(12).
- Yang, Y., Khorshidi, H. A., & Aickelin, U. (2024). A review on over-sampling techniques in classification of multi-class imbalanced datasets: insights for medical problems. *Frontiers in Digital Health*, 6(1430245). <https://doi.org/10.3389/fdgth.2024.1430245>