

PENDEKATAN KOMPARATIF ALGORITMA MACHINE LEARNING UNTUK PREDIKSI KEMISKINAN GLOBAL

COMPARATIVE APPROACHES TO MACHINE LEARNING ALGORITHMS FOR GLOBAL POVERTY PREDICTION

Darvi Mailisa Putri^{1§}, Dina Friska²

¹Prodi Matematika, Fakultas Sains dan Teknologi UIN Imam Bonjol Padang [darvimailisa@uinib.ac.id]

²Prodi Matematika, Fakultas Sains dan Teknologi UIN Imam Bonjol Padang [dinafriskaa@gmail.com]

[§]Corresponding Author

Received 21th Oct 2024; Accepted 21th Nov 2024; Published 10th Des 2024;

Abstrak

Kemiskinan masih menjadi masalah sosial yang butuh perhatian khusus untuk ditangani. Banyak dampak yang ditimbulkan, diantaranya permasalahan pertumbuhan otak anak, meningkatnya penyakit jangka panjang, dan meningkatnya konflik sosial dan keamanan. Maka perlu usaha untuk mengatasi kasus kemiskinan secara efektif dengan pendekatan yang inovatif dan berbasis data yaitu penggunaan Algoritma machine learning. Algoritma ini dapat menganalisis data kemiskinan, mengidentifikasi pola, dan memprediksi risiko kemiskinan dengan lebih akurat. Penelitian ini fokus menganalisis performa algoritma machine learning yaitu Decision Trees dan Naïve Bayes. Hasil penelitian menunjukkan algoritma Decision Trees memiliki akurasi lebih baik (84,2%) dibandingkan akurasi algoritma Naïve Bayes (78,9%). Namun performa algoritma Naïve Bayes dalam memprediksi berbagai kelas lebih stabil dibanding algoritma Decision Trees berdasarkan nilai presisi, recall, F1-score dan specificity.

Kata Kunci: Kemiskinan, Decision Trees, Naïve Bayes

Abstract

Poverty remains a significant social issue that requires targeted attention for effective resolution. The consequences of this situation are numerous and far-reaching. They include adverse effects on child development, an increase in chronic illnesses, and a rise in social and security-related conflicts. It is therefore necessary to implement effective measures to address poverty in an innovative and data-driven manner, utilising machine learning algorithms. The algorithm is capable of analyzing poverty data, identifying patterns, and predicting poverty risks with greater accuracy. This study focuses on analyzing the performance of two machine learning algorithms, namely Decision Trees and Naïve Bayes. The findings of the study indicate that the Decision Trees algorithm exhibits superior accuracy (84.2%) in comparison to the Naïve Bayes algorithm, which displays an accuracy of 78.9%. However, the performance of the Naïve Bayes algorithm in predicting various classes is more stable than that of the Decision Trees algorithm, as evidenced by the values of precision, recall, F1-score, and specificity.

Keywords: Poverty, Decision Trees, Naïve Baiyes

1. Pendahuluan

Kemiskinan adalah salah satu masalah global yang kompleks dan mempengaruhi kehidupan jutaan orang di dunia. Laporan terbaru dari Bank Dunia menunjukkan bahwa 9.2% orang di dunia hidup dalam kemiskinan ekstrem pada tahun 2021, dengan pendapatan kurang dari \$2,15 per hari [1]. Meskipun dalam beberapa dekade terakhir telah terjadi penurunan dalam jumlah orang miskin di seluruh dunia, kemajuan tersebut cenderung tidak merata, dan kemiskinan masih tinggi di beberapa wilayah dan negara, terutama di Asia Selatan dan Sub-Sahara Afrika [2].

Tidak hanya kekurangan pendapatan yang menyebabkan kemiskinan, tetapi kemiskinan juga mencakup ketidakmampuan seseorang untuk mendapatkan akses ke sumber daya dasar seperti makanan, air bersih, layanan kesehatan, dan pendidikan [3]. Lebih dari 1,3 miliar orang di seluruh dunia terkena dampak kemiskinan, yang mencakup aspek seperti standar hidup, kesehatan, dan pendidikan, menurut penelitian yang dilakukan oleh United Nations Development Programme. Mayoritas dari mereka menetap di negara-negara berpenghasilan rendah dan menengah [4].

Kemiskinan disebabkan oleh sejumlah faktor yang sangat beragam dan saling terkait. Faktor ekonomi seperti pertumbuhan ekonomi yang lambat, ketimpangan pendapatan, dan kurangnya kesempatan kerja adalah penyebab umum kemiskinan. Ditambah elemen sosial dan budaya seperti diskriminasi etnis dan gender, serta akses yang tidak merata terhadap pendidikan dan layanan kesehatan, menyebabkan situasi tersebut

menjadi lebih buruk [5]. Selain itu, siklus kemiskinan sering diperkuat oleh korupsi, ketidakstabilan politik, dan kebijakan publik yang buruk [6].

Kemiskinan memiliki dampak yang luas dan mempengaruhi banyak aspek kehidupan sosial. Pertumbuhan otak dan tingkat kematian anak menjadi akibat dari kemiskinan [7], [8], [2]. Ditambah penyakit jangka panjang meningkat dan juga terbatasnya akses terhadap pendidikan berkualitas tinggi [9], [2]. Akibatnya, kemiskinan memperburuk siklus kemiskinan antar generasi. Dalam jangka panjang, kemiskinan dapat meningkatkan ketidakstabilan sosial dan politik, meningkatkan kemungkinan konflik sosial dan ketidakamanan [10].

Upaya untuk mengatasi kemiskinan secara efektif membutuhkan pendekatan yang inovatif dan berbasis data. Data kemiskinan tergolong ke dalam data time series. Model ARIMA merupakan model dasar untuk menganalisis data time series, namun data harus stasioner terhadap mean dan varians agar menghasilkan model yang optimal [11] [12]. Selain itu, asumsi white noise juga menjadi hal perlu diperhatikan pada model ARIMA [13]. Pendekatan terbaru yang sedang berkembang adalah penggunaan Algoritma machine learning. Algoritma ini dapat menganalisis data time series, mengidentifikasi pola, dan memprediksi risiko dengan lebih akurat. Selain itu, sebagai sub-bidang kecerdasan buatan, menawarkan potensi besar dalam mengolah berbagai jenis data [14].

Algoritma machine learning juga

digunakan untuk menganalisis faktor-faktor yang berkontribusi terhadap kemiskinan. Misalnya, Balasubramanian et al. (2020) menggunakan model random forest dan support vector machines untuk mengeksplorasi keterkaitan antara kemiskinan dan faktor-faktor seperti akses terhadap pendidikan, layanan kesehatan, dan infrastruktur di India. Penelitian ini menemukan bahwa kombinasi berbagai jenis data dan algoritma machine learning dapat meningkatkan kemampuan untuk memprediksi dan memahami kemiskinan di berbagai konteks geografis dan sosial.

Berdasarkan latar belakang ini, penelitian lebih lanjut akan diterapkan algoritma machine learning pada prediksi data tingkat kemiskinan. Algoritma yang digunakan adalah algoritma Decision Trees dan Naïve Bayes. Selanjutnya, pada akhir penelitian akan dianalisis kinerja algoritma machine learning melalui perbandingan hasil akurasi dan performa yang diperoleh dari kedua algoritma.

2. Landasan Teori

Algoritma Machine Learning

Machine Learning adalah cabang dari kecerdasan buatan (AI) yang memungkinkan sistem untuk belajar dari data tanpa pemrograman eksplisit [14]. Pada penelitian diterapkan tiga algoritma *Machine Learning* yaitu *Decision Trees*, *Naïve Bayes*, dan *Support Vector Machine*. Berikut penjelasan masing-masing algoritma.

Decision Tree (DT)

Decision Trees (DT) adalah bagian algoritma *Machine Learning* yang menggunakan

struktur pohon untuk membangun model regresi atau klasifikasi [16]. DT menjadi salah satu metode kuat yang biasa digunakan di berbagai bidang, seperti pembelajaran mesin, pemrosesan gambar, dan identifikasi pola [17]. Simpul-simpul dan cabang terdiri dari setiap pohon. Setiap node mewakili fitur dalam kategori yang akan diklasifikasikan dan setiap subset mendefinisikan nilai yang dapat diambil oleh node tersebut [18]. Proses algoritma dimulai dengan menciptakan cabang untuk setiap nilai, memilih atribut sebagai akar, dan mengulangi proses ini untuk setiap cabang sampai semua kasus di cabang memiliki kelas yang sama. Terakhir, atribut dipilih berdasarkan nilai gain tertinggi dari semua atribut yang ada.

$$\begin{aligned} \text{Information Gain}(S, A) &= \text{Entropy}(S) - \\ &\quad \sum \frac{|S_v|}{|S|} \text{Entropy}(S_v) \\ \text{Entropy}(S) &= - \sum_{i=1}^c P_i \log_2(P_i) \end{aligned} \quad (1)$$

Keterangan:

- S : Dataset.
- A : Fitur yang digunakan untuk membagi dataset.
- S_v : Subset dari S di mana fitur A memiliki nilai v .
- c : Banyak kelas.
- p_i : Proporsi elemen dalam kelas ke- i .

Naïve Bayes (NB)

Naïve Bayes (NB) adalah algoritma klasifikasi yang sederhana namun memiliki hasil akurat yang didasarkan pada penerapan Teorema Bayes. Teorema ini menggambarkan hubungan antara probabilitas kondisi dan probabilitas marginal dari dua kejadian. Asumsi yang dibawa pada NB adalah bahwa fitur-fitur dalam dataset bersifat independen satu sama lain atau dapat

dikatakan dengan istilah “Naïve”. Artinya kehadiran atau ketiadaan suatu fitur tidak mempengaruhi fitur lainnya. Proses klasifikasi bekerja dengan menghitung probabilitas posterior untuk setiap kelas, dan kemudian memilih kelas dengan probabilitas tertinggi sebagai prediksi [19].

$$P(C|X) = \frac{P(X|C)P(C)}{P(X)} \quad (2)$$

Keterangan:

X : Data dengan kelas belum diketahui.

C : Hipotesis data (spesifik kelas).

$P(C|X)$: Probabilitas hipotesis berdasarkan kondisi (*posterior probability*).

$P(X|C)$: Probabilitas berdasarkan kondisi pada hipotesis (*conditional probability*).

$P(C)$: Probabilitas hipotesis (*prior probability*).

$P(X)$: Probabilitas hipotesis S (*evidence*).

Matriks Akurasi

Hal yang penting dalam pembahasan *Machine Learning* adalah matriks akurasi atau dikenal juga dengan *confusion matrix*. Matriks akurasi berguna untuk mendapatkan pengetahuan yang mendalam tentang performa model dalam melakukan klasifikasi data. Matriks ini memberikan gambaran apakah model memprediksi kelas yang benar atau salah, melalui perhitungan jumlah prediksi benar dan salah yang dihasilkan oleh model, lalu dibandingkan dengan label sebenarnya. Terdapat empat nilai atau komponen utama pada matriks akurasi untuk masalah klasifikasi biner, yaitu True Positif (TP), False Positive (FP), False Negative (FN) dan True Negative (TN) yang ilustrasinya ditampilkan pada Tabel 1. Berdasarkan komponen utama pada matriks akurasi, pada tahap evaluasi dapat dilakukan perhitungan nilai accuracy, precision,

recall, f1-score dan specificity pada model

Tabel 1. Ilustrasi Matriks Akurasi

	Prediksi Positif	Prediksi Negatif
Kelas Positif	True Positive (TP)	False Negative (FN)
Kelas Negatif	False Positive (FP)	True Negative (TN)

Keterangan:

TP : Jumlah sampel positif yang diprediksi dengan benar sebagai positif oleh model.

TN : Jumlah sampel negatif yang diprediksi dengan benar sebagai negative oleh model.

FP : Jumlah sampel negatif yang salah diprediksi sebagai positif oleh model.

FN : Jumlah sampel positif yang salah diprediksi sebagai negatif oleh model.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (4)$$

$$Recall = \frac{TP}{TP + FN} \quad (5)$$

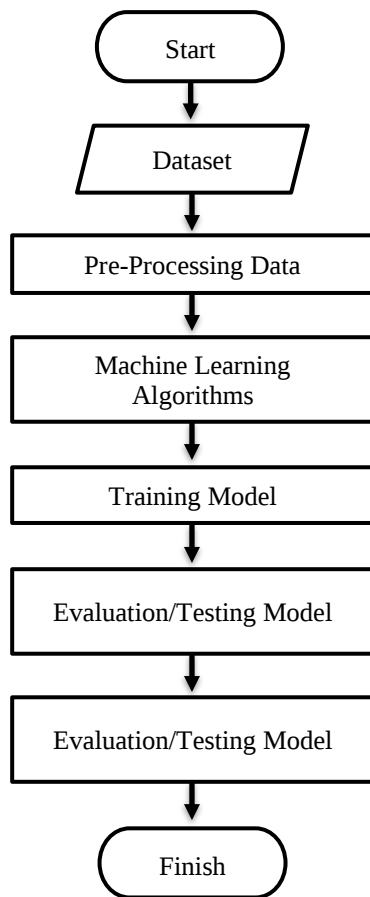
$$Precision = \frac{TP}{TP + FP} \quad (6)$$

$$F1 - Score = 2 \times \frac{(Precision \times Recall)}{(Precision + Recall)} \quad (7)$$

$$Specificity = \frac{TN}{TN + FP} \quad (8)$$

Perancangan dan Implementasi

Perancangan proses klasifikasi dengan menggunakan algoritma *Machine Learning*, yaitu *Decision Trees*, *Naïve Bayes*, dan *Support Vector Machine* ditampilkan pada Gambar 1.



Gambar 2. Flow Chart Penelitian

Proses awal dimulai dengan penginputan dataset yang akan dikaji pada penelitian yaitu terdapat enam variabel yang menjadi indikator *Multidimensional Poverty Headcount Ratio* (Rasio Jumlah Penduduk Miskin Multidimensi), diantaranya;

1. Monetary (Keuangan)

Sebuah rumah tangga dikatakan miskin jika pendapatan atau pengeluaran, dalam paritas daya beli dolar AS tahun 2017, kurang dari US\$2,15 per orang per hari. Hal ini didasarkan pada data mikro yang tersedia dari Studi Pendapatan Luksemburg.

2. Educational Attainment (Pencapaian Pendidikan)

Sebuah rumah tangga dikatakan tidak mampu jika tidak ada orang dewasa (usia

setara kelas 9 atau lebih tua) yang telah menyelesaikan pendidikan dasar.

3. Educational enrolment (Pendaftaran Pendidikan)

Sebuah rumah tangga dianggap miskin jika setidaknya ada satu anak usia sekolah hingga usia (setara) kelas 8 yang tidak terdaftar di sekolah.

4. Electricity (Listrik)

Sebuah rumah tangga dikatakan miskin jika tidak memiliki akses listrik.

5. Sanitation (Sanitasi)

Sebuah rumah tangga dikatakan miskin jika tidak memiliki akses terhadap sanitasi dengan standar terbatas.

6. Drinking Water (Air Mimum)

Sebuah rumah tangga dikatakan terdeprivasi jika tidak memiliki akses terhadap air minum berstandar terbatas.

Proses selanjutnya adalah melakukan pre-processing data, guna mempersiapkan kualitas data yang akan mempengaruhi performa model. Hal ini menjadi tahapan yang krusial dalam kerja *Machine Learning*. Sebelum Algoritma *Machine Learning* diterapkan, dataset dibagi menjadi dua bagian yaitu sebagian besar sebagai data training (75%) dan sisanya sebagai data testing (25%). Proses testing dapat dilakukan jika proses training telah selesai. Maka hasil akhir akan mendapatkan output berupa klasifikasi rasio jumlah penduduk miskin multidimensi. Output ini akan digunakan untuk proses evaluasi dengan mengukur akurasi yang dihasilkan algoritma *Machine Learning*.

3. Hasil Dan Pembahasan

Proses awal yang dilakukan pada dataset adalah pre-processing data, guna mendapatkan hasil yang optimal. Pre-processing data yang diterapkan adalah imputasi, yaitu mengisi *missing value* dengan nilai mean atau median. Mean digunakan saat data berdistribusi normal, namun saat data tidak berdistribusi normal, maka gunakan median untuk mengisi *missing value*. Pada penelitian ini, dataset berdistribusi normal yang dibuktikan dengan hasil statistik uji *Shapiro-Wilk normality test* ($W = 0.94756$) atau nilai *p-value* sebesar 0.1229.

Tabel 3. Evaluasi *Confusion Matrix* Algoritma DT

Aktual\Prediksi	5	4	3	2	1
5 (sangat rendah)	26	0	0	0	0
4 (rendah)	0	1	0	0	0
3 (menengah)	2	0	0	0	0
2 (tinggi)	0	1	0	2	1
1 (sangat tinggi)	0	1	0	1	3

Tabel 4. Evaluasi *Confusion Matrix* Algoritma NB

Aktual\Prediksi	5	4	3	2	1
5 (sangat rendah)	21	0	5	0	0
4 (rendah)	0	1	0	0	0
3 (menengah)	0	1	1	0	0
2 (tinggi)	0	0	0	3	1
1 (sangat tinggi)	0	0	0	1	4

Analisis Nilai Precision, Recall, F1-Score, Specificity

Analisis performa algoritma DT dan NB akan dilihat berdasarkan nilai precision, recall, F1-score, dan specificity yang diperoleh pada kedua algoritma. Berikut hasil perbandingan nilai pada berbagai kelas.

Tabel 5. Perbandingan Performa Algoritma DT dan NB

Nilai	Algoritma	Kelas				
		5	4	3	2	1
Precision	DT	0,928	0,333	0,000	0,667	0,750
	NB	1,000	0,500	0,166	0,750	0,800
Recall	DT	1,000	1,000	0,000	0,500	0,600
	NB	0,807	1,000	0,500	0,750	0,800
F1-Score	DT	0,962	0,499	0,000	0,571	0,666
	NB	0,893	0,666	0,249	0,750	0,800
Specificity	DT	0,833	0,944	1,000	0,970	0,969
	NB	1,000	0,972	0,861	0,970	0,969

Berdasarkan empat nilai yang dihasilkan dari kedua algoritma, berikut analisis lebih lanjut yang dapat diberikan.

Algoritma *Decision Trees*

1. Akurasi; 84,2% adalah nilai akurasi yang cukup tinggi (baik) secara keseluruhan.
2. Precision dan Recall

Kelas 5 : Presisi 0,928 dan Recall 1,0 menunjukkan model ini sangat baik dalam mendeteksi kelas 5.

Kelas 4 : Presisi 0,333 dan Recall 1,0 menunjukkan banyak False Positives, namun kelas 4 berhasil terdeteksi sepenuhnya.

Kelas 3 : Presisi dan Recall 0, menunjukkan bahwa kelas ini tidak terdeteksi sama sekali.

Kelas 2 dan 1 : Presisi dan Recall cukup baik, meskipun ada ruang untuk peningkatan.

3. F1-Score; Menunjukkan keseimbangan yang baik untuk kelas 5, tapi performa menurun pada kelas 3 dan 4.
4. Specificity; Specificity tinggi pada semua kelas, artinya model dapat mengenali kelas negatif dengan baik.

Metode Naïve Bayes

1. Akurasi; 78,9%, akurasi lebih rendah dibanding algoritma DT.
2. Precision dan Recall
 Kelas 5 : Presisi sempurna (1,0) tetapi Recall 0,807 menunjukkan ada beberapa False Negatives.
 Kelas 4 dan 3 : Performanya lebih seimbang dibanding metode lain, khususnya pada kelas 3 dengan nilai Presisi 0,166 dan Recall 0,5.
 Kelas 2 dan 1 : Nilai Presisi dan Recall baik, menunjukkan keseimbangan dalam mendeteksi dua kelas ini.
3. F1-Score; Menunjukkan keseimbangan yang lebih baik di semua kelas, dengan kelas 3 memiliki F1-Score 0,249 (meski masih rendah).
4. Specificity; Specificity tinggi di hampir semua kelas, artinya model dapat mengenali kelas negatif dengan baik.

4. Kesimpulan Dan Saran

Kesimpulan

Berdasarkan hasil dan pembahasan yang telah

dilakukan, maka kesimpulan yang dapat diberikan pada analisis kinerja algoritma Machine Learning dalam memprediksi tingkat kemiskinan global adalah:

1. Algoritma *Decision Trees* memiliki keseimbangan antara akurasi dan deteksi kelas. Performa pada kelas 5 sangat baik, tetapi model gagal mendeteksi kelas 3. Ini bisa menjadi pilihan baik jika fokus utama adalah kinerja yang baik pada sebagian besar kelas kecuali kelas 3.
2. Algoritma Naïve Bayes menunjukkan keseimbangan terbaik antara berbagai kelas. Meskipun akurasinya lebih rendah (78,9%), kelas 4 dan 3 terdeteksi meskipun dengan nilai presisi dan recall yang lebih rendah. F1-Score dan specificity juga menunjukkan bahwa model ini lebih stabil untuk semua kelas.

Saran

Penelitian menunjukkan bahwa performa ketiga algoritma memberikan gambaran yang cukup baik, namun masih terdapat kekurangan, khususnya pada penanganan kelas-kelas yang kurang terdeteksi. Beberapa saran yang bisa dilakukan kedepannya adalah penanganan ketidakseimbangan data, seperti oversampling kelas minoritas (misalnya dengan SMOTE) atau undersampling kelas mayoritas.

5. Ucapan Terima Kasih

Terimakasih kepada Program Studi Matematika Fakultas Sains dan Teknologi Universitas Islam Negeri Imam Bonjol Padang.

Daftar Pustaka

- [1] B. World, *Correcting Course*, vol. 24, no. 5. 2022.
- [2] Unicef, *On my mind*, vol. 182, no. 4. 2021.
- [3] A. Sen, "Evaluative Reason :," *Oxford: Oxford University Press*, p. 5, 1999.
- [4] U. N. D. P. UNDP and O. P. and H. D. I. OPHI, "Global MPI 2020 – Charting pathways out of multidimensional poverty: Achieving the SDGs.," *United Nations Development Programme (UNDP) and Oxford Poverty and Human Development Initiative (OPHI)*, no. July, pp. 1–52, 2020.
- [5] N. Kabeer, "Gender, poverty, and inequality: a brief history of feminist contributions in the field of international development," *Gender and Development*, vol. 23, no. 2, pp. 189–205, 2015, doi: 10.1080/13552074.2015.1062300.
- [6] Acemoglu & Robinson, "Why Nations Fail : The Origins of Power , Prosperity and Poverty (review) Review Reviewed Work (s): Why Nations Fail : The Origins of Power , Prosperity and Poverty by Daron Acemoglu and James A . Robinson Review by : HIMANSHU JHA Published by : ISEAS," *Asean Economic Bulletin*, vol. 29, no. 2, pp. 168–170, 2012.
- [7] J. L. Hanson *et al.*, "Family poverty affects the rate of human infant brain growth," *PLoS ONE*, vol. 8, no. 12, 2013, doi: 10.1371/journal.pone.0080954.
- [8] J. Luby *et al.*, "NIH Public Access Author Manuscript JAMA Pediatr. Author manuscript; available in PMC 2014 April 28. Published in final edited form as: JAMA Pediatr. 2013 December; 167(12): 1135–1142. doi:10.1001/jamapediatrics.2013.3139. The Effects of Poverty on Child," *JAMA Pediatr*, vol. 12, no. 167, pp. 1135–1142, 2013, doi: 10.1001/jamapediatrics.2013.3139.
- [9] F. A. Suhendar, R. V. Sari, T. Pangesti, Z. M. G. Putra, and A. P. A. Santoso, "The Impact of Poverty on Education," vol. 8, no. 2, pp. 39–47, 2024, doi: 10.1007/978-3-319-22807-5_4.
- [10] T. O'Brien, "The bottom billion: Why the poorest countries are failing and what can be done about it: Some insights for the Pacific?," *Economic Round-up*, vol. Winter, pp. 95–118, 2007.
- [11] D. M. Putri and Aghsilni, "Estimasi Model Terbaik Untuk Peramalan Harga Saham PT. Polychem Indonesia Tbk. dengan ARIMA," *MAP Journal: Mathematics and Applications*, vol. 1, pp. 1–12, Dec. 2019.
- [12] D. M. Putri, L. H. Hasibuan, R. A. Nur, and E. Asfa'ani, "OPTIMASI PREDIKSI CURAH HUJAN KOTA PADANG DENGAN MODEL ARIMA".
- [13] R. A. Nur and D. M. Putri, "Penerapan ARIMA Pada Data Curah Hujan di Stasiun Meteorologi Kelas II Minangkabau Padang Pariaman," *JOSTECH*, vol. 4, no. 1, pp. 77–86, Mar. 2024, doi: 10.15548/jostech.v4i1.8361.
- [14] N. Jean, M. Burke, M. Xie, W. M. Davis, D. B. Lobell, and S. Ermon, "Combining satellite imagery and machine learning to predict poverty," *Science*, vol. 353, no. 6301, pp. 790–794, 2016, doi: 10.1126/science.aaf7894.
- [15] D. Kreuzberger and N. Kühl, "Machine Learning Operations (MLOps): Overview , Definition , and Architecture," vol. 11, no. March, 2023, doi: 10.1109/ACCESS.2023.3262138.
- [16] G. S. Linoff and M. J. A. Berry, "Data mining techniques: for marketing, sales, and customer relationship management. Wiley," Third., 2011.
- [17] B. Harmadi, "PENERAPAN ALGORITMA DECISION TREE PADA PREDIKSI RISIKO TERSEKANG STROKE," 2022.
- [18] Iykra, "Mengenal Decision Tree dan Manfaatnya," Medium. [Online]. Available: <https://medium.com/iykra/mengenal-decision-tree-dan-manfaatnya-b98cf3cf6a8d>
- [19] P. L. Meyer, *Introductoy Probability and Statistical aplication*, Second. Addison Wesley, 1970.