

PENINGKATAN MODEL KLASIFIKASI PADA KAB/KOTA DI INDONESIA MENGGUNAKAN METODE DECISION TREE

Supta Danil^{1*}, Nining Rahaningsih², Raditya Danar Dana³, Mulyawan⁴

^{1,2,3} Sekolah Tinggi Manajemen Informatika dan Komputer (STMIK) IKMI Cirebon; Jl. Perjuangan No.10B, Karyamulya, Kec. Kesambi, Kota Cirebon, Jawa Barat 45135; Telp. (0231) 490480

Received: 5 Maret 2025
Accepted: 27 Maret 2025
Published: 14 April 2025

Keywords:

Decision Tree;
Klasifikasi;
Tingkat Kemiskinan;
Confusion Matrik.

Correspondent Email:

suptadaniel@gmail.com

Kemiskinan masih menjadi permasalahan signifikan di Indonesia, terutama dalam hal ketidaktepatan sasaran dalam pemerataan ekonomi. Oleh karena itu, penelitian ini bertujuan untuk mengembangkan model klasifikasi kemiskinan di tingkat kabupaten/kota di Indonesia menggunakan algoritma Decision Tree. Penelitian ini mengangkat beberapa rumusan masalah, antara lain pengembangan model klasifikasi, pengukuran performa model, dan analisis pengaruh pemilihan fitur terhadap akurasi model. Dataset yang digunakan bersumber dari Kaggle, terdiri dari 514 data dengan variabel seperti pengeluaran per kapita, Indeks Pembangunan Manusia (IPM), dan akses terhadap sanitasi layak.

Proses penelitian mencakup tahapan preprocessing data, meliputi seleksi atribut, pembersihan data, dan transformasi atribut kategorikal menjadi numerik. Model klasifikasi yang dihasilkan menunjukkan akurasi hingga 87%, dengan analisis yang menyoroti pengeluaran per kapita dan akses terhadap sanitasi sebagai faktor utama yang memengaruhi tingkat kemiskinan. Evaluasi kinerja model dilakukan menggunakan matriks kebingungan, presisi, recall, dan F1-score, yang menunjukkan performa baik dalam membedakan kategori "miskin" dan "tidak miskin".

Poverty is still a significant problem in Indonesia, especially in terms of inaccurate targeting in economic equality. Therefore, this research aims to develop a poverty classification model at the district/city level in Indonesia using the Decision Tree algorithm. This research raises several problems, including the development of a classification model, measurement of model performance, and analysis of the effect of feature selection on model accuracy. The dataset used is sourced from Kaggle, consisting of 514 data with variables such as per capita expenditure, Human Development Index (HDI), and access to proper sanitation.

The research process includes data preprocessing, including attribute selection, data cleaning, and transformation of categorical attributes into numerical ones. The resulting classification model showed an accuracy of 87%, with the analysis highlighting per capita expenditure and access to sanitation as the main factors affecting poverty levels. Evaluation of the model's performance was conducted using the confusion matrix, precision, recall, and F1-score, which showed good performance in distinguishing between the "poor" and "non-poor" categories.

1. PENDAHULUAN

Kemiskinan masih menjadi permasalahan kompleks dan multidimensional di Indonesia, meskipun angka kemiskinan telah menurun dalam beberapa dekade terakhir. Berbagai upaya penanggulangan kemiskinan telah dilakukan, namun efektivitasnya masih perlu ditingkatkan, salah satunya melalui ketepatan sasaran program. Metode *data mining*, seperti *Decision Tree*, menawarkan potensi untuk mengidentifikasi karakteristik rumah tangga miskin berdasarkan faktor-faktor yang mempengaruhinya [1]. Penelitian sebelumnya [2] menunjukkan keberhasilan penerapan *Decision Tree* dalam klasifikasi kemiskinan di Kabupaten Sleman. Akan tetapi, masih ada ruang untuk peningkatan akurasi dan efektivitas model dengan mempertimbangkan pemilihan atribut, penanganan data yang tidak seimbang, dan penggunaan algoritma *ensemble learning* [3]. Penelitian ini bertujuan untuk mengembangkan model klasifikasi kemiskinan di tingkat kabupaten/kota di Indonesia menggunakan algoritma *Decision Tree* yang lebih akurat, efektif, dan mudah diinterpretasi. Model ini diharapkan dapat mengidentifikasi faktor-faktor kunci penyebab kemiskinan dan memberikan informasi berharga bagi perumusan kebijakan penanggulangan kemiskinan yang tepat sasaran, yang selama ini cenderung fokus pada analisis di tingkat provinsi atau regional [4]. Selain itu, pemanfaatan algoritma *Decision Tree* diharapkan memberikan keunggulan dalam hal interpretasi dan visualisasi, sehingga faktor-faktor dominan dalam menentukan kemiskinan dapat diidentifikasi dengan lebih mudah. Penelitian ini diharapkan dapat memberikan kontribusi dalam upaya penanggulangan kemiskinan di Indonesia melalui penyediaan model klasifikasi yang akurat dan mudah diinterpretasi.

2. TINJAUAN PUSTAKA

2.1. Konsep Kemiskinan

Kemiskinan didefinisikan sebagai kondisi kekurangan multidimensional yang ditandai dengan rendahnya pendapatan, akses terbatas terhadap layanan dasar, dan kerentanan terhadap guncangan. Pengukuran kemiskinan dapat dilakukan melalui pendekatan kuantitatif, seperti penghitungan garis kemiskinan

berdasarkan pengeluaran per kapita, maupun pendekatan kualitatif yang mempertimbangkan aspek sosial dan budaya [5].

2.2. Data Mining dan Klasifikasi

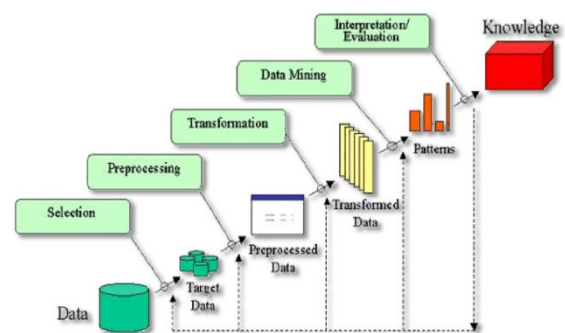
Data mining adalah proses penemuan pola dan informasi berharga dari sejumlah besar data [1]. Klasifikasi merupakan salah satu tugas utama dalam *data mining*, yang bertujuan untuk mengelompokkan data ke dalam kategori yang telah ditentukan.

2.3. Algoritma Decision Tree

Decision Tree adalah algoritma klasifikasi yang populer karena kemampuannya dalam menangani data dengan berbagai tipe atribut dan menghasilkan model yang mudah diinterpretasi [6]. Algoritma ini bekerja dengan membangun struktur pohon yang membagi data berdasarkan atribut-atribut yang paling relevan, hingga mencapai simpul daun yang mewakili kategori klasifikasi. Keunggulan *Decision Tree* terletak pada kemudahan visualisasi dan interpretasi model, serta kemampuan untuk menangani data kategorikal dan numerik. Beberapa penelitian telah menunjukkan efektivitas *Decision Tree* dalam klasifikasi kemiskinan.

3. METODE PENELITIAN

3.1. Teknik Analisis Data



Gambar 1 KDD (Knowledge Discovery in Database)

Penelitian ini menggunakan metode Knowledge Discovery in Database (KDD) dengan tahapan sebagai berikut:

a. Seleksi Data (Data Selection):

Dalam penelitian ini, seleksi data melibatkan pemilihan atribut yang relevan untuk analisis aktivitas dan tingkat kemiskinan kab/kota di Indonesia. Contohnya, Pengeluaran per

Kapita Disesuaikan (Ribuan Rupiah/Orang/Tahun) dan PDRB atas Dasar Harga Konstan menurut Pengeluaran (Rupiah), serta 10 variabel kategorik, yaitu Provinsi, Kab/Kota, Persentase Penduduk Miskin (P0) Menurut Kabupaten/Kota (Persen), Rata-rata Lama Sekolah Penduduk 15+ (Tahun), Indeks Pembangunan Manusia, Umur Harapan Hidup (Tahun), Persentase rumah tangga yang memiliki akses terhadap sanitasi layak, Persentase rumah tangga yang memiliki akses terhadap air minum layak, Tingkat Pengangguran Terbuka, dan Tingkat Partisipasi Angkatan Kerja sebagai atribut yang menjadi fokus penelitian.

- b. Pembersihan Data (Data Cleaning): Pada tahap pembersihan data, dilakukan eliminasi data yang tidak konsisten atau tidak relevan. Misalnya, data-data yang memiliki kesalahan atau format yang tidak sesuai dihapus dari dataset, sehingga dataset menjadi bersih dan siap untuk proses selanjutnya.
- c. Transformasi Data (Data Transformation): Proses transformasi data melibatkan konversi atribut kategorikal menjadi bentuk numerik dan menghapus kolom non-numerik "Provinsi" dan "Kab/Kota".
- d. Data Mining: Tahap data mining pada penelitian ini mencakup penggunaan algoritma *Decision Tree* untuk mengklasifikasi tingkat kemiskinan di Indonesia berdasarkan miskin dan tidak miskin.
- e. Evaluasi: Dalam konteks penelitian, evaluasi dilakukan secara berkelanjutan untuk menilai validitas data mining. Misalnya, menggunakan metrik evaluasi seperti *accuracy*, *precision*, *recall*, dan F1-score untuk menentukan seberapa baik model *Decision Tree* dapat mengklasifikasi tingkat kemiskinan kab/kota di Indonesia.

Confusion matrix adalah sebuah tabel yang digunakan dalam machine learning untuk mengevaluasi kinerja model klasifikasi. Tabel ini memberikan visualisasi detail tentang seberapa baik model dalam memprediksi kelas-kelas yang berbeda dalam dataset.

Dengan confusion matrix, dapat dilihat dengan jelas:

- a. True Positive (TP) : Jumlah kasus yang benar-benar miskin dan diprediksi dengan benar oleh model sebagai miskin.
- b. True Negative (TN) : Jumlah kasus yang benar-benar tidak miskin dan diprediksi dengan benar oleh model sebagai tidak miskin.
- c. False Positive (FP) : Jumlah kasus yang sebenarnya tidak miskin tetapi diprediksi salah oleh model sebagai miskin.
- d. False Negative (FN)² : Jumlah kasus yang sebenarnya miskin tetapi diprediksi salah oleh model sebagai tidak miskin.

Dalam konteks penelitian tentang tingkat kemiskinan, confusion matrix membantu kita untuk memahami seberapa baik model dalam mengidentifikasi daerah miskin dan tidak miskin. Ini penting untuk menentukan sasaran program penanggulangan kemiskinan secara lebih tepat dan efisien.

Berikut adalah code yang akan menghasilkan confusion metrik dengan model klasifikasi *Decision Tree*.

4. HASIL DAN PEMBAHASAN

4.1. Data Selection

Pada tahap ini, data yang digunakan berasal dari dataset "Klasifikasi Tingkat Kemiskinan di Indonesia" yang diunduh dari Kaggle. Dataset terdiri dari 514 data dengan atribut seperti pengeluaran per kapita, IPM, akses sanitasi, dan tingkat pengangguran. Proses seleksi data bertujuan untuk memilih atribut yang relevan. Atribut "Provinsi" dan "Kabupaten/Kota" dihapus karena tidak relevan dalam analisis. Atribut seperti *persentase rumah tangga dengan akses sanitasi* dan *pengeluaran per kapita* dipertahankan karena relevansinya yang

3.2. Confusion Matrix

tinggi terhadap tingkat kemiskinan (Nugroho, 2022).

```
# Hapus kolom non-numerik 'Provinsi' dan 'Kab/Kota'
df = df.drop(['Provinsi', 'Kab/Kota'], axis=1)
df.head()
```

Gambar 2 Data Selection

4.2. Data Cleaning

Pada gambar 3 data diperiksa untuk menemukan nilai yang kosong atau tidak valid. Baris yang memiliki nilai kosong dihapus untuk memastikan kualitas data yang digunakan dalam analisis. Tahap ini penting untuk menghindari bias atau error dalam hasil klasifikasi.

```
# Hapus baris dengan nilai NaN di kolom target ('Klasifikasi Kemiskinan')
df = df.dropna(subset=['Klasifikasi Kemiskinan'])
df.head()
```

Gambar 3 Data Cleaning

4.3. Transformation Data

Transformasi data merupakan tahapan penting dalam *preprocessing* data yang bertujuan untuk mengubah format atau struktur data agar sesuai dengan kebutuhan algoritma yang akan digunakan. Dalam penelitian ini, transformasi data dilakukan untuk mempersiapkan data sebelum membangun model klasifikasi kemiskinan menggunakan algoritma *Decision Tree*.

Transformasi data ini juga mencakup normalisasi nilai untuk meningkatkan performa algoritma.

```
# Buat DataFrame untuk menyimpan tipe data sebelum transformasi
before_transform = pd.DataFrame(df.dtypes, columns=['Sebelum Transformasi'])

# Transformasi Data: Mengubah tipe data atribut menjadi numerik
for column in df.columns:
    if df[column].dtype == 'object':
        df[column] = pd.to_numeric(df[column], errors='coerce')

# Buat DataFrame untuk menyimpan tipe data sesudah transformasi
after_transform = pd.DataFrame(df.dtypes, columns=['Sesudah Transformasi'])

# Gabungkan kedua DataFrame
tipe_data = pd.concat([before_transform, after_transform], axis=1)

# Atur index menjadi nama atribut
tipe_data.index.name = 'Atribut'

# Tampilkan tabel
print(tipe_data.to_markdown(numalign="left", stralign="left"))
```

Gambar 4 Transformation Data

Gambar 4 menunjukkan proses transformasi data yang dilakukan sebelum membangun model klasifikasi kemiskinan. Proses ini bertujuan untuk mempersiapkan data agar sesuai dengan kebutuhan algoritma *Decision Tree*.

Perubahan utama yang dilakukan adalah mengubah tipe data beberapa kolom dari tipe 'object' menjadi 'float64'. Tipe data 'object' biasanya digunakan untuk data teks atau data kategorikal, sedangkan 'float64' digunakan

untuk data numerik. Algoritma *Decision Tree* bekerja dengan data numerik, sehingga transformasi ini diperlukan agar data dapat diproses oleh algoritma.

4.4. Data Mining

Splitting data adalah proses membagi dataset menjadi dua atau lebih subset. Ini adalah langkah penting dalam machine learning karena memungkinkan untuk melatih, mengevaluasi, dan menguji model secara terpisah dengan data yang berbeda, sehingga bisa mendapatkan gambaran yang lebih akurat tentang kinerja model.

```
# Pisahkan data menjadi fitur dan target
X = df.drop('Klasifikasi Kemiskinan', axis=1)
# Menggunakan semua kolom kecuali 'Klasifikasi Kemiskinan' sebagai fitur
y = df['Klasifikasi Kemiskinan']

# Bagi dataset menjadi data training dan testing (80:20)
X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2, random_state=42)
print("Data berhasil dibagi menjadi training dan testing.")
```

Data berhasil dibagi menjadi training dan testing.

Gambar 5 Data Mining

Gambar 5 menunjukkan code Algoritma *Decision Tree* digunakan untuk membangun model klasifikasi. Data dibagi menjadi 80% data latih dan 20% data uji untuk mengevaluasi kemampuan generalisasi model.

Dengan *splitting data* yang tepat, performa model diuji pada data uji, memberikan akurasi sebesar 87%. Ini menunjukkan model memiliki performa yang baik dalam mengklasifikasikan tingkat kemiskinan [12].

4.5. Evaluation

Evaluation adalah proses mengukur kinerja suatu model machine learning. Evaluasi ini bertujuan untuk mengetahui seberapa baik model dapat melakukan tugasnya, seperti membuat prediksi atau klasifikasi, pada data baru yang belum pernah dilihat sebelumnya.

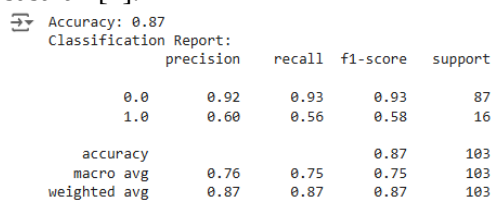
```
# Evaluasi model menggunakan accuracy, F1 score
accuracy = accuracy_score(y_test, y_pred)
print(f"Accuracy: {accuracy:.2f}")
print("Classification Report:")
print(classification_report(y_test, y_pred))
```

Gambar 6 Evaluation

Gambar 6 menunjukkan kode program untuk mengevaluasi kinerja model *Decision Tree* dalam memprediksi tingkat kemiskinan. Evaluasi dilakukan untuk mengukur kemampuan model mengklasifikasikan data baru yang belum pernah dilihat sebelumnya. Dua metrik utama yang dihitung adalah akurasi dan laporan klasifikasi. Akurasi menunjukkan persentase prediksi yang benar, dengan nilai antara 0 hingga 1, di mana 1 menunjukkan akurasi sempurna. Laporan klasifikasi

memberikan rincian kinerja untuk setiap kelas, yaitu "Tidak Miskin (0)" dan "Miskin (1)", mencakup presisi, recall, F1-score, dan support. Presisi mengukur ketepatan prediksi positif, recall menilai kemampuan model mendeteksi data positif, F1-score memberikan keseimbangan antara presisi dan recall, dan support menunjukkan jumlah data yang dianalisis.

Evaluasi kinerja model sangat penting untuk memahami kelemahan dan kekuatan model dalam memberikan hasil yang tepat sasaran [2].



Accuracy: 0.87				
Classification Report:				
	precision	recall	f1-score	support
0.0	0.92	0.93	0.93	87
1.0	0.60	0.56	0.58	16
accuracy			0.87	103
macro avg	0.76	0.75	0.75	103
weighted avg	0.87	0.87	0.87	103

Gambar 7 Hasil Evaluation

Gambar 7 hasil evaluasi model Decision Tree menunjukkan akurasi sebesar 0,87 dengan rincian laporan klasifikasi yang menyortir presisi, recall, F1-score, dan support untuk kelas "Tidak Miskin (0)" dan "Miskin (1)". Evaluasi ini relevan dengan yang membahas optimasi algoritma Decision Tree untuk meningkatkan akurasi klasifikasi [10], serta penelitian [11] yang menerapkan algoritma Decision Tree untuk analisis tingkat kemiskinan di Indonesia. Kinerja model Decision Tree yang terukur pada setiap kelas memberikan wawasan penting, sebagaimana [7], dalam membandingkan efektivitas model berbasis algoritma decision tree pada data kemiskinan. Sumber lain, mendukung pentingnya laporan klasifikasi sebagai alat untuk memahami lebih dalam performa model pada setiap kelas. Evaluasi ini dilakukan menggunakan pendekatan yang sejalan dengan metode yang dipaparkan oleh [10], yang menilai pengaruh karakteristik data dalam klasifikasi status kemiskinan di wilayah Indonesia Timur.

4.6. Confusion Matrix

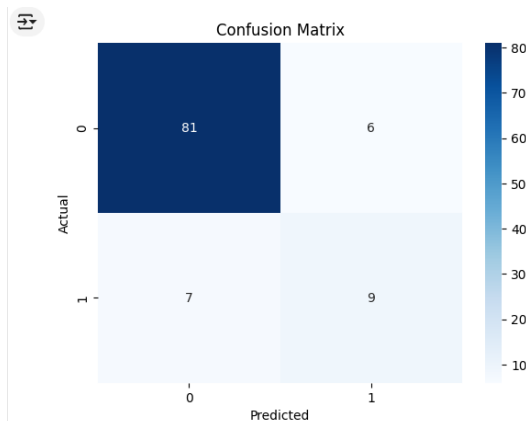
Confusion matrix yang dihasilkan dalam penelitian ini menunjukkan bahwa model Decision Tree memiliki performa yang baik dalam mengklasifikasikan data ke dalam kategori "Miskin" dan "Tidak Miskin". Berdasarkan tabel confusion matrix, terdapat 81

kasus *True Negative* (data sebenarnya "Tidak Miskin" dan diprediksi dengan benar), 6 kasus *False Positive* (data sebenarnya "Tidak Miskin" tetapi salah diprediksi sebagai "Miskin"), 7 kasus *False Negative* (data sebenarnya "Miskin" tetapi salah diprediksi sebagai "Tidak Miskin"), dan 9 kasus *True Positive* (data sebenarnya "Miskin" dan diprediksi dengan benar).

Table 1 Confusion Matrix

Actual / Predicted	0 (Tidak Miskin)	1 (Miskin)
0 (Tidak Miskin)	81	6
1 (Miskin)	7	9

Dari hasil ini, terlihat bahwa model memiliki kemampuan yang sangat baik dalam mengenali kategori "Tidak Miskin", tetapi masih memiliki kelemahan dalam mengklasifikasikan individu atau rumah tangga miskin, seperti yang tercermin dari nilai *False Negative* yang lebih tinggi dibandingkan *False Positive*. Kelemahan ini dapat memengaruhi efektivitas model dalam membantu pemerintah menargetkan bantuan sosial kepada kelompok yang benar-benar membutuhkan. Decision Tree sering kali lebih akurat dalam mengklasifikasikan kelas mayoritas dibandingkan kelas minoritas, terutama jika dataset tidak seimbang [9]. menekankan pentingnya evaluasi *confusion matrix* dalam analisis klasifikasi untuk memahami kesalahan prediksi yang terjadi. Dengan menggunakan evaluasi mendalam seperti ini, penyesuaian model, seperti optimasi parameter atau penggunaan algoritma ensemble seperti *Random Forest*, dapat dilakukan untuk memperbaiki kelemahan model [8]. yang menunjukkan bahwa penggunaan splitting criteria yang lebih optimal pada algoritma pohon keputusan dapat meningkatkan performa klasifikasi, termasuk mengurangi kesalahan prediksi seperti *False Negative*.



Gambar 8 Hasil Confusion Matrix

Dengan demikian, *confusion matrix* tidak hanya memberikan gambaran akurasi model, tetapi juga menjadi alat penting untuk mengevaluasi dan memperbaiki model klasifikasi kemiskinan, sehingga dapat memberikan hasil yang lebih tepat sasaran.

5. KESIMPULAN

- a. Penelitian ini berhasil mengimplementasikan algoritma Decision Tree untuk memprediksi tingkat kemiskinan dengan akurasi sebesar 87%. Evaluasi model menunjukkan performa yang baik untuk kelas "Tidak Miskin", dengan presisi 92%, recall 93%, dan F1-score 93%. Namun, kinerja pada kelas "Miskin" masih perlu ditingkatkan karena presisi, recall, dan F1-score masing-masing hanya mencapai 60%, 56%, dan 58%. Temuan ini mengindikasikan bahwa model lebih unggul dalam mengklasifikasikan data dengan distribusi yang lebih dominan (kelas mayoritas), tetapi mengalami kesulitan pada kelas minoritas. Hal ini menjadi tantangan utama dalam aplikasi model Decision Tree pada data yang tidak seimbang.
- b. Untuk penelitian selanjutnya, disarankan untuk mengeksplorasi teknik penanganan data tidak seimbang, seperti oversampling menggunakan metode Synthetic Minority Oversampling Technique (SMOTE) atau undersampling. Selain itu, pengoptimalan model

dapat dilakukan dengan menerapkan teknik seperti hyperparameter tuning menggunakan grid search atau metode optimasi lainnya. Riset lanjutan juga dapat mengeksplorasi algoritma lain, seperti Random Forest atau Gradient Boosting, untuk dibandingkan dengan kinerja Decision Tree. Terakhir, perluasan dataset dengan fitur tambahan yang relevan dapat meningkatkan akurasi prediksi dan memberikan pemahaman lebih mendalam terkait faktor yang memengaruhi kemiskinan.

UCAPAN TERIMA KASIH

Saya dengan tulus menyampaikan rasa terima kasih kepada Ermila atas dukungan luar biasa yang diberikan dalam bentuk data primer, laporan resmi, dan informasi lapangan yang menjadi landasan utama penelitian ini. Dukungan ini sangat membantu kami dalam mengidentifikasi pola tingkat Kemiskinan di Indonesia. Kami juga berterima kasih kepada STMIK IKMI Cirebon atas penyediaan fasilitas penelitian, termasuk perangkat lunak dan ruang diskusi, yang mendukung proses analisis data menggunakan algoritma Decision Tree. Apresiasi setinggi-tingginya kami sampaikan kepada rekan-rekan sejawat yang telah memberikan masukan konstruktif selama tahapan praproses data, transformasi, dan evaluasi hasil klasterisasi. Tidak lupa, kami menghargai peran berbagai pihak yang mendukung di balik layar, termasuk kolega yang memberikan saran praktis terkait penerapan metode Knowledge Discovery in Databases (KDD) dan keluarga yang senantiasa memberikan motivasi dan semangat selama penelitian berlangsung. Semua bentuk kontribusi ini menjadi fondasi keberhasilan penelitian ini, yang kami harapkan dapat memberikan manfaat nyata bagi pengelolaan sektor Ekonomi untuk lebih tepat sasaran dalam menyalurkan bantuan dan menjadi referensi dalam pengembangan kebijakan pertanian yang lebih efisien dan terarah di masa depan.

DAFTAR PUSTAKA

- [1] Kaunang, F. (2019). Penerapan algoritma j48 decision tree untuk analisis tingkat kemiskinan di indonesia. *Cogito Smart Journal*, 4(2), 348-357.
<https://doi.org/10.31154/cogito.v4i2.141.348-357>
- [2] Esananda, S., Nugroho, B., & Anggraeny, F. (2021). Implementasi fase boosting pada algoritma c5.0 dalam menentukan prestasi akademik siswa. *Prosiding Seminar Nasional Informatika Bela Negara*, 2, 1-6.
<https://doi.org/10.33005/santika.v2i0.67>
- [3] Nugroho, A. (2022). Analisa splitting criteria pada decision tree dan random forest untuk klasifikasi evaluasi kendaraan. *Jsitik Jurnal Sistem Informasi Dan Teknologi Informasi Komputer*, 1(1), 41-49.
<https://doi.org/10.53624/jsitik.v1i1.154>
- [4] Bahaiddin, A., Fatmawati, A., & Sari, F. (2021). Analisis clustering provinsi di indonesia berdasarkan tingkat kemiskinan menggunakan algoritma k-means. *Jurnal Manajemen Informatika Dan Sistem Informasi*, 4(1), 1-8.
<https://doi.org/10.36595/misi.v4i1.216>
- [5] A. Rohmatullah, D. Rahmalia, and M. S. Pradana, "Klasterisasi Data Pertanian di Kabupaten Lamongan Menggunakan Algoritma K-Means Dan Fuzzy C-Means," *J. Ilm. Teknosains*, vol. V, no. 2, pp. 86–93, 2019.
- [6] Finaliamartha, D., Supriyadi, D., & Fitriana, G. (2022). Penerapan metode jaringan syaraf tiruan backpropagation untuk prediksi tingkat kemiskinan di provinsi jawa tengah. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 9(4), 751-760.
<https://doi.org/10.25126/jtiik.2022934806>
- [7] Rahmadeni, R. and Nurjannah, N. (2022). Model tingkat kemiskinan di kabupaten/kota provinsi riau: menggunakan regresi data panel. *Kubik Jurnal Publikasi Ilmiah Matematika*, 6(2), 98-109.
<https://doi.org/10.15575/kubik.v6i2.13598>
- [8] Rachma, C. A. (2022). *Implementasi Algoritma K-Neraest Neighbor dalam penentuan Klasifikasi Tingkat Kedalaman Kemiskinan Provinsi Jawa Timur* (Doctoral dissertation, Universitas Islam Negeri Maulana Malik Ibrahim).
- [9] Taufiq, N. (2022). Penciri kemiskinan ekstrem di 35 kabupaten prioritas penanganan kemiskinan ekstrem. *Seminar Nasional Official Statistics*, 2022(1), 895-904.
<https://doi.org/10.34123/semnasoffstat.v2022i1.1258>
- [10] Amida, O. and Sitorus, J. (2021). Penerapan regresi logistik biner multilevel dalam analisis pengaruh karakteristik individu, rumah tangga, dan wilayah terhadap status kemiskinan balita di kepulauan maluku dan pulau papua. *Seminar Nasional Official Statistics*, 2020(1), 967-977.
<https://doi.org/10.34123/semnasoffstat.v2020i1.569>
- [11] Arifin, T. (2020). Optimasi decision tree menggunakan particle swarm optimization untuk klasifikasi sel pap smear. *Jatisi (Jurnal Teknik Informatika Dan Sistem Informasi)*, 7(3), 572-579.
<https://doi.org/10.35957/jatisi.v7i3.361>
- [12] S. S. Elfaretta, A. A. Arifiyanti, and A. S. Fitri, "Klasifikasi Calon Pendorong Darah Potensial Menggunakan Algoritma Decision Tree Di Utd Pmi Kota Surabaya," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3, 2024, doi: 10.23960/jitet.v12i3.4957.