

ANALISIS KLASIFIKASI RISIKO GEMPA BUMI DI INDONESIA BERDASARKAN MAGNITUDO DAN KEDALAMAN MENGUNAKAN ALGORITMA *DECISION TREE*

Salsabila Alfath Annisaa', Novita Ramadhini,
Adiaz Salsabilah Putri, Ken Ditha Tania, Allsela Meiriza

Sistem Informasi, Universitas Sriwijaya

Jl. Raya Palembang – Prabumulih No.KM. 32, Indralaya Indah, Kec.

Indaralaya, Kabupaten Ogan Ilir, Sumatera Selatan, Indonesia

salsabialfatha@gmail.com

ABSTRAK

Indonesia merupakan negara dengan aktivitas kegempaan sangat tinggi akibat letaknya di pertemuan tiga lempeng tektonik besar. Tingginya potensi ancaman ini menimbulkan permasalahan terkait perlunya analisis risiko gempa yang cepat dan akurat untuk mendukung sistem mitigasi bencana. Oleh karena itu, penelitian ini bertujuan untuk mengklasifikasikan tingkat risiko gempa bumi di Indonesia berdasarkan parameter magnitudo dan kedalaman. Metode yang digunakan adalah algoritma *Decision Tree* terhadap 131.833 data gempa dari BMKG periode 2008–2025. Tahapan penelitian ini meliputi *data collection*, *preprocessing*, *labelling*, *splitting*, *modelling*, *rule extraction*, dan *evaluation*. Hasil klasifikasi menunjukkan bahwa magnitudo $> 4,501$ termasuk kategori Risiko Sedang (315 data), sedangkan magnitudo $\leq 4,501$ termasuk Risiko Rendah (10.726 data). Model ini menghasilkan akurasi sebesar 99,96% dan *recall* 100%. Tingginya performa ini membuktikan bahwa algoritma *Decision Tree* sangat efektif untuk klasifikasi risiko gempa dan berpotensi mendukung upaya mitigasi bencana nasional.

Kata kunci : Klasifikasi, Risiko Gempa Bumi, Data Mining, *Decision Tree*.

1. PENDAHULUAN

Indonesia dikenal sebagai salah satu negara dengan aktivitas kegempaan paling tinggi di dunia [1]. Hal ini disebabkan posisinya yang berada di pertemuan tiga lempeng tektonik besar, yaitu Indo-Australia, Eurasia, dan Pasifik [1]. Setiap tahun, ribuan gempa tercatat oleh Badan Meteorologi, Klimatologi, dan Geofisika (BMKG), dengan berbagai tingkat magnitudo, kedalaman, dan lokasi [2]. Fenomena ini umumnya dipicu oleh proses subduksi maupun pergerakan sesar yang masih aktif [3]. Tingginya frekuensi gempa ini tidak hanya menjadi fenomena geologis, tetapi juga memunculkan dampak serius mulai dari kerugian materi, kerusakan bangunan dan fasilitas umum, hingga memakan korban jiwa [4]. Karena itu, memahami bagaimana karakteristik gempa di Indonesia menjadi hal yang sangat penting sebagai upaya mengurangi risiko bencana dan meningkatkan kesiapsiagaan masyarakat [5].

Meskipun data gempa bumi di Indonesia telah direkam secara sistematis oleh BMKG, pemanfaatan data tersebut untuk penilaian tingkat risiko belum sepenuhnya optimal [6]. Sebagian besar analisis masih bersifat deskriptif dan belum memanfaatkan pendekatan analitik modern seperti machine learning untuk mengklasifikasikan tingkat kerawanan berdasarkan parameter seismik seperti magnitudo dan kedalaman [7]. Kondisi ini menghambat pemerintah dan masyarakat dalam memperoleh informasi risiko yang terstruktur dan mudah dipahami untuk mendukung perencanaan mitigasi bencana [7]. Oleh karena itu, penting untuk memiliki sistem pemetaan

dan klasifikasi risiko gempa yang akurat agar bisa mendukung upaya mitigasi dan penanggulangan darurat secara efektif [8].

Berbagai penelitian sebelumnya telah memanfaatkan metode *machine learning* dalam analisis gempa bumi di Indonesia. Beberapa studi menerapkan model klasifikasi untuk memprediksi magnitudo maupun tingkat bahaya gempa dengan menggunakan algoritma seperti *Random Forest*, *XGBoost*, dan *LightGBM*. Temuan tersebut menunjukkan bahwa model-model tersebut, khususnya *LightGBM* dan *Random Forest*, mampu memberikan performa yang baik dalam analisis seismologi dan prediksi magnitudo [9]. Namun demikian, penelitian lainnya memperlihatkan bahwa sebagian besar studi yang menggunakan algoritma *K-Means* masih terbatas pada proses klusterisasi dan belum menilai tingkat risiko setiap kejadian gempa, serta belum mengaitkan hasil tersebut dengan tingkat risiko yang dapat digunakan dalam upaya mitigasi bencana [5]. Temuan ini memang membantu untuk melihat pola kejadian gempa, namun belum cukup untuk memberikan gambaran tentang tingkat kerawannya sehingga belum bisa langsung dimanfaatkan dalam pengambilan keputusan. Kondisi ini menunjukkan adanya *research gap*, karena penelitian sebelumnya lebih menekankan pada proses klusterisasi dan belum mengembangkan metode klasifikasi yang mampu menentukan tingkat risiko berdasarkan parameter penting seperti magnitudo dan kedalaman. Selain itu, pemanfaatan algoritma *machine*

learning untuk kategori risiko yang jelas dan mudah dipahami juga masih sangat terbatas.

Berdasarkan latar belakang dan permasalahan tersebut, terlihat bahwa setiap algoritma memiliki keunggulan dan keterbatasan masing-masing. Oleh karena itu, penelitian ini bertujuan untuk mengklasifikasikan tingkat risiko gempa bumi di Indonesia berdasarkan parameter magnitudo dan kedalaman. Sebagai rencana penyelesaian masalah, penelitian ini mengusulkan penggunaan algoritma *Decision Tree* sebagai metode klasifikasi. *Decision Tree* dipilih karena mampu menghasilkan model yang mudah dibaca, transparan dalam proses pengambilan keputusan, serta efektif dalam mengelompokkan data ke dalam kategori tertentu. Algoritma ini juga terbukti mampu mengidentifikasi pola penting pada data gempa historis secara konsisten dan efisien [10]. Dengan rencana penyelesaian ini, informasi yang dihasilkan diharapkan lebih mudah dipahami dan dapat dimanfaatkan langsung dalam mendukung mitigasi serta perencanaan penanggulangan bencana.

2. TINJAUAN PUSTAKA

Penelitian terkait analisis risiko gempa bumi di Indonesia menjadi topik yang penting karena Indonesia berada di dekat beberapa lempeng tektonik aktif dan memiliki tingkat aktivitas gempa yang tinggi sehingga memiliki banyak aktivitas gempa. Untuk memahami tingkat risiko yang ditimbulkan oleh gempa bumi, analisis parameter diperlukan. Data dari Badan Meteorologi, Klimatologi, dan Geofisika (BMKG) menunjukkan bahwa gempa bumi yang terjadi di Indonesia memiliki magnitudo dan kedalaman yang berbeda [2]. Beberapa penelitian sebelumnya telah menggunakan teknik *data mining* dan *machine learning* dalam menganalisis data gempa bumi. Penelitian dalam lainnya menunjukkan bahwa *clustering* dapat digunakan untuk mengelompokkan data gempa berdasarkan fitur tertentu [10]. Ini dapat membantu memahami pola distribusi gempa di suatu area. Sama hal juga didalam penelitian [5] menggunakan algoritma *K-Means* untuk mengelompokkan data gempa berdasarkan kedalaman, yang menunjukkan bahwa faktor kedalaman sangat penting untuk menganalisis karakteristik gempa bumi. Meskipun metode *clustering* dapat menemukan pola pengelompokan data, metode tersebut belum menghasilkan klasifikasi tingkat risiko yang jelas secara langsung.

Jumlah penelitian yang memanfaatkan metode *machine learning* untuk menganalisis data gempa bumi secara lebih mendalam telah meningkat dalam beberapa tahun terakhir sebagai akibat dari kemajuan metode ini. Untuk menghasilkan model analisis yang lebih akurat, pendekatan prediktif berbasis *machine learning* mulai digunakan untuk memanfaatkan parameter gempa seperti magnitudo, kedalaman, dan lokasi geografis. Untuk memprediksi magnitudo gempa, penelitian tersebut menggunakan berbagai algoritma pengajaran mesin, seperti *Random Forest*,

Support Vector Regression, *XGBoost*, *LightGBM*, dan *Multi-Layer Perceptron*. Hasil penelitian menunjukkan bahwa parameter kedalaman dan geolokasi dapat digunakan secara efektif dalam proses pemodelan untuk menghasilkan tingkat prediksi yang cukup baik [9]. Sejalan dengan perkembangan tersebut, penelitian lain telah menggunakan algoritma *Decision Tree* untuk mengklasifikasikan tingkat risiko gempa di Indonesia dengan memanfaatkan karakteristik spasial dan temporal dari data gempa bumi. Metode ini menunjukkan bahwa hubungan antar parameter gempa dapat dianalisis secara lebih sistematis melalui proses pembentukan *Decision Tree*, juga dikenal sebagai pohon keputusan. Hasil penelitian tersebut menunjukkan bahwa algoritma *Decision Tree* mampu meningkatkan tingkat keamanan gempa di Indonesia [10].

Berdasarkan uraian tersebut, penelitian ini melakukan analisis klasifikasi risiko gempa bumi di Indonesia dengan menggunakan algoritma *Decision Tree*. Parameter magnitudo dan kedalaman dipilih karena merupakan indikator penting dalam menentukan tingkat kekuatan serta potensi dampak gempa bumi. Dengan menggunakan algoritma *Decision Tree*, diharapkan penelitian ini dapat menghasilkan model klasifikasi yang mampu memberikan gambaran yang lebih baik tentang bagaimana gempa bumi dapat berdampak pada masyarakat Indonesia.

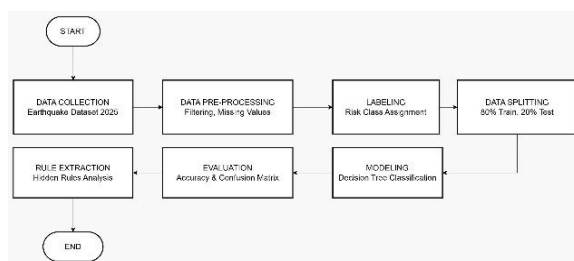
Dalam melakukan analisis risiko gempa bumi tahap awal dimulai dari *data collection* melalui data resmi gempa, yang dimana data ini meliputi magnitudo, kedalaman, dan tempat terjadinya gempa. Setelah data diperoleh, tahap *preprocessing* data dilakukan mencakup pembersihan data, penyaringan atribut yang relevan, serta penanganan nilai yang hilang (*missing values*) agar data menjadi lebih terstruktur saat dianalisis [11]. Setelah proses *labelling* selesai, data kemudian dibagi menjadi dua bagian, 80% data latihan (*training data*) dan 20% data uji (*testing data*). Tujuan dari pembagian data ini adalah untuk menguji kemampuan model yang dibangun untuk melakukan klasifikasi pada data yang belum pernah dipelajari sebelumnya. Selanjutnya, proses *modeling* dilakukan dengan menggunakan algoritma *Decision Tree* untuk membangun model klasifikasi tingkat resiko gempa.

Kemudian, tahap *evaluation* dilakukan untuk menilai kinerja model dalam klasifikasi dengan menggunakan metrik seperti *accuracy* dan *confusion matrix*. Setelah tahap *evaluation* selesai, model dianalisis lebih lanjut melalui tahap *rule extraction* aturan untuk mendapatkan aturan keputusan yang terbentuk dari struktur *Decision Tree*. Dengan demikian, hubungan antar parameter gempa dalam menentukan tingkat risiko menjadi lebih jelas.

3. METODE PENELITIAN

Pada Metodologi penelitian merupakan tahapan yang menjelaskan mengenai gambaran umum dari

metode serta tahapan proses klasifikasi risiko gempa bumi yang akan digunakan dalam penelitian. Proses penelitian dilakukan dengan memanfaatkan algoritma *Decision Tree* sebagai metode klasifikasi melalui pendekatan *data mining* untuk mengidentifikasi pola dan tingkat risiko gempa [11]. Proses ini dilakukan melalui tahapan *data mining* yang sistematis, mulai dari tahap pengumpulan data, pra-pemrosesan, penambangan data (*modelling*), hingga evaluasi pola dan presentasi pengetahuan [12]. Setiap tahap dalam metodologi ini dirancang secara terstruktur menggunakan parameter magnitudo dan kedalaman agar mampu menghasilkan model prediksi risiko gempa yang akurat, relevan, dan dapat diinterpretasikan. Alur keseluruhan dari proses penelitian ini dapat dilihat pada Gambar 1.



Gambar 1. Diagram alur proses

Dari gambar 1, dapat dilihat dapat bahwa alur penelitian ini diawali dengan tahap *data collection* untuk mengumpulkan data mentah, yang kemudian dilanjutkan dengan tahap *data preprocessing* yang mencakup proses *filtering* dan penanganan *missing values* guna menghasilkan data yang bersih. Setelah data siap, dilakukan tahap *labelling* melalui *risk class assignment* sebelum data masuk ke tahap *data splitting* untuk membagi dataset menjadi 80% data *train* dan 20% data *test*. Proses dilanjutkan dengan Modeling menggunakan *Decision Tree Classification* untuk membentuk model klasifikasi yang akurat berdasarkan pola data yang telah diproses. Rangkaian alur proses ini diakhiri dengan tahap tahap *evaluation* menggunakan *accuracy* dan *confusion matrix* yang diikuti oleh *rule extraction* untuk melakukan analisis aturan tersembunyi (*hidden rules analysis*).

3.1. Data Collection

Tahap pertama dari penelitian ini adalah pengumpulan data mentah yang akan diolah untuk memperoleh informasi yang relevan sesuai kebutuhan analisis. Pada penelitian ini, pengumpulan data dilakukan dengan mengambil dataset sekunder berupa katalog gempa bumi di Indonesia dari situs Kaggle yang bersumber dari BMKG. Data yang digunakan memuat catatan kejadian gempa bumi yang secara spesifik difokuskan pada periode tahun 2025. Atribut utama yang diekstrak dari dataset ini untuk mendukung proses analisis adalah parameter magnitudo dan kedalaman (*depth*) [11].

3.2. Data Preprocessing

Tahap ini bertujuan untuk melakukan pra-pemrosesan pada data mentah yang sering kali memiliki *noise*, ketidaksempurnaan, atau atribut yang tidak relevan. Proses pembersihan data merupakan langkah penting untuk mengidentifikasi dan menangani nilai yang hilang atau tidak konsisten agar keandalan data tetap terjaga [11]. Tahapan pra-pemrosesan yang dilakukan dalam penelitian ini meliputi penyaringan (*filtering*) data khusus untuk kejadian tahun 2025, seleksi atribut yang hanya mempertahankan kolom magnitudo dan kedalaman, serta pembersihan baris data yang kosong (*missing values handling*) agar dataset siap dan optimal untuk pemrosesan lebih lanjut [12].

3.3. Labelling

Tahap pelabelan bertujuan untuk menciptakan atau mendefinisikan kelas target pada sekumpulan data yang belum memiliki klasifikasi tingkat bahaya secara resmi. Pelabelan dilakukan secara kondisional menggunakan fungsi operator pada perangkat lunak dengan membuat atribut baru bernama 'Label_Risiko'. Atribut target ini dibentuk dengan menetapkan aturan ambang batas kekuatan gempa untuk mengkategorikan setiap baris data kejadian menjadi tiga kelas utama, yaitu Risiko Tinggi, Risiko Sedang, dan Risiko Rendah [11].

3.4. Data Splitting

Data splitting merupakan tahapan penting untuk memecah dataset yang telah berlabel menjadi dua proporsi bagian yang berbeda sebelum masuk ke tahap pemodelan. Pada penelitian ini, pemecahan data dilakukan dengan menerapkan rasio pembagian 80:20. Sebanyak 80% dari total dataset digunakan sebagai data latih (*training data*) untuk melatih model pada sebagian besar data, sementara 20% sisanya dialokasikan sebagai data uji (*testing data*) untuk mengevaluasi kemampuan model dalam memprediksi data yang belum pernah digunakan pada proses pelatihan [13].

3.5. Modelling

Tahap pemodelan dilakukan dengan menerapkan algoritma *Decision Tree* pada sekumpulan data latih untuk melakukan proses klasifikasi. Algoritma ini bekerja dengan membangun model berdasarkan variabel-variabel yang telah ditentukan, seperti magnitudo dan kedalaman, untuk mengklasifikasikan data ke dalam kategori risiko yang sesuai [13]. Hasil dari tahapan ini adalah terbentuknya sebuah model klasifikasi berupa struktur pohon keputusan yang merepresentasikan logika prediksi risiko gempa bumi [11].

3.6. Rule Extraction

Ekstraksi aturan (*rule extraction*) adalah proses mengonversi visualisasi bentuk pohon keputusan yang dihasilkan pada tahap *modelling* menjadi pernyataan

logika kondisional (*IF-THEN*) yang dapat dibaca secara langsung. Tahap ini memungkinkan peneliti untuk mengungkap pola tersembunyi dari fenomena gempa bumi dan menemukan angka ambang batas desimal spesifik yang diidentifikasi oleh algoritma dalam memisahkan masing-masing tingkat risiko [11].

3.7. Evaluation

Tahapan evaluasi merupakan proses akhir yang digunakan untuk mengukur dan memvalidasi seberapa baik kinerja model klasifikasi saat dihadapkan pada data uji. Evaluasi dalam penelitian ini memanfaatkan metode *Confusion Matrix* yang berisi informasi perbandingan detail antara hasil prediksi yang dilakukan oleh sistem kecerdasan buatan dengan data aktualnya [11]. Melalui *Confusion Matrix*, peneliti dapat memperoleh informasi mengenai nilai akurasi (*accuracy*), *precision*, *recall*, dan *F1-score* yang dihasilkan model untuk menganalisis keandalannya [13]. Beberapa Rumus dari *Accuracy*, *Precision*, dan *Recall* sebagai berikut :

- Accuracy*

$$\frac{(TP+TN)}{TP+TN+FP+FN} \times 100\% \quad (1)$$
- Precision*

$$\frac{TP}{TP+FP} \times 100\% \quad (2)$$
- Recall*

$$\frac{TP}{TP+FN} \times 100\% \quad (3)$$

4. HASIL DAN PEMBAHASAN

Memuat Penelitian ini menggunakan dataset yang berjumlah 131,833 data gempa bumi yang diambil dari website dataset Kaggle.

4.1. Data Collection

Dataset yang didapatkan berasal dari repositori milik BMKG yang dapat diakses secara publik (https://www.kaggle.com/datasets/kekavigi/earthquakes-in-indonesia?select=katalog_gempa_v2.tsv) dengan jumlah record sebanyak 131,833 dimulai dari tahun 2008-2025 dan atribut yang dimiliki terdiri atas *event_id*, *date_time*, *latitude*, *longitude*, *magnitude*, *mag_type*, *depth*, *phasecount*, *azimuth_gap*, *location*, seperti yang ditampilkan pada gambar 2.

Row No.	eventID	datetime	latitude	longitude	magnitude	mag_type	depth	ph
1	bmg2008vlye	Nov 1, 2008 ...	-0.604	98.896	2.990	MLv	20	6
2	bmg2008vlag	Nov 1, 2008 ...	-6.612	129.387	5.508	mb	30	62
3	bmg2008vlaj	Nov 1, 2008 ...	-3.651	127.991	3.540	MLv	5	4
4	bmg2008vlbt	Nov 1, 2008 ...	-4.199	128.097	2.424	MLv	5	5
5	bmg2008vlcd	Nov 1, 2008 ...	-4.092	128.200	2.410	MLv	10	5
6	bmg2008vlcw	Nov 1, 2008 ...	-3.755	127.382	2.939	MLv	10	4
7	bmg2008vlcf	Nov 1, 2008 ...	-3.889	128.243	2.219	MLv	10	4
8	bmg2008vlfr	Nov 1, 2008 ...	0.487	98.328	3.853	MLv	16	7
9	bmg2008vlja	Nov 1, 2008 ...	-4.264	127.662	3.243	MLv	10	5
10	bmg2008vlji	Nov 1, 2008 ...	-3.433	128.389	2.317	MLv	10	4
11	bmg2008vlkj	Nov 1, 2008 ...	-3.942	127.447	2.594	MLv	9	6
12	bmg2008vlkr	Nov 1, 2008 ...	-4.430	127.447	3.150	MLv	10	6
13	bmg2008vlks	Nov 1, 2008 ...	-3.375	128.016	2.786	MLv	10	6

Gambar 2. Dataset gempa bumi Indonesia

4.2. Data Preprocessing

Tahap pra-pemrosesan data dimulai dengan membaca dataset menggunakan operator Read CSV.

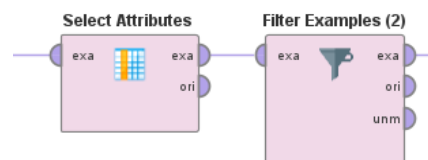
Karena penelitian ini difokuskan pada kejadian gempa terkini, dilakukan penyaringan data khusus untuk tahun 2025 menggunakan operator *Filter Examples*. Hasil dari proses *filtering* ini mereduksi dataset menjadi 13.970 data kejadian gempa. Selanjutnya, operator *Select Attributes* diterapkan untuk menyeleksi variabel-variabel yang relevan dengan pemodelan, dimana atribut *magnitude* dan kedalaman (*depth*) dipilih sebagai variabel prediktor utama. Untuk memastikan kualitas data, digunakan kembali operator *Filter Examples* (2) guna menangani *missing value* (data kosong) sehingga data yang akan diproses selanjutnya benar-benar bersih.



Gambar 3. Baca file CSV

Row No.	eventID	datetime	latitude	longitude	magnitude	mag_type	depth	ph
1	bmg2025easr	Jan 1, 2025 1...	2.058	126.953	3.276	MLv	12	20
2	bmg2025eaba	Jan 1, 2025 1...	0.318	121.791	2.564	MLv	178	22
3	bmg2025eac7	Jan 1, 2025 1...	-8.791	129.062	3.983	MLv	276	22
4	bmg2025eadl	Jan 1, 2025 1...	-8.005	117.968	4.092	MLv	11	88
5	bmg2025easer	Jan 1, 2025 2...	-8.165	116.411	2.922	MLv	11	30
6	bmg2025eash	Jan 1, 2025 3...	-10.156	116.756	3.044	MLv	7	30
7	bmg2025easb	Jan 1, 2025 6...	-9.726	119.473	2.430	MLv	28	15
8	bmg2025easo	Jan 1, 2025 6...	2.316	99.008	2.917	MLv	136	31
9	bmg2025easg	Jan 1, 2025 7...	-0.212	100.370	1.918	MLv	12	16
10	bmg2025easr	Jan 1, 2025 1...	-8.542	118.397	3.154	MLv	111	53
11	bmg2025easa	Jan 1, 2025 1...	1.559	126.351	3.327	MLv	21	29
12	bmg2025easb	Jan 1, 2025 1...	-2.230	138.958	3.786	MLv	36	16
13	bmg2025eash	Jan 1, 2025 3...	6.780	126.406	3.645	MLv	11	16

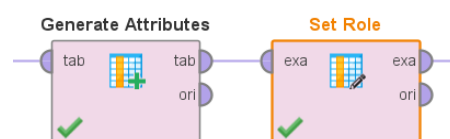
Gambar 4. Dataset 2025 setelah *filtering*



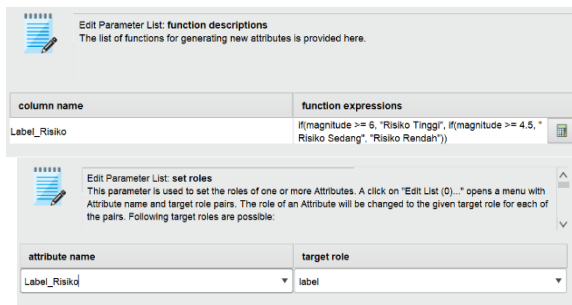
Gambar 5. Tahapan data preprocessing

4.3. Labelling

Setelah data dibersihkan, langkah berikutnya adalah menetapkan kelas target. Pada tahapan ini, ditambahkan operator *Generate Attributes* untuk menciptakan atribut baru berupa kelas target yang dinamakan *Label_Risiko*. Setelah atribut tersebut dibuat, operator *Set Role* digunakan untuk mengubah peran atribut *Label_Risiko* menjadi label, sehingga algoritma *machine learning* dapat mengenalinya sebagai variabel yang akan diprediksi, seperti yang ditampilkan pada gambar 6 dan gambar 7.



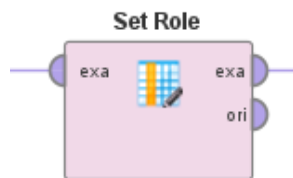
Gambar 6. Tahapan labelling



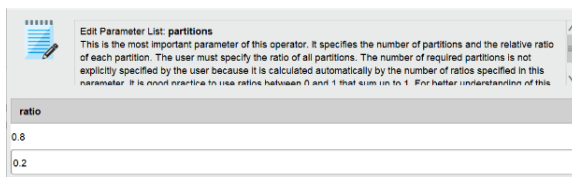
Gambar 7. Tahapan input parameter labelling

4.4. Data Splitting

Untuk menguji keandalan model yang akan dibangun, dataset berjumlah 13.970 record tersebut dibagi menjadi dua bagian menggunakan operator *Split Data*. Rasio pembagian yang digunakan adalah 80% sebagai *Training Data* (data latih) untuk membangun model, dan 20% sebagai *Testing Data* (data uji) untuk mengevaluasi kinerja model yang telah dibentuk, seperti yang ditampilkan pada gambar 8 dan gambar 9.



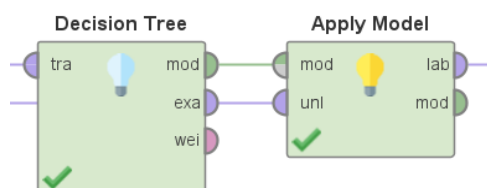
Gambar 8. Tahapan data splitting



Gambar 9. Tahapan input parameter data splitting

4.5. Modelling

Tahap pemodelan (*modelling*) dilakukan dengan menerapkan algoritma *Decision Tree*. Operator *Decision Tree* dihubungkan dengan data latih (*training data*) dari hasil proses *Split Data*. Setelah pohon keputusan terbentuk dari proses pembelajaran data latih, model tersebut kemudian diterapkan pada data uji menggunakan operator *Apply Model* untuk memprediksi label risiko gempa bumi berdasarkan magnitudo dan kedalaman, seperti yang ditampilkan pada gambar 10.



Gambar 10. Proses Decision Tree

4.6. Evaluation

Kinerja model algoritma *Decision Tree* dievaluasi menggunakan operator *Performance* untuk melihat seberapa baik model dapat mengklasifikasikan data. Berdasarkan hasil pengujian (*PerformanceVector*), model *Decision Tree* menghasilkan tingkat akurasi yang sangat tinggi, yaitu sebesar 99,96% dapat terlihat pada gambar 11 dan gambar 12.

The image shows the 'PerformanceVector (Performance)' window. It displays the accuracy as 99.96% and a confusion matrix table.

	True Risiko Rendah	True Risiko Sedang	True Risiko Tinggi	Class
pred. Risiko Rendah	2681	0	0	100.00%
pred. Risiko Sedang	0	79	1	98.75%
pred. Risiko Tinggi	0	0	0	0.00%
Class Recall	100.00%	100.00%	0.00%	

Gambar 11. Hasil confusion matrix

The image shows the 'PerformanceVector' window. It displays the accuracy as 99.96% and a confusion matrix table.

True:	Risiko Rendah	Risiko Sedang	Risiko Tinggi
Risiko Rendah:	2681	0	0
Risiko Sedang:	0	79	1
Risiko Tinggi:	0	0	0

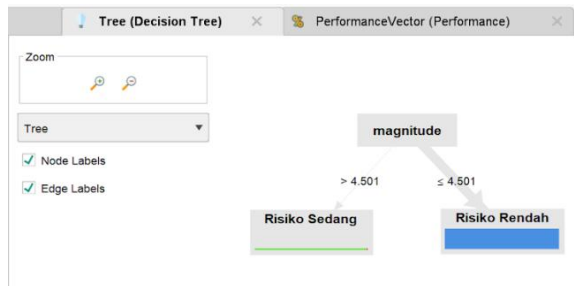
Gambar 12. Deskripsi hasil confusion matrix

Rincian klasifikasi dari evaluasi tersebut dapat dilihat melalui *Confusion Matrix* berikut:

- True Risiko Rendah: Dari data yang aslinya berisiko rendah, model memprediksi 2.681 data dengan benar sebagai Risiko Rendah, 0 sebagai Risiko Sedang, dan 0 sebagai Risiko Tinggi. (Nilai *Class Recall* 100%).
- True Risiko Sedang: Dari data yang aslinya berisiko sedang, model memprediksi 79 data dengan benar sebagai Risiko Sedang, namun terdapat 1 data yang meleset diprediksi sebagai Risiko Tinggi, dan 0 sebagai Risiko Rendah. (Nilai *Class Recall* 100%).
- True Risiko Tinggi: Model memprediksi 0 untuk seluruh kemungkinan pada kategori ini.

4.7. Rule Extraction

Hasil visualisasi algoritma *Decision Tree* tidak hanya memberikan prediksi kelas, melainkan juga mengekstraksi aturan (*hidden rules*) yang dapat dipahami untuk analisis ambang batas (*threshold*). Pembagian akar keputusan bergantung penuh pada atribut magnitudo, dapat dilihat pada gambar 13 dan gambar 14.

Gambar 13. Hasil visualisasi *Decision Tree*Gambar 14. Deskripsi hasil *Decision Tree*

Berdasarkan *tree* yang terbentuk, ditemukan aturan klasifikasi sebagai berikut:

- Jika *magnitude* > 4.501, maka gempa diklasifikasikan ke dalam Risiko Sedang. Pada simpul ini, distribusi datanya adalah: Risiko Rendah = 0, Risiko Sedang = 315, Risiko Tinggi = 4.
- Jika *magnitude* ≤ 4.501, maka gempa diklasifikasikan ke dalam Risiko Rendah. Pada simpul ini, distribusi datanya dominan pada kelas yang aman: Risiko Rendah = 10.726, Risiko Sedang = 0, Risiko Tinggi = 0

Hal ini menyimpulkan bahwa nilai magnitudo sebesar 4.501 menjadi titik potong (*threshold*) utama yang sangat membedakan klasifikasi potensi risiko kerusakan gempa dalam dataset tahun 2025.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian dan pengujian yang telah dilakukan menggunakan algoritma *Decision Tree* pada dataset gempa bumi Indonesia, dapat disimpulkan bahwa model klasifikasi ini memiliki performa yang sangat tinggi dengan tingkat akurasi mencapai 99,96%. Nilai *recall* dan *precision* untuk kategori Risiko Rendah menunjukkan angka sempurna sebesar 100%. Pada kategori Risiko Sedang, nilai *precision* mencapai 100%, sedangkan *recall* sebesar 98,75% karena terdapat satu data Risiko Sedang yang terklasifikasi sebagai Risiko Tinggi oleh sistem. Melalui proses pemodelan ini, sistem secara otomatis menentukan variabel *magnitude* sebagai faktor penentu utama dengan ambang batas sebesar 4,501. Dari hasil ekstraksi aturan, ditemukan bahwa gempa dengan *magnitude* > 4,501 dikategorikan sebagai Risiko Sedang dengan total sebanyak 315 sampel data, sementara gempa dengan *magnitude* ≤ 4,501 dikategorikan sebagai Risiko Rendah dengan frekuensi yang jauh lebih dominan yaitu sebanyak 10.726 sampel data. Secara keseluruhan, penggunaan metode *Decision Tree* ini terbukti sangat efektif dalam memetakan tingkat kerawanan gempa secara cepat dan akurat.

Meskipun demikian, perlu diperhatikan adanya ketidakseimbangan kelas pada data penelitian, terutama pada kategori Risiko Tinggi yang jumlahnya sangat terbatas, sehingga pengujian lanjutan dan penyeimbangan data disarankan untuk meningkatkan keandalan model dalam konteks mitigasi serta perencanaan tanggap darurat bencana di Indonesia. Namun demikian, hasil ini tetap perlu divalidasi lebih lanjut menggunakan data yang lebih seimbang agar generalisasi model dapat dipastikan.

DAFTAR PUSTAKA

- [1] A. Prasetyo and M. M. Effendi, "Analisis gempa bumi di Indonesia dengan metode clustering," *Bull. Inf. Technol.*, vol. 4, no. 3, pp. 338–343, 2023.
- [2] BMKG. (2025). *Buletin Gempabumi dan Tsunami Tahun 2024* (Publikasi Februari 2025). Stasiun Geofisika Lampung.
- [3] Y. Selvina, "PEMETAAN DAN KLASIFIKASI WILAYAH RAWAN GEMPA DI INDONESIA DENGAN METODE K-MEANS DAN LIGHTGBM," *J. Inform. dan Tek. Elektro Terap.*, vol. 13, no. 3S1, 2025.
- [4] BNPB. (2023). *Laporan Penanggulangan Bencana Tahun 2023: Evaluasi Dampak Gempabumi*
- [5] M. Ubaidillah and Z. Fatah, "Implementasi RapidMiner pada Klasterisasi Gempa Bumi di Indonesia Berdasarkan Kedalaman Menggunakan K-Means," *J. Ilm. Multidisiplin Ilmu*, vol. 1, no. 6, pp. 84–91, 2024.
- [6] A. Nurusyifa, M. Valeri, and A. S. Rahman, "PEMETAAN INDEKS BAHAYA GEMPA BUMI DAN PEMBUATAN SHAKEMAP GEMPA BUMI DKI JAKARTA," *Bul. Meteorol. Klimatologi dan Geofis.*, vol. 3, no. 6, pp. 8–19, 2023.
- [7] R. D. W. I. DHARMAWAN, J. Sartohadi, and T. H. Purwanto, "Pemodelan Risiko Gempabumi Multiskenario di Daerah Istimewa Yogyakarta Menggunakan Deep Learning," 2023, *Universitas Gadjah Mada*.
- [8] S. H. Alavi, A. Bahrami, M. Mashayekhi, and M. Zolfaghari, "Optimizing interpolation methods and point distances for accurate earthquake hazard mapping," *Buildings*, vol. 14, no. 6, p. 1823, 2024.
- [9] I. Maulita and A. M. Wahid, "Prediksi Magnitudo Gempa Menggunakan Random Forest, Support Vector Regression, XGBoost, LightGBM, dan Multi-Layer Perceptron Berdasarkan Data Kedalaman dan Geolokasi," *J. Pendidik. Dan Teknol. Indones.*, vol. 4, no. 5, pp. 221–232, 2024.
- [10] M. Prasetyo, H. Sulistiani, O. Y. Inonu, K. Magda, and B. Santosa, "Klasifikasi Tingkat Risiko Gempa di Indonesia Menggunakan Pola Spasial dan Temporal Berbasis Decision Tree," *Bull. Comput. Sci. Res.*, vol. 5, no. 5, pp. 1059–

- 1066, 2025.
- [11] N. Akbar, H. Alghifari, N. Abdillah, and O. Dahwanu, "Klasifikasi Risiko Gempa Bumi menggunakan metode Decision Tree," *DEVICE J. Inf. Syst. Comput. Sci. Inf. Technol.*, vol. 6, no. 2, pp. 683–694, 2025.
- [12] Alwendi, A. Mandopa, and L. Budiarti, "APLIKASI DATA MINING UNTUK MENENTUKAN LULUSAN MAHASISWA TEPAT WAKTU," *J. Adv. Res. Informatics*, vol. 3, no. 1, pp. 59–65, 2024.
- [13] Irawan, A., & Ningsih, D. H. U. (2026). IMPLEMENTASI ALGORITMA DECISION TREE DALAM SISTEM PENDUKUNG KEPUTUSAN SELEKSI CALON KARYAWAN PADA PT G-TEXTILE INDONESIA MENGGUNAKAN FRAMEWORK DJANGO PYTHON. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 10(1), 1302–1308.