

Analisis Perbandingan Varian Naïve Bayes dalam Klasifikasi Kelayakan Penerima BLT-DD Di Desa Taraju

Neng Sri Lathifah Zulfa^{1*}, Iffah Athifah²

^{1,2}Teknik, Informatika, Universitas Islam Al-Ihya Kuningan, Kuningan, Indonesia

E-mail: ^{1*}srilathifah@gmail.com, ²iffahathifah3@gmail.com

(* : corresponding author)

Abstrak

Program BLT-DD bertujuan memberikan bantuan tunai kepada masyarakat desa yang berada dalam kondisi ekonomi kurang mampu. Agar penyaluran bantuan tepat sasaran, diperlukan mekanisme klasifikasi yang efektif dalam menentukan kelayakan penerima. Penelitian ini membandingkan kinerja tiga varian algoritma Naïve Bayes yaitu *Gaussian*, *Bernoulli*, dan *Complement* dalam klasifikasi kelayakan penerima BLT-DD di Desa Taraju. Setiap algoritma diuji dengan teknik *10-fold cross-validation* untuk memastikan keandalan evaluasi. Berdasarkan hasil pengujian, algoritma *Bernoulli Naïve Bayes* mencatat akurasi tertinggi sebesar 91%, diikuti *Gaussian Naïve Bayes* dengan 90%, dan *Complement Naïve Bayes* sebesar 64%. Temuan ini menunjukkan bahwa pendekatan *Bernoulli Naïve Bayes* lebih optimal dalam konteks klasifikasi data BLT-DD. Kebaruan dari penelitian ini terletak pada penerapan perbandingan sistematis tiga varian *Naïve Bayes* untuk klasifikasi kelayakan bantuan sosial berbasis data desa, yang masih minim eksplorasi dalam studi-studi sebelumnya.

Kata kunci: Bantuan Langsung Tunai, K-fold cross validation, Klasifikasi, Naïve Bayes

Abstract

The BLT-DD program aims to provide direct cash assistance to rural communities in economically disadvantaged conditions. To ensure accurate targeting, an effective classification mechanism is required to determine beneficiary eligibility. This study compares the performance of three Naïve Bayes algorithm variants—*Gaussian*, *Bernoulli*, and *Complement*—in classifying BLT-DD recipients in Taraju Village. Each algorithm was evaluated using *10-fold cross-validation* to ensure reliable assessment. The results show that the *Bernoulli Naïve Bayes* algorithm achieved the highest accuracy at 91%, followed by *Gaussian Naïve Bayes* at 90%, and *Complement Naïve Bayes* at 64%. These findings indicate that the *Bernoulli Naïve Bayes* approach is more optimal for BLT-DD data classification. The novelty of this research lies in the systematic comparison of three Naïve Bayes variants for classifying eligibility in village-based social assistance programs, an area still underexplored in previous studies.

Keywords: Classification, Direct Cash Assistance, K-fold cross-validation, Naïve Bayes

1. PENDAHULUAN

Di banyak negara, termasuk Indonesia, kemiskinan tetap menjadi masalah besar yang bersifat kompleks secara global. Badan Pusat Statistik (BPS) mencatat bahwa kondisi ini tercatat pada bulan Maret 2024, sebanyak 25,87 juta jiwa atau 9,03% dari populasi Indonesia tergolong sebagai penduduk miskin [1]. Salah satu wilayah yang memiliki tingkat kemiskinan cukup tinggi adalah Kabupaten Kuningan, Provinsi Jawa Barat. Berdasarkan data BPS Kabupaten Kuningan, per Maret 2024, persentase penduduk miskin adalah 11,88% atau sekitar 131.830 jiwa, turun dari 12,12% di tahun sebelumnya. Kondisi ini belum mencerminkan pemerataan kesejahteraan. Hal ini ditunjukkan oleh meningkatnya Indeks Kedalaman Kemiskinan (P1) dari 1,87 menjadi 2,02 dan Indeks Keparahan Kemiskinan (P2) dari 0,42 menjadi 0,53, yang mengindikasikan memburuknya situasi kemiskinan secara struktural [2]. Fakta ini mengindikasikan adanya kemiskinan struktural yang tidak cukup ditangani hanya dengan mengurangi jumlah penduduk miskin, tetapi juga memerlukan intervensi yang menysasar kualitas hidup dan kesetaraan akses. Program BLT-DD diluncurkan oleh pemerintah sebagai strategi untuk mendukung kesejahteraan masyarakat miskin dan berpenghasilan rendah [3].

Desa Taraju terletak di Kuningan, Jawa Barat yang terbagi ke dalam 4 Rukun Warga (RW) dan dihuni oleh total 4.152 jiwa yang tersebar dalam 1.334 kepala keluarga. Pemilihan Desa Taraju sebagai objek penelitian didasarkan pada fakta bahwa desa tersebut telah menerapkan

program BLT-DD. Namun, implementasi program tersebut masih belum berjalan secara maksimal. Permasalahan utama yang dihadapi adalah kesulitan dalam menentukan warga yang benar-benar layak menerima bantuan. Temuan wawancara mengindikasikan bahwa keterbatasan alokasi bantuan, kurangnya pembaruan data, serta mekanisme seleksi yang mengandalkan musyawarah desa yang dalam beberapa kasus bersifat subjektif mengakibatkan penyaluran bantuan tidak sepenuhnya sesuai dengan kriteria yang ditetapkan.

Untuk menjawab tantangan tersebut, teknologi *data mining* dapat digunakan sebagai pendekatan yang lebih objektif dan terukur, dengan cara menggali pola serta informasi tersembunyi dari data berukuran besar [4]. Salah satu teknik yang relevan adalah klasifikasi, yaitu metode pembelajaran terawasi untuk mengelompokkan data ke dalam kelas tertentu berdasarkan atribut dari data berlabel [5]. Klasifikasi *Naïve Bayes* menggunakan rumus *Bayes* untuk memprediksi probabilitas suatu kelas dengan menggabungkan informasi sebelumnya dan data baru yang dianalisis [6].

Beberapa penelitian sebelumnya menunjukkan efektivitas algoritma ini dalam klasifikasi penerima bantuan sosial. Darussalam et al. mencatat akurasi 96,47%, *precision* 100,00%, *recall* 94,44% [7]. Penelitian lain oleh Huriah dan Nuris memperoleh akurasi 95,43%, *precision* 97,87%, *recall* 93,88% [8]. Di sisi lain, Putri et al. yang menggunakan *K-Nearest Neighbor* hanya memperoleh akurasi 87,31% [9], yang masih lebih rendah dibandingkan hasil menggunakan *Naïve Bayes*.

Meskipun demikian, sebagian besar studi tersebut hanya menggunakan satu jenis algoritma, sehingga belum memberikan gambaran menyeluruh terhadap efektivitas varian-varian *Naïve Bayes* dalam menangani data dengan karakteristik berbeda. Penelitian ini berupaya mengisi celah tersebut dengan membandingkan tiga varian *Naïve Bayes* yaitu *Gaussian*, *Bernoulli*, dan *Complement* yang masing-masing memiliki keunggulan terhadap tipe data tertentu. Proses evaluasi dilakukan melalui teknik *10-fold cross-validation* untuk menghasilkan model yang lebih stabil dan dapat digeneralisasi secara lebih baik. Pendekatan ini diharapkan dapat berkontribusi dalam penyusunan sistem seleksi penerima BLT-DD yang lebih akurat, objektif, dan transparan.

2. METODE PENELITIAN

Pendekatan *mixed methods* digunakan dalam penelitian ini, menggabungkan metode kualitatif dan kuantitatif guna memahami fenomena secara lebih mendalam untuk memperoleh pemahaman menyeluruh terhadap permasalahan [10]. Penelitian kualitatif adalah studi yang menggambarkan suatu peristiwa secara menyeluruh dan mendalam, dengan data dikumpulkan melalui studi dokumen, pengamatan langsung, dan wawancara [11], sedangkan pendekatan kuantitatif mengacu pada pemanfaatan data numerik secara sistematis guna menjawab rumusan masalah dan memahami fenomena sosial yang diteliti [12] menggunakan algoritma *Naïve Bayes*. Fokus utama pendekatan ini adalah melakukan analisis data dan mengukur performa klasifikasi berdasarkan metrik evaluatif, yakni *accuracy*, *precision*, *recall*, dan *F1-score*.

Objek kajian dalam penelitian ini adalah data penduduk Desa Taraju, Kecamatan Sindangagung, Kabupaten Kuningan tahun 2024 sebanyak 1.334 kepala keluarga, dengan memanfaatkan data dimana data dikumpulkan dari dokumen atau arsip resmi milik kantor desa yang memuat informasi kependudukan serta kondisi sosial ekonomi masyarakat. Data ini mencakup berbagai atribut terkait yang dirangkum dalam Tabel 1.

Tabel 1. Variabel Dalam Penelitian

No	Variabel	Keterangan
1.	No	Data angka yang menjelaskan urutan dan jumlah data.
2.	Nama Kepala Keluarga	Nama kepala keluarga.
3.	Jenis Kelamin	Informasi jenis kelamin kepala keluarga.
4.	Usia	Usia kepala keluarga.

5.	Tanggungan Keluarga	Menjelaskan jumlah tanggungan keluarga.
6.	Pendidikan	Pendidikan terakhir kepala keluarga.
7.	Status Ekonomi	Data angka yang menjelaskan pendapatan perbulan kepala keluarga.
8.	Pekerjaan	Informasi jenis pekerjaan kepala keluarga.
9.	Kepemilikan Aset	Data angka yang menjelaskan jumlah aset yang dimiliki oleh keluarga.
10.	Mempunyai Penyandang Disabilitas	Informasi apakah dalam keluarga tersebut ada penyandang disabilitas.
11.	Alamat	Informasi alamat tempat tinggal.

Setelah data terkumpul, dilakukan serangkaian proses pengolahan data untuk mendukung klasifikasi menggunakan algoritma *Naïve Bayes*. Setiap tahap dilakukan secara sistematis untuk memastikan kualitas data dan akurasi model yang dibangun.

2.1 Pembersihan Data (*Cleaning Data*)

Data diperiksa dari nilai kosong, data ganda, atau ketidakkonsistenan. Proses ini dilakukan menggunakan skrip *Python* untuk meningkatkan efisiensi dan akurasi.

2.2 Integrasi Data (*Data Integration*)

Data diperoleh dari berbagai sumber seperti arsip desa, data RT/RW, dan catatan penerima bantuan tahun sebelumnya, lalu digabungkan menjadi satu basis data utama.

2.3 Seleksi Data (*Data Selection*)

Dalam proses klasifikasi, tidak semua data penduduk digunakan. Beberapa variabel seperti No, Nama Kepala Keluarga, dan Alamat dihapus karena bersifat unik dan tidak relevan. Variabel target adalah status penerima bantuan yang diklasifikasikan menjadi dua kelas, yaitu Layak dan Tidak Layak, berdasarkan data penerimaan bantuan pada tahun 2021, 2023, dan 2024.

2.4 Evaluasi Atribut dengan *Information Gain*

Untuk mengevaluasi hubungan antara atribut (variabel bebas) dan target, digunakan metode *Information Gain*, yang mengukur pengaruh masing-masing atribut terhadap keputusan klasifikasi. Perhitungan detail nilai *gain* untuk setiap variabel dilihat pada persamaan (1) [13], dan hasil perhitungan di tampilkan pada Tabel 2.

$$Gain(y,A) = Entropy(y) - \sum_{c \in Vals(A)} \frac{y_c}{y} entropy(y_c)$$

(1)

Tabel 2 menampilkan nilai *gain* dari setiap variabel yang diperoleh melalui perhitungan.

Tabel 2. Perhitungan Gain

Variabel	Gain
Jenis Kelamin	0.0104
Usia	0.0271
Tanggungan Keluarga	0.0054
Pendidikan	0.0155
Status Ekonomi	0.0190
Pekerjaan	0.0090
Kepemilikan Aset	0.0177
Mempunyai Penyandang Disabilitas	0.0027

Berdasarkan Tabel 2, variabel usia memiliki nilai gain tertinggi sebesar 0,0271, sedangkan nilai *gain* terendah terdapat pada variabel mempunyai penyandang disabilitas. Variabel dengan gain mendekati nol seperti disabilitas, tanggungan keluarga, dan pekerjaan dieliminasi dari proses

klasifikasi. Dengan demikian, penelitian ini hanya memanfaatkan lima variabel utama, yaitu jenis kelamin, usia, pendidikan, status ekonomi, dan kepemilikan aset.

2.5 Transformasi Data (*Data Transformation*)

Proses transformasi dilakukan secara bertahap untuk mengubah data kategorikal (berbentuk teks) menjadi format numerik yang dapat diproses oleh algoritma *Naïve Bayes*, yang hanya menerima input angka. Langkah pertama adalah mengidentifikasi variabel-variabel kategorikal dalam dataset. Selanjutnya, dilakukan pengkodean dengan memberikan angka unik pada setiap kategori. Untuk variabel dengan urutan logis seperti pendidikan dan status ekonomi, digunakan pengkodean ordinal agar tetap mencerminkan tingkatannya. Transformasi variabel kategorikal dari bentuk teks ke bentuk numerik ditampilkan pada Tabel 3.

Tabel 3. Transformasi Variabel Kategorikal dengan Skema Pengkodean

Variabel	Sebelum	Sesudah
Jenis Kelamin	Laki-Laki	0
	Perempuan	1
Pendidikan	Tidak/Belum Sekolah	0
	SD/Sederajat	1
	SLTP/Sederajat	2
	SLTA/Sederajat	3
	Diploma I/II/III/Strata I	4
	Strata II	5
Status Ekonomi	Miskin	0
	Menengah	1
	Kaya	2
Pekerjaan	Tidak Bekerja	0
	Pensiunan	1
	Buruh	2
	IRT	3
	Supir	4
	Pedagang	5
	Wiraswasta	6
	PNS	7
	Karyawan Swasta	8
	Honorar	9
	Guru Honor	10
	Dosen	11
	Karyawan BUMN	12
	Karywn Swasta	13
	Polisi	14
Kepemilikan Aset	<Rp. 500.000	0
	>Rp. 500.000	1
Status	Layak	0
	Tidak Layak	1

Langkah selanjutnya adalah menerapkan skema pengkodean ke dalam dataset, di mana seluruh nilai teks diubah menjadi angka sesuai kategori masing-masing. Transformasi ini

dilakukan secara otomatis menggunakan skrip *Python* untuk meningkatkan efisiensi dan meminimalkan kesalahan manual. Setelah itu, dilakukan validasi dengan memeriksa nilai unik di setiap kolom guna memastikan tidak ada kesalahan atau kategori yang terlewat. Data yang telah valid kemudian disimpan dalam format akhir dan siap digunakan untuk tahap selanjutnya, yaitu pembagian data dan pembuatan model klasifikasi. Dengan demikian, seluruh data telah berada dalam format numerik dan dapat langsung diproses oleh algoritma *Naïve Bayes*. Gambar 1 memperlihatkan hasil transformasi data dalam bentuk numerik setelah seluruh proses di atas dilakukan.

	JENIS KELAMIN	USIA	TANGGAPAN KELUARGA	PENDIDIKAN	STATUS EKONOMI	PEKERJAAN	KEPERILIKAN ASET	PENDAPATAN PENYANDANG DISABILITAS	STATUS
0	0	32	4	2	0	2	1	0	1
1	1	61	1	1	0	0	0	0	0
2	0	57	3	1	0	2	1	0	1

Gambar 1. Hasil Transformasi Data

2.6 Pembagian Data

Sebanyak 1.334 data kepala keluarga dibagi secara acak menjadi 80% (1.067 data) dan 20% untuk pengujian (267 data). Data pelatihan digunakan sebagai dasar pembuatan model, sedangkan data pengujian digunakan untuk memvalidasi hasil model.

2.7 Penambahan Data (*Data Mining*)

Pada tahap ini, algoritma *Naïve Bayes* digunakan untuk mengklasifikasikan kelayakan penerima bantuan BLT-DD berdasarkan atribut yang telah ditentukan. Penelitian mengimplementasikan tiga varian *Naïve Bayes* dan menguji masing-masing model pada dataset yang sama, yang sebelumnya telah dibagi menjadi data pelatihan dan data pengujian. Model dijalankan menggunakan data yang telah diberi label guna mengidentifikasi keterkaitan antar atribut. Kinerja model kemudian diuji melalui data uji dengan mencocokkan prediksi dan label sebenarnya. Evaluasi performa dilakukan berdasarkan metrik *accuracy*, *precision*, *recall*, dan *F1-score* guna menentukan model yang paling optimal [13]. Pembagian data dalam penelitian ini dilakukan dengan proporsi 80% data pelatihan (1.067 data) dan 20% data pengujian (267 data).

2.8 Evaluasi Pola (*Pattern Evaluation*)

Setelah model *Naïve Bayes* menghasilkan hasil klasifikasi, dilakukan evaluasi terhadap akurasi dan ketepatan model. Pengujian pada penelitian ini menggunakan metode *cross validation*. Hasilnya dibandingkan dengan data realisasi penerima bantuan tahun sebelumnya.

2.9 Penyajian Pengetahuan (*Knowledge Presentation*)

Evaluasi performa algoritma dilakukan menggunakan *confusion matrix* yaitu evaluasi yang efektif dan rinci untuk mengukur kinerja algoritma klasifikasi, serta memiliki peran penting dalam bidang ilmu data [14]. Evaluasi model melalui confusion matrix menghasilkan ukuran seperti seperti *accuracy*, *precision*, *recall*, dan *F1-score*. Diterapkan guna mengevaluasi kinerja model dalam proses klasifikasi kelayakan penerima bantuan dengan tepat. Untuk menjamin evaluasi model yang lebih konsisten, digunakan pendekatan *k-fold cross-validation* dengan membagi dataset menjadi *k* bagian yang memiliki ukuran seragam. Masing-masing bagian bergantian menjadi data uji, sedangkan sisanya untuk pelatihan. Proses diulang *k* kali, dan hasilnya dirata-rata. Umumnya, *k* = 5 atau 10 [15].

Hasil evaluasi disajikan dalam bentuk tabel, grafik, dan interpretasi naratif untuk memberikan pemahaman menyeluruh mengenai performa masing-masing model.

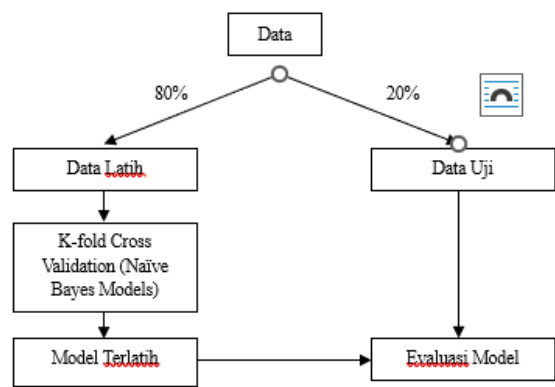
3. HASIL DAN PEMBAHASAN

3.1 Prosedur Evaluasi Model

Dataset didistribusikan ke dalam data latih (80%) dan data uji (20%) setelah tahap prapemrosesan selesai dilakukan. Algoritma *Gaussian*, *Bernoulli*, dan *Complement* digunakan untuk mengklasifikasikan kelayakan penerima BLT-DD. Untuk memilih model terbaik dari

masing-masing algoritma, dilakukan validasi menggunakan metode *10-fold cross-validation* terhadap data latih. Dalam metode ini, data latih dipartisi menjadi 10 bagian, dengan masing-masing bagian secara bergiliran difungsikan sebagai data validasi dan sisanya untuk melatih model.

Setelah mendapatkan model terbaik dari setiap algoritma, pengujian dilakukan menggunakan data uji yang belum pernah dipakai selama pelatihan. Tujuannya untuk mengukur kemampuan model dalam memprediksi data baru dengan akurat dan menghindari ketergantungan berlebihan pada data latih. Gambar 2 menyajikan gambaran umum dari proses implementasi metode yang digunakan.



Gambar 2. Diagram Alur Implementasi Model Naïve Bayes dengan K-Fold Cross-Validation

Hasil evaluasi *Gaussian Naïve Bayes* yang dilakukan dengan teknik *10-fold cross-validation* dapat dilihat pada Tabel 4.

Tabel 4. Pengujian *10-fold Gaussian Naïve Bayes*

K-Fold	Akurasi (%)
1	87,85 %
2	90,65 %
3	87,85 %
4	85,98 %
5	89,71 %
6	95,32 %
7	89,71 %
8	89,62 %
9	88,67 %
10	93,39 %
Rata-rata	89,88 %

Pada pengujian yang dilakukan, *Gaussian Naïve Bayes* menunjukkan hasil akurasi yang cukup tinggi dan stabil dengan rata-rata sebesar 89,88%. Tabel 5 menyajikan hasil pengujian algoritma *Bernoulli Naïve Bayes* menggunakan metode *10-fold cross-validation*.

Tabel 5. Pengujian *10-fold Bernoulli Naïve Bayes*

K-Fold	Akurasi (%)
1	91,58 %
2	93,45 %

K-Fold	Akurasi (%)
3	88,78 %
4	87,85 %
5	89,71 %
6	93,45 %
7	90,65 %
8	93,39 %
9	90,56 %
10	93,39 %
Rata-rata	91,28 %

Pada pengujian yang dilakukan, *Bernoulli Naïve Bayes* menunjukkan hasil akurasi yang cukup tinggi dan stabil dengan rata-rata sebesar 91,28%. Tabel 6 menyajikan hasil pengujian algoritma *Complement Naïve Bayes* menggunakan metode *10-fold cross-validation*.

Tabel 6. Pengujian *10-fold Complement Naïve Bayes*

K-Fold	Akurasi (%)
1	67,28 %
2	63,55 %
3	65,42 %
4	58,87 %
5	53,27%
6	63,55%
7	64,48 %
8	65,09 %
9	70,75 %
10	69,81 %
Rata-rata	64,21%

Pada pengujian yang dilakukan, hasil akurasi dari algoritma *Complement Naïve Bayes* menunjukkan rata-rata akurasi tertinggi, yaitu sebesar 64,21%.

3.2 Evaluasi Klasifikasi dengan *Confusion Matrix*

Sebanyak 267 data pengujian yang tidak digunakan dalam pelatihan dievaluasi untuk mengukur performa model. Hasilnya dianalisis melalui *confusion matrix* yang mencerminkan empat aspek utama.

True Positive (TP): Data diklasifikasikan Layak dan benar Layak

False Positive (FP) : Data diklasifikasikan Layak, tetapi sebenarnya Tidak Layak

False Negative (FN) : Data diklasifikasikan Tidak Layak, padahal sebenarnya Layak

True Negative (TN) : Data diklasifikasikan Tidak Layak dan benar Tidak Layak

Berdasarkan keempat komponen tersebut, pengukuran performa model klasifikasi dapat dihitung menggunakan rumus-rumus pada persamaan (2) hingga (5) [16]:

$$accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (2)$$

$$precision = \frac{TP}{TP + FP} \quad (3)$$

$$recall = \frac{TP}{TP + FN} \quad (4)$$

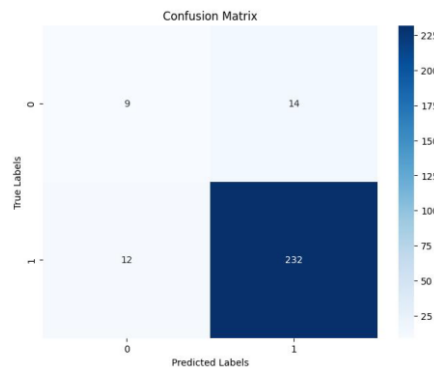
$$F1 - score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (5)$$

3.3 Evaluasi *Gaussian Naïve Bayes*

Berdasarkan confusion matrix, model menghasilkan 9 data *True Positive* (TP), 14 *False Negative* (FN), 12 *False Positive* (FP), dan 232 *True Negative* (TN).

- TP menggambarkan kasus di mana data diklasifikasikan sebagai Layak dan memang benar Layak, dengan jumlah 9 data.
- FP menunjukkan data yang diklasifikasikan sebagai Layak, namun sebenarnya Tidak Layak, dengan jumlah 14 data.
- FN mengindikasikan data yang diklasifikasikan sebagai Tidak Layak, padahal sebenarnya Layak, sebanyak 12 data.
- TN merepresentasikan data yang diklasifikasikan sebagai Tidak Layak dan memang benar Tidak Layak, dengan total 232 data.

Gambar 3 menyajikan visualisasi *confusion matrix* dari algoritma *Gaussian Naïve Bayes* yang menggambarkan distribusi hasil klasifikasi secara visual.



Gambar 3. *Confusion Matrix* Pada *Gaussian Naïve Bayes*

Evaluasi kinerja algoritma *Gaussian Naïve Bayes* dilakukan berdasarkan hasil pengujian, dan hasil pengukuran performa tersebut disajikan pada Tabel 7.

Tabel 7. Hasil Pengukuran Kinerja *Gaussian Naïve Bayes*

Class	Precision	Recall	F1-Score	Support
0	0.43	0.39	0.41	23
1	0.94	0.95	0.95	244
Accuracy			0.90	267
Macro Avg	0.69	0.67	0.68	267
Weighted Avg	0.90	0.90	0.90	267

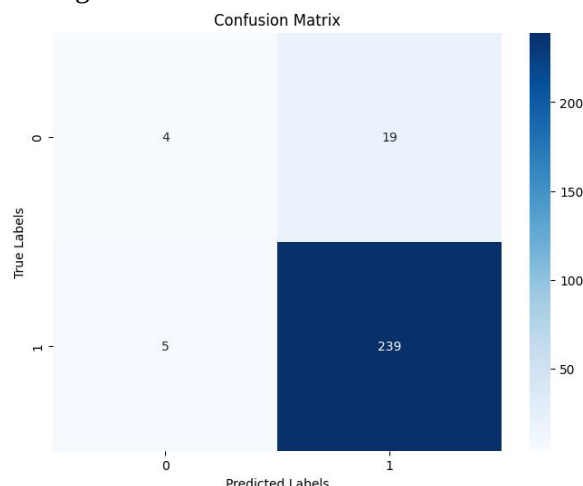
Berdasarkan Tabel 7, diperoleh nilai *accuracy* sebesar 90%. Untuk metrik *precision*, masing-masing kelas “Layak” dan “Tidak Layak” memiliki nilai sebesar 43% dan 94%. Sementara itu, nilai *recall* untuk kelas “Layak” adalah 39% dan untuk kelas “Tidak Layak” sebesar 95%. Adapun nilai *F1-Score* yang dihasilkan untuk kelas “Layak” dan “Tidak Layak” masing-masing sebesar 41% dan 95%. Ketiga metrik tersebut memperoleh rata-rata nilai masing-masing: 69% untuk *precision*, 67% untuk *recall*, dan 68% untuk *F1-score*.

3.4 Evaluasi *Bernoulli Naïve Bayes*

Berdasarkan hasil evaluasi menggunakan *confusion matrix*, diperoleh jumlah data *True Positive* (TP) sebanyak 4, *False Negative* (FN) sebanyak 19, *False Positive* (FP) sebanyak 5, dan *True Negative* (TN) sebanyak 239.

- TP menunjukkan bahwa data diklasifikasikan sebagai Layak dan memang benar Layak, dengan jumlah 4 data.
- FP menunjukkan bahwa data diklasifikasikan sebagai Layak, namun sebenarnya Tidak Layak, sebanyak 5 data.
- FN menunjukkan bahwa data diklasifikasikan sebagai Tidak Layak, padahal sebenarnya Layak, sebanyak 19 data.
- TN menunjukkan bahwa data diklasifikasikan sebagai Tidak Layak dan memang benar Tidak Layak, dengan jumlah 239 data.

Visualisasi *confusion matrix* pada algoritma *Bernoulli Naïve Bayes* ditampilkan pada Gambar 4 yang memberikan gambaran visual terkait distribusi hasil klasifikasi.



Gambar 4. *Confusion Matrix* pada *Bernoulli Naïve Bayes*

Tabel 8 menunjukkan hasil evaluasi performa algoritma *Bernoulli Naïve Bayes* yang dilakukan berdasarkan hasil pengujian.

Tabel 8. Hasil Pengukuran Kinerja *Bernoulli Naïve Bayes*

Class	Precision	Recall	F1-Score	Support
0	0.44	0.17	0.25	23
1	0.93	0.98	0.95	244
Accuracy			0.91	267
Macro Avg	0.69	0.58	0.60	267
Weighted Avg	0.88	0.91	0.89	267

Berdasarkan Tabel 8, diperoleh nilai *accuracy* sebesar 91%. Untuk metrik *precision*, kelas “Layak” memiliki nilai sebesar 44%, sedangkan kelas “Tidak Layak” sebesar 93%. Nilai *recall* yang dihasilkan untuk kelas “Layak” adalah 17% dan untuk kelas “Tidak Layak” sebesar 98%. Adapun nilai *F1-Score* masing-masing adalah 25% untuk kelas “Layak” dan 95% untuk kelas “Tidak Layak”. Ketiga metrik tersebut memperoleh rata-rata nilai masing-masing: 69% untuk *precision*, 58% untuk *recall*, dan 60% untuk *F1-score*.

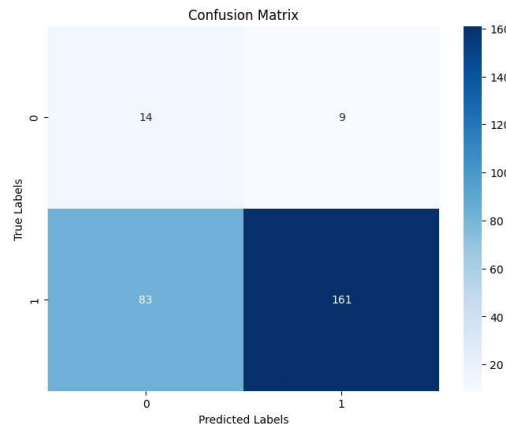
3.5 Evaluasi *Complement Naïve Bayes*

Berdasarkan hasil evaluasi menggunakan *confusion matrix*, diperoleh jumlah data *True Positive* (TP) sebanyak 14, *False Negative* (FN) sebanyak 9, *False Positive* (FP) sebanyak 83, dan *True Negative* (TN) sebanyak 161.

- TP menunjukkan bahwa data yang diklasifikasikan sebagai *Layak* memang benar merupakan data *Layak*, dengan jumlah 14 data.
- FP menunjukkan bahwa data yang diklasifikasikan sebagai *Layak*, ternyata sebenarnya adalah *Tidak Layak*, sebanyak 83 data.

- c. FN menunjukkan bahwa data diklasifikasikan sebagai *Tidak Layak*, padahal sebenarnya *Layak*, dengan jumlah 9 data.
- d. TN menunjukkan bahwa data diklasifikasikan sebagai *Tidak Layak* dan memang benar *Tidak Layak*, sebanyak 161 data.

Visualisasi *confusion matrix* pada algoritma *Complement Naïve Bayes* ditampilkan pada Gambar 5 yang memberikan gambaran visual terkait distribusi hasil klasifikasi.



Gambar 5. *Confusion Matrix* Pada *Complement Naïve Bayes*

Tabel 9 menampilkan evaluasi kinerja algoritma *Complement Naïve Bayes* yang dievaluasi berdasarkan hasil pengujian.

Tabel 9. Hasil Pengukuran Kinerja *Complement Naïve Bayes*

Class	Precision	Recall	F1-Score	Support
0	0.14	0.61	0.23	23
1	0.95	0.66	0.78	244
Accuracy			0.66	267
Macro Avg	0.55	0.63	0.51	267
Weighted Avg	0.88	0.66	0.73	267

Berdasarkan Tabel 9, diperoleh nilai *accuracy* sebesar 66%. Untuk metrik *precision*, kelas Layak memiliki nilai sebesar 14%, sedangkan kelas Tidak Layak sebesar 95%. Nilai *recall* untuk kelas Layak adalah 61%, dan untuk kelas Tidak Layak sebesar 66%. Adapun nilai *F1-Score* masing-masing adalah 23% kelas Layak dan 78% kelas Tidak Layak. Ketiga metrik evaluasi menghasilkan rata-rata nilai sebesar 55% untuk *precision*, 63% untuk *recall*, dan 51% untuk *F1-score*.

3.6 Perbandingan Kinerja Ketiga Algoritma

Berdasarkan hasil evaluasi klasifikasi, performa ketiga algoritma *Naïve Bayes* dapat dibandingkan berdasarkan indikator evaluasi kinerja seperti *accuracy*, *precision*, *recall*, dan *F1-score*. Tabel 10 berikut menyajikan ringkasan hasil pengukuran kinerja ketiga algoritma.

Tabel 10. Hasil Klasifikasi Tiga Tipe *Naïve Bayes*

Algorithm	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Gaussian Naïve Bayes	90%	69%	67%	68%
Bernoulli Naïve Bayes	91%	69%	58%	60%
Complement Naïve Bayes	66%	55%	63%	51%

Hasil evaluasi menunjukkan bahwa algoritma Bernoulli Naïve Bayes mencapai performa tertinggi dengan akurasi sebesar 91% mengungguli *Gaussian Naïve Bayes* (90%) dan *Complement Naïve Bayes* (66%). Keunggulan algoritma Bernoulli tidak terlepas dari karakteristik

data yang digunakan, yang sebagian besar merupakan data kategorikal. Data tersebut telah dikonversi ke dalam format biner menggunakan teknik *one-hot encoding*, yaitu metode yang merepresentasikan setiap kategori sebagai vektor angka 0 dan 1 agar dapat diproses oleh model kecerdasan buatan secara lebih efisien [17].

Bernoulli Naïve Bayes memang dirancang untuk menangani fitur biner, sehingga mampu mengenali pola secara lebih efektif dalam data penerima BLT-DD yang didominasi oleh nilai 0 dan 1. Sebaliknya, *Gaussian Naïve Bayes* mengasumsikan distribusi normal yang kurang sesuai untuk data diskret, meskipun tetap menghasilkan akurasi tinggi karena dominasi kelas “Tidak Layak”. Sementara itu, *Complement Naïve Bayes* menunjukkan performa terendah, karena algoritma ini umumnya lebih optimal untuk data teks dan distribusi kelas yang tidak seimbang, sedangkan data dalam penelitian ini relatif seimbang dan bukan berbentuk teks. Dengan demikian, kesesuaian antara karakteristik algoritma dan struktur data menjadi faktor utama yang menentukan efektivitas model klasifikasi.

4. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian, faktor yang paling memengaruhi kelayakan penerima BLT-DD adalah usia, status ekonomi, pendidikan, kepemilikan aset, dan jenis kelamin. Ketiga algoritma *Naïve Bayes* yang digunakan menunjukkan performa yang bervariasi, di mana *Bernoulli Naïve Bayes* memberikan akurasi tertinggi sebesar 91%, dibandingkan *Gaussian* (90%) dan *Complement* (66%). Keunggulan *Bernoulli* disebabkan kecocokannya dengan data biner hasil dari *one-hot encoding*, yang umum digunakan dalam klasifikasi data sosial ekonomi.

Temuan ini menunjukkan bahwa *Bernoulli Naïve Bayes* layak dijadikan metode utama dalam proses klasifikasi kelayakan penerima bantuan. Secara praktis, hasil ini dapat diterapkan untuk membantu pemerintah desa dalam menyeleksi calon penerima BLT-DD dengan pendekatan yang lebih objektif, efisien, dan transparan.

Sebagai tindak lanjut, algoritma *Bernoulli Naïve Bayes* layak untuk diintegrasikan ke dalam platform sistem pendukung keputusan berbasis web yang menyediakan fitur input data warga, klasifikasi otomatis, dan output rekomendasi penerima bantuan. Sistem ini juga dapat dilengkapi panel admin untuk pelacakan riwayat dan dokumentasi keputusan, sehingga meningkatkan akuntabilitas dan kemudahan dalam penyaluran bantuan.

DAFTAR PUSTAKA

- [1] Badan Pusat Statistik, “Persentase Penduduk Miskin (P0) Menurut Provinsi dan Daerah” [Online]. Available: <https://www.bps.go.id/id/statistics-table/2/MTkyIzI=/persentase-penduduk-miskin--maret-2023.html>
- [2] Badan Pusat Statistik, “Persentase penduduk miskin Kabupaten Kuningan pada Maret 2024.” [Online]. Available: <https://kuningankab.bps.go.id/id/pressrelease/2024/10/07/2901>
- [3] M. Maspawati *et al*, “Pengaruh Bantuan Langsung Tunai Dana Desa (BLT DD) Terhadap Kesejahteraan Masyarakat Desa Parenring, Kecamatan Lilirilau, Kabupaten Soppeng,” *J. Admin. Soc. Sci.*, vol. 4, no. 2, pp. 82–96, 2023, doi: 10.55606/jass.v4i2.351.stiayappimakassar.ac.id/index.php/jass/article/view/351/357
- [4] A. A. Maryoosh and E. M. Hussein, “A Review: Data Mining Techniques and Its Applications,” *Int. J. Comput. Sci. Mob. Appl.*, vol. 10, no. 3, pp. 1–14, 2022, doi: 10.47760/IJCSMA.2022.V10I03.001
- [5] A. Jain *et al*, “A Review: Data Mining Classification Techniques,” *Proc. 3rd Int. Conf. Intell. Eng. Manag. ICIEM 2022*, pp. 636–642, 2022, doi: 10.1109/ICIEM54221.2022.9853036.
- [6] A. Çınar, “Multi-Class Classification with the Gaussian Naive Bayes Algorithm,” *J. Data Appl.*, vol. 0, no. 2, pp. 1–13, 2024, doi: 10.26650/joda.1389471.
- [7] L. N. Darussalam *et al*, “Algoritma Naive Bayes Untuk Meningkatkan Model Klasifikasi Penerima Program Indonesia Pintar Di Sdn 2 Purwawinangun,” *JITET.*, vol. 13, no. 1, pp.

- 1129–1237, 2022, doi: 10.23960/jitet.v13i1.5882.
- [8] D. A. Huriah and N. D. Nuris, “Klasifikasi Penerima Bantuan Sosial Umkm Menggunakan Algoritma Naïve Bayes,” *JATI: Jurnal Mahasiswa Teknik Informatika*, vol. 7, no. 1, pp. 360–365, 2023, doi: 10.36040/jati.v7i1.6300.
- [9] A. Rahma Putri *et al*, “Klasifikasi Penerima Bantuan Langsung Tunai Dana Desa Menggunakan Algoritma K-Nearest Neighbor,” *JATI*, vol. 8, no. 4, pp. 6123–6131, 2024, doi: 10.36040/jati.v8i4.10194.
- [10] R. Justan *et al*, “Penelitian Kombinasi (Mixed Methods),” *ULIL ALBAB J. Ilm. Multidisiplin*, vol. 3, no. 2, pp. 253–263, Jan. 2024, doi: 10.56799/JIM.V3I2.2772.
- [11] A. U. B. Anelda *et al*, “Kualitatif: Memahami Karakteristik Penelitian Sebagai Metodologi,” *Jurnal Pendidikan Dasar*, vol. 11, no. 2, pp. 341–348, 2023, doi: 10.46368/JPD.V11I2.902.
- [12] M. Waruwu *et al*, “Metode Penelitian Kuantitatif: Konsep, Jenis, Tahapan dan Kelebihan,” *J. Ilm. Profesi Pendidik*, vol. 10, no. 1, pp. 917–932, 2025, doi: 10.29303/jipp.v10i1.3057.
- [13] M. I. Prasetyowati *et al*, “Determining threshold value on information gain feature selection to increase speed and prediction accuracy of random forest,” *J. Big Data*, vol. 8, no. 1, 2021, doi: 10.1186/s40537-021-00472-4.
- [14] S. dan Sathyanarayanan and B. R. Tantri, “Confusion Matrix-Based Performance Evaluation Metrics,” *Afr. J. Biomed. Res.*, vol. 27, no. 4S, pp. 4023–4031, 2024, doi: 10.53555/AJBR.v27i4S.4345.
- [15] T. J. Bradshaw *et al*, “A Guide to Cross-Validation for Artificial Intelligence in Medical Imaging,” *Radiol. Artif. Intell.*, vol. 5, no. 4, 2023, doi: 10.1148/ryai.220232.
- [16] Jude Chukwura Obi, “A comparative study of several classification metrics and their performances on data,” *World J. Adv. Eng. Technol. Sci.*, vol. 8, no. 1, pp. 308–314, 2023, doi: 10.30574/wjaets.2023.8.1.0054.
- [17] C. Liu *et al*, “Comparative evaluation of encoding techniques for workflow process remaining time prediction for cloud applications,” *J. Cloud Comput.*, vol. 14, no. 1, 2025, doi: 10.1186/s13677-025-00763-8.