



Klasifikasi teks dakwaan narkotika: komparasi *Naive Bayes* dan *Logistic Regression*

Steven F Handoko¹, Rudi Latuperissa²Email: 1682022105@student.uksw.edu, rudi.latuperissa@uksw.edu¹Program Studi Sistem Informasi, Fakultas Teknologi Informasi, Universitas Kristen Satya Wacana²Fakultas Teknologi Informasi, Universitas Kristen Satya Wacana

Diterima: 09 Juni 2026 | Direvisi: 07 September 2026 | Disetujui: 08 September 2026

©2020 Program Studi Teknik Informatika Fakultas Ilmu Komputer,
Universitas Muhammadiyah Riau, Indonesia

Abstrak

Tindak pidana narkotika di Indonesia menghasilkan akumulasi dokumen putusan hukum yang masif, khususnya di Pengadilan Negeri Medan sebagai wilayah dengan beban perkara tertinggi. Analisis dokumen dakwaan primair secara manual membutuhkan waktu lama dan rentan subjektivitas, sehingga diperlukan pendekatan komputasi. Penelitian ini bertujuan untuk mengklasifikasikan teks dakwaan primair kasus narkotika ke dalam kelas vonis ringan (label 0) dan vonis berat (label 1) menggunakan metode *naive bayes* dan *logistic regression*. Kebaruan metodologis penelitian ini terletak pada penentuan ambang batas biner objektif menggunakan nilai median durasi hukuman dari 500 dokumen sampel. Pemrosesan teks dilakukan melalui tahapan *text preprocessing* dan ekstraksi fitur berbasis pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF), dilanjutkan dengan pengujian menggunakan skema *stratified 10-fold cross validation*. Hasil eksperimen menunjukkan kesenjangan performa yang signifikan, di mana *logistic regression* mendominasi dengan mencatatkan tingkat akurasi (*classification accuracy*), *f1-score*, dan *recall* yang konsisten sebesar 91,6%, serta nilai AUC sebesar 0,978 (kategori *excellent classification*). Sebaliknya, *naive bayes* hanya mencapai akurasi dan *recall* sebesar 73,0% serta AUC 0,784, dengan angka misklasifikasi yang tinggi mencapai 135 dokumen akibat asumsi independensi mutlak antar-kata dan keterbatasan data latih. Keunggulan *logistic regression* didukung oleh regularisasi *ridge* yang efektif mereduksi efek tumpang tindih informasi pada matriks data hukum yang padat frasa terikat. Dengan demikian, model *logistic regression* berbasis TF-IDF lebih direkomendasikan sebagai instrumen komputasi untuk membantu melakukan klasifikasi biner terhadap dokumen putusan hukum tindak pidana narkotika.

Kata kunci: *dakwaan primair, klasifikasi teks, logistic regression, naive bayes, TF-IDF*

Classification of narcotics indictment texts: a comparison of Naive Bayes and Logistic Regression

Abstract

Narcotics crimes in Indonesia generate a massive accumulation of legal judgment documents, particularly at the Medan District Court, which handles the highest case load. Analyzing primary indictment documents manually is time-consuming and prone to subjectivity, necessitating a computational approach. This study aims to classify primary indictment texts of narcotics cases into light sentence (label 0) and heavy sentence (label 1) classes using *naive bayes* and *logistic regression* methods. The methodological novelty of this research lies in the determination of an objective binary threshold calculated from the median sentence duration of 500 sample documents. Text processing involves text preprocessing and feature extraction based on *Term Frequency-Inverse Document Frequency* (TF-IDF) weighting, followed by evaluation using a *stratified 10-fold cross-validation* scheme. Experimental results show a significant performance gap, with *logistic regression* dominating by recording a consistent classification accuracy (CA), *f1-score*, and *recall* of 91.6%, along with an AUC of 0.978 (*excellent classification category*). Conversely, *naive bayes* only achieves an accuracy and *recall* of 73.0% and an AUC of 0.784, producing a high misclassification rate of 135 documents due to the strict assumption of word independence and limited training data. The superiority of *logistic regression* is driven by *ridge* regularization, which effectively reduces the overlapping information in legal data matrices densely packed with bound phrases. Thus, the TF-IDF-based *logistic regression* model is highly recommended as a computational tool to assist in the binary classification of narcotics court judgment documents.

Keywords: , *primary indictment, text classification, logistic regression, naive bayes, TF-IDF*

1. PENDAHULUAN

Tindak pidana narkoba merupakan salah satu kejahatan luar biasa yang menyumbang beban perkara terbesar dalam sistem peradilan di Indonesia. Berdasarkan Indonesia *drug report* 2025, provinsi Sumatera utara merupakan wilayah dengan narapidana dan tahanan kasus narkoba tertinggi, sejumlah 19.378 orang [1]. Tingginya prevalensi ini berdampak langsung pada volume perkara yang harus ditangani oleh institusi peradilan di wilayah tersebut, termasuk Pengadilan Negeri Medan. Setiap proses persidangan menghasilkan tumpukan dokumen putusan hukum yang masif. Dokumen putusan ini terdiri beberapa bagian, salah satunya adalah dakwaan primair. Dakwaan primair memegang peran krusial karena memuat konstruksi peristiwa pidana dan kualifikasi perbuatan terdakwa yang nantinya akan sangat mempengaruhi berat atau ringannya vonis pidana yang dijatuhkan oleh majelis hakim. Namun, proses analisis dokumen hukum secara manual membutuhkan waktu yang lama dan rentan terhadap subjektivitas. Oleh karena itu, diperlukan pendekatan komputasi untuk mengekstraksi informasi dan mengenali pola dari data teks hukum yang tidak terstruktur tersebut.

Beberapa penelitian terdahulu telah berupaya mengimplementasikan teknik *text mining* dan klasifikasi otomatis pada berbagai jenis opini publik berbahasa Indonesia. Penelitian yang dilakukan oleh Misrun dkk. mengaplikasikan metode *naive bayes classifier* yang dipadukan dengan pembobotan kata *term frequency-inverse document frequency* (TF-IDF) untuk menganalisis sentimen komentar pada media sosial *YouTube*. Eksperimen tersebut menunjukkan bahwa pemanfaatan pengujian berbasis *data splitting* dengan proporsi 90% data latih dan 10% data uji mampu menghasilkan nilai akurasi sebesar 78% [2]. Hasil ini membuktikan bahwa algoritma *naive bayes* memiliki efektivitas yang baik dalam mengenali pola teks opini terstruktur, meskipun fokus penelitian tersebut masih terbatas pada analisis satu algoritma tunggal tanpa adanya pembandingan.

Penelitian lain mengenai optimalisasi performa algoritma tunggal juga dilakukan pada platform *Twitter*. Shiddicky dan Agustian memanfaatkan algoritma *Logistic Regression* untuk mengklasifikasikan data teks opini masyarakat terhadap kebijakan publik di media sosial *Twitter*. Dalam implementasi awalnya, model *logistic regression* mencatatkan nilai akurasi yang cukup tinggi sebesar 82%, namun nilai *F1-Score* yang dihasilkan kurang optimal akibat adanya masalah ketidakseimbangan distribusi data *training (imbalanced data)*. Untuk mengatasi kendala tersebut, penelitian mereka menerapkan proses pemotongan data (*slicing data*) serta optimasi parameter (*tuning hyperparameters*). Melalui tahapan optimasi ini, hasil pengujian akhir menunjukkan kinerja metrik yang jauh lebih stabil dan berimbang dengan nilai akurasi 67% serta *F1-Score* 60%, di mana performa tersebut terbukti lebih handal dibandingkan dengan algoritma lainnya seperti *naive bayes*, *support vector machine* (SVM), dan *long short-term memory* (LSTM) [3].

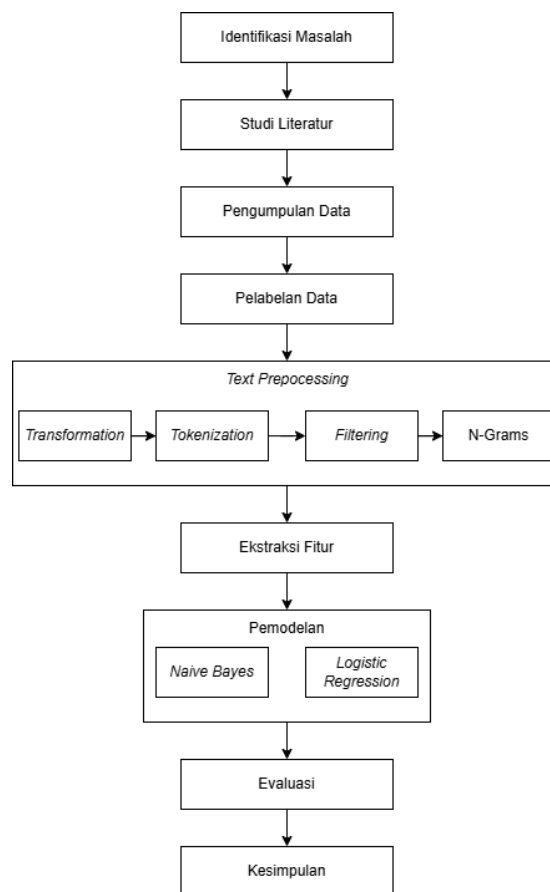
Untuk melihat perbandingan performa, penelitian komparatif dikembangkan oleh para peneliti lain, yaitu Ramadhani dan Suryono yang melakukan studi komparasi antara algoritma *naive bayes* dan *logistic regression* untuk menganalisis data sentimen teks pada media sosial X. Hasil eksperimen mereka menunjukkan bahwa *logistic regression* mampu mencapai nilai akurasi sedikit lebih tinggi yaitu sebesar 91%, dibandingkan dengan *naive bayes* yang memperoleh akurasi 90% [4]. Namun, penelitian tersebut mencatat adanya kendala berupa *imbalanced data*, di mana performa metrik seperti *precision* dan *recall* menjadi rendah akibat dominasi label mayoritas, sehingga model cenderung bias saat mempelajari kelas minoritas. Hal ini mengindikasikan perlunya pendekatan seleksi fitur atau penentuan ambang batas klasifikasi yang lebih ketat agar model tidak terjebak pada ketidakseimbangan kelas data.

Evaluasi komparatif yang lebih mendalam pada struktur komentar teks tidak terstruktur berbahasa Indonesia juga dikaji oleh Monica dan Purwanto. Penelitian mereka menguji kembali kinerja *naive bayes* dan *logistic regression* dengan mengintegrasikan teknik pembersihan teks (*preprocessing*) yang optimal serta pembobotan TF-IDF untuk menghasilkan performa klasifikasi opini dengan akurasi maksimum mencapai 84,88% [5]. Studi tersebut menegaskan bahwa algoritma berbasis linear seperti *Logistic Regression* sering kali memperlihatkan konsistensi yang lebih stabil pada data teks yang lebar dan memiliki dimensi tinggi, asalkan parameter model dioptimalkan secara tepat untuk mereduksi *noise*.

Sebagaimana studi literatur yang telah diuraikan, evaluasi komparatif antara algoritman *naive bayes* dan *logistic regression* sebagian besar diimplementasikan pada data opini teks media sosial. Oleh karena itu, penelitian ini menawarkan kebaruan dengan menerapkan kedua metode tersebut pada domain hukum formal, menggunakan dokumen teks dakwaan primair kasus narkoba dari Pengadilan Negeri Medan. Sisi kebaruan metodologis dalam penelitian ini terletak pada penerapan strategi penentuan kelas biner yang tidak lagi bersifat subjektif, melainkan dikalkulasikan secara matematis menggunakan nilai median vonis hukuman penjara untuk memisahkan kelas vonis ringan dan vonis berat. Penelitian ini bertujuan untuk mengklasifikasikan dokumen teks dakwaan primair kasus narkoba ke dalam dua kelas vonis pidana menggunakan metode *naive bayes* dan *logistic regression*, serta menganalisis perbandingan tingkat akurasi dan metrik evaluasi dari kedua algoritma tersebut dalam menangani data teks hukum formal.

2. METODE PENELITIAN

Metodologi penelitian merupakan kerangka kerja yang disusun secara sistematis dan metodis untuk memandu seluruh jalannya penelitian, mencakup seluruh tahapan perancangan kerangka berpikir, penyusunan matriks, hingga proses evaluasi kesimpulan [6] sehingga penyelesaian masalah yang dirumuskan dapat berjalan secara terarah demi mencapai hasil yang valid. Rangkaian tahapan terstruktur yang dilakukan dalam penelitian ini disajikan secara rinci pada gambar 1.



Gambar 1. Alur Penelitian

2.1. Pengumpulan Data

Pengumpulan data meliputi peninjauan terhadap kategori data yang digunakan untuk klasifikasi serta mekanisme yang diterapkan dalam mengumpulkan data tersebut. Sumber data penelitian ini berasal dari situs resmi direktori putusan mahkamah agung republik indonesia, yang didapatkan melalui proses pengunduhan berkas putusan hukum secara manual. Selanjutnya dokumen diekstraksi menggunakan bahasa pemrograman *python* dengan memanfaatkan *library* *PyPDF2*, *pandas*, dan *regular expression* untuk mengenali pola teks secara spesifik guna mendapatkan bagian dakwaan primair dan amar putusan.

2.2. Pelabelan Data

Tahap selanjutnya adalah memberi label terhadap seluruh data yang dikumpulkan sebelumnya. Label yang diberikan adalah label 0 dan label 1. Label 0 atau vonis ringan menunjukkan dokumen putusan dengan masa hukuman penjara yang kurang dari atau sama dengan nilai ambang batas (*threshold*). Label 1 atau vonis berat menunjukkan dokumen putusan dengan masa hukuman penjara yang berada di atas nilai ambang batas tersebut. Pelabelan biner dataset ini dilakukan secara objektif dengan menghitung nilai statistik median dari durasi hukuman 500 dokumen menggunakan *microsoft excel* untuk menetapkan ambang batas [7]. Pemilihan median sebagai nilai ambang batas didasarkan pada karakteristik distribusi data durasi vonis yang tidak selalu bersifat normal dan rentan terhadap pengaruh nilai ekstrem (*outlier*). Berbeda dengan nilai rata-rata yang dapat terdistorsi secara signifikan oleh vonis hukuman yang sangat tinggi seperti hukuman seumur hidup, median memberikan titik pembagi yang lebih stabil dan representatif terhadap kondisi aktual distribusi vonis di lapangan. Berdasarkan kalkulasi statistik terhadap 500 dokumen sampel, diperoleh nilai median durasi hukuman sebesar 6,125 tahun, yang selanjutnya ditetapkan sebagai ambang batas biner untuk memisahkan kelas vonis ringan (label 0) dan kelas vonis berat (label 1).

2.3. Text Preprocessing

Text preprocessing merupakan tahapan penting dalam pengolahan data teks yang bertujuan untuk membersihkan data dari deretan karakter, simbol, maupun kata yang tidak relevan guna meningkatkan efisiensi dan akurasi saat proses pemodelan klasifikasi [8]. Pada penelitian ini, tahapan praproses teks diterapkan menggunakan bantuan *widget preprocess text* yang tersedia pada *software orange data mining*. Dalam tahap ini ada beberapa langkah yang dilakukan, yaitu :

1. *Transformation*, merupakan tahapan pengubahan seluruh teks ke bentuk huruf kecil (*lowercase*) serta mengeliminasi karakter penekanan kata (*remove accents*).
2. *Tokenization*, merupakan tahapan memisahkan untaian kalimat panjang menjadi token kata mandiri berdasarkan karakter alfanumerik memanfaatkan ekspresi reguler (*regex*) bawaan sistem dengan pola penentu `\w+`.
3. *Filtering*, merupakan tahapan pembersihan komponen teks secara mendalam dengan menerapkan daftar kata hentian kustom (*custom stopwords*) bahasa Indonesia yang dirancang khusus untuk membuang elemen data kat administratif, nama aparat penegak hukum, nama terdakwa, lokasi geografis, serta runtutan birokrasi institusi peradilan yang tidak memuat esensi perkara hukum.
4. *Document frequency*, merupakan tahapan mengonfigurasi seleksi kapasitas kata secara relatif pada parameter ambang batas bawah guna menyingkirkan token yang terlalu langka atau terlalu umum dalam dataset.

2.4. Ekstraksi Fitur

Ekstraksi fitur merupakan proses transformasi data teks yang telah melalui tahap praproses menjadi representasi vektor numerik agar dapat diproses oleh algoritma komputasi. Dalam implementasinya, data teks tersebut diubah menggunakan teknik *Term Frequency-Inverse Document Frequency* (TF-IDF) yang menghitung frekuensi kata dalam sebuah dokumen sekaligus memberikan bobot yang lebih tinggi pada kata-kata yang jarang muncul [9]. Pada penelitian ini, proses ekstraksi fitur dieksekusi melalui *widget bag of words* pada *software orange data mining* dengan menerapkan metode *term frequency - inverse document frequency* (TF-IDF). Parameter-parameter teknis yang dikonfigurasi di dalam *widget* tersebut dijabarkan sebagai berikut :

1. *Term frequency* menggunakan opsi *count*, di mana bobot awal dihitung berdasarkan jumlah frekuensi kemunculan murni dari setiap token kata secara absolut di dalam masing-masing berkas dakwaan.
2. *Document frequency* menggunakan opsi *inverse document frequency* (IDF) untuk mengalikan nilai hitungan kata murni dengan bobot kebalikan distribusi dokumen, sehingga memberikan nilai bobot yang lebih tinggi pada token kata spesifik yang informatif dan mereduksi bobot token yang terlalu universal di seluruh dataset.
3. *Regularization*, menggunakan konfigurasi *none*, yang menandakan bahwa tidak ada transformasi normalisasi vektor tambahan yang diterapkan pada matriks bobot setelah nilai TF-IDF selesai dikalkulasi.

2.5. Naive Bayes

Naive bayes merupakan algoritma klasifikasi probabilistik yang bekerja berdasarkan teorema bayes dengan asumsi independensi yang kuat antar-fitur (kata) [10]. Persamaan matematika dasar teorema bayes disajikan pada persamaan 1.

$$P(C|X) = (P(X|C) \cdot P(C))/P(X) \quad (1)$$

Di mana X adalah data dengan kelas yang tiridentifikasi, yaitu kata yang terkandung di dalam dakwaan primair. C adalah hipotesis yang menunjukkan afiliasi data dengan kategori kelas tertentu, yaitu kelas vonis hukuman penjara. $P(C|X)$ adalah probabilitas posterior dari kelas target (vonis hukuman penjara) setelah fitur dakwaan diketahui. $P(X|C)$ merupakan probabilitas *likelihood*, yaitu peluang kemunculan fitur kata tertentu pada suatu kelas vonis. $P(C)$ adalah probabilitas *prior*, yaitu peluang kemunculan awal dari masing-masing kelas vonis di dalam dataset atau probabilitas hipotesa C sebelum mempertimbangkan data (prior probabilitas). Dan $P(X)$ adalah probabilitas *predictor prior*, yaitu konstanta peluang kemunculan kombinasi kata-kata dakwaan tersebut secara universal di dalam dataset.

Dalam implementasinya menggunakan perangkat lunak *orange data mining*, algoritma ini secara bawaan telah menerapkan teknik *laplace smoothing* dengan parameter konstanta nilai satu ($\alpha = 1$) [11]. Teknik ini berfungsi secara otomatis untuk mengatasi kendala *zero probability*, yaitu kondisi ketika terdapat kosakata baru pada dokumen uji yang tidak pernah terekam di data latihan sehingga dapat menyebabkan hasil akhir perkalian peluang algoritma langsung menjadi nol. Melalui penyesuaian tersebut, stabilitas performa model dalam mengenali variasi kata pada teks dakwaan primair dapat tetap terjaga dengan baik.

2.6. Logistic Regression

Logistic Regression merupakan algoritma klasifikasi linier yang digunakan untuk memprediksi probabilitas variabel terikat biner berdasarkan fungsi logistik atau fungsi sigmoid [12]. Pada penelitian ini, *Logistic Regression* digunakan untuk memprediksi peluang dokumen putusan hukum masuk ke dalam kategori vonis berat (label 1), di mana nilai komplemennya secara otomatis menentukan klasifikasi untuk kategori vonis ringan (label 0) berdasarkan nilai bobot fitur *Term Frequency - Inverse Document Frequency* (TF-IDF) dari teks dakwaan primair. Persamaan matematis fungsi sigmoid pada *Logistic Regression* disajikan pada persamaan 2.

$$f(x) = 1/(1 + e^{-(\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_n x_n)}) \quad (2)$$

Di mana $f(x)$ mewakili nilai output probabilitas biner dengan rentang nilai antara 0 hingga 1 yang menunjukkan peluang dokumen putusan masuk ke dalam kelas vonis berat (label 1). e mewakili bilangan eksponensial alami (bilangan Euler) yang bernilai konstan sebesar 2,718. β_0 mewakili nilai konstanta (*intercept*) awal pada model regresi ketika seluruh bobot fitur kata bernilai nol. $\beta_1, \beta_2, \dots, \beta_n$ mewakili koefisien regresi yang menunjukkan arah dan kekuatan pengaruh kontribusi untuk masing-masing token kata. $x_1, x_2, \dots, x_n$ mewakili nilai bobot numerik TF-IDF dari setiap token kata tunggal hasil ekstraksi fitur dokumen dakwaan primair, mulai dari kata pertama hingga kata ke- n .

Nilai output probabilitas kontinu dari fungsi sigmoid tersebut kemudian dikonversi menjadi keputusan kelas biner oleh sistem berdasarkan nilai ambang batas standar (*default threshold*) sebesar 0,5 [13]. Dokumen dakwaan dengan nilai probabilitas mencapai 0,5 atau lebih akan otomatis dimasukkan ke dalam kelas vonis berat (label 1), sedangkan nilai di bawah 0,5 masuk ke kelas vonis ringan (label 0). Selain itu, untuk mencegah model mengalami kesalahan analisis (*overfitting*) akibat tingginya dimensi matriks kata hasil ekstraksi TF-IDF, konfigurasi *logistic regression* pada *orange* menggunakan jenis regularisasi *Ridge* (L2) yang berfungsi untuk mengontrol dan mereduksi bobot fitur kata yang kurang informatif.

2.7. Implementasi dan evaluasi model

Implementasi dan evaluasi model merupakan tahapan krusial untuk mengintegrasikan seluruh komponen pemrosesan data ke dalam lingkungan sistem sekaligus memvalidasi performa algoritma secara objektif. Pada penelitian ini, seluruh alur pemrosesan data mulai dari impor data putusan hukum, proses pembersihan teks (*text preprocessing*), ekstraksi fitur numerik berbasis pembobotan *term frequency - inverse document frequency* (TF-IDF), hingga pelatihan algoritma klasifikasi diimplementasikan menggunakan perangkat lunak *orange data mining*. Selanjutnya, evaluasi keandalan model dieksekusi secara komputasi melalui *widget test and score* dengan menerapkan skema *stratified 10-fold cross validation* untuk mengukur tingkat akurasi prediksi dari algoritma *naive bayes* dan *logistic regression*. Melalui skema ini, keseluruhan 500 dokumen sampel dibagi secara acak namun proporsional ke dalam 10 bagian (*fold*) yang masing-masing berisi 50 dokumen. Di setiap bagian tersebut, sifat *stratified* memastikan distribusi kelas tetap berimbang, yaitu terdiri dari 25 dokumen vonis ringan dan 25 dokumen vonis berat. Proses pengujian dieksekusi sebanyak 10 kali perulangan, di mana pada setiap iterasi, 9 bagian (450 dokumen atau 90%) dialokasikan sebagai data latih dan 1 bagian (50 dokumen atau 10%) bertindak sebagai data uji secara bergantian. Nilai metrik evaluasi akhir yang disajikan merupakan hasil rata-rata akumulatif dari 10 tahapan pengujian paralel tersebut guna menghindari bias pembagian data tunggal [14]. Evaluasi akhir keberhasilan model dipetakan secara terukur menggunakan visualisasi instrumen *confusion matrix* yang menyajikan parameter nilai tingkat akurasi (*accuracy*), presisi (*precision*), *recall*, serta nilai perbandingan seimbang *F1-score* [15].

3. HASIL DAN PEMBAHASAN

3.1. Pengumpulan dan Pelabelan Data

Tahap awal penelitian dilakukan dengan mengumpulkan 500 dokumen putusan kasus tindak pidana narkotika dari Pengadilan Negeri Medan dalam rentang tahun 2022 hingga 2026. Berkas dakwaan primair yang berhasil diekstraksi kemudian melalui proses pelabelan biner berbasis nilai median durasi hukuman nyata [16], yang menghasilkan nilai ambang batas sebesar 6,125 tahun. Dokumen dengan amar putusan di bawah atau sama dengan 6,125 tahun dikategorikan sebagai vonis ringan (label 0), sedangkan dokumen dengan hukuman di atas 6,125 tahun dikategorikan sebagai vonis berat (label 1). Distribusi dataset yang berimbang ini menjadi fondasi utama dalam melatih model klasifikasi agar tidak bias terhadap salah satu kelas target [4].

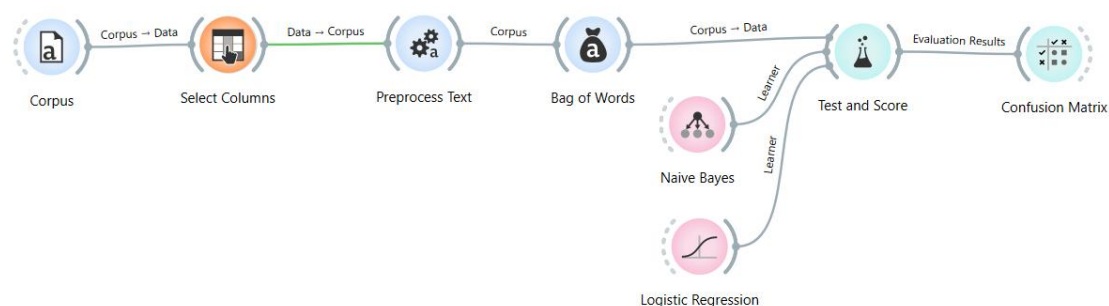
Tabel 1. Tabel Sampel Pelabelan Data

Nama File	Dakwaan Primair	Amar Putusan	Vonis	Label
putusan_1648_pid.sus_2025_pn_mdn_20260203214512.pdf	...tanpa hak atau melawan hukum menawarkan untuk dijual, menjual, membeli, menerima, menjadi perantara dalam jual beli, menukar atau menyerahkan Narkotika Golongan I... Perbuatan Terdakwa sebagaimana diatur dan diancam pada pasal 114 ayat (1) UU RI Nomor 35 Tahun 2009 Tentang Narkotika.	...2.Menjatuhkan pidana kepada Terdakwa oleh karena itu dengan pidana penjara selama 5 (lima) tahun dan 6 (enam) bulan...	5,5	0
putusan_1854_pid.sus_2025_pn_mdn_20260203214554.pdf	...Bahwa Terdakwa tidak memiliki izin dari pihak berwenang untuk dijual, menjual, membeli, menerima, menjadi perantara dalam jual beli, menukar, atau menyerahkan narkotika bukan tanaman yang mengandung Metamfetamina jenis sabu. Perbuatan Terdakwa sebagaimana diuraikan tersebut di atas, diatur dan diancam pidana Pasal 114 ayat (1) UU No. 35 tahun 2009 tentang Narkotika.	...2.Menjatuhkan pidana terhadap Terdakwa oleh karena itu dengan pidana penjara selama 6 (enam) Tahun dikurangi selama Terdakwa berada dalam tahanan dan pidana denda sejumlah Rp. 300.000.000,00 (tiga ratus juta rupiah) dengan ketentuan pembayarannya dapat dilakukan dengan cara ...	6	0

putusan_1635_pid.sus_2025_pn_mdn_20260203215129.pdf	...tanpa hak atau melawan hukum menawarkan untuk dijual, menjual, membeli, menerima, menjadi perantara dalam jual beli, menukar atau menyerahkan Narkotika Golongan I berupa narkotika jenis sabu dengan berat 16,85 (enam belas koma delapan puluh lima) gram netto... Perbuatan Terdakwa sebagaimana diatur dan diancam pidana Pasal 114 ayat (2) UU RI No. 35 Tahun 2009 Tentang Narkotika Jo Pasal 55 ayat (1) ke-1 KUHPidana;	...2.Menjatuhkan pidana kepada Terdakwa oleh karena itu dengan pidana penjara selama 8 (delapan) tahun...	8	1
putusan_2019_pid.sus_2025_pn_mdn_20260203215657.pdf	...Menawarkan Untuk Dijual, Menjual, Membeli, Menjadi Perantara Dalam Jual Beli, Menukar, Menyerahkan, Atau Menerima Narkotika Golongan I Sebagaimana Dimaksud Pada Ayat (1) Yang Dalam Bentuk Tanaman Beratnya Melebihi 1 (Satu) Kilogram Atau Melebihi 5 (Lima) Batang Pohon Atau Dalam Bentuk Bukan Tanaman Beratnya 5 (Lima) Gram... Sebagaimana Perbuatan Terdakwa diatur dan diancam pidana menurut Pasal 114 ayat (2) UURI No. 35 tahun 2009 Tentang Narkotika.	...2.Menjatuhkan pidana terhadap Terdakw a oleh karena itu dengan pidana penjara selama 15 (lima belas) Tahun serta denda sejumlah Rp500.000.000. 00 (lima ratus juta rupiah) yang harus dibayar dalam jangka waktu 2 (dua) tahun sejak putusan berkekuatan hukum tetap...	15	1
putusan_1991_pid.sus_2025_pn_mdn_20260203215849.pdf	...Dari hasil analisis bahwa barang bukti milik Terdakwa Jetman Hutahean adalah benar mengandung Metamfetamina dan terdaftar dalam Golongan I (satu) nomor urut 61 Lampiran I Undang-Undang Republik Indonesia No. 35 Tahun 2009 tentang Nakrotika. Perbuatan ia Terdakwa sebagaimana diatur dan diancam pidana melanggar Pasal 114 ayat (1) UURI Nomor 35 Tahun 2009 Tentang Narkotika.	...2.Menjatuhkan pidana penjara terhadap Terdakwa selama 6 (enam) tahun dan denda sebesar Rp1.000.000.000,00 (satu miliar rupiah) dengan ketentuan apabila denda tersebut tidak dibayar maka diganti dengan Pidana Penjara selama 3 (tiga) bulan...	6	0

3.2. Implementasi Pemodelan Klasifikasi

Implementasi pemodelan klasifikasi biner pada penelitian ini dirancang menggunakan arsitektur aliran data (*workflow*) terintegrasi pada perangkat lunak *orange data mining*. Aliran data secara menyeluruh disajikan pada gambar 2, yang memperlihatkan bagaimana data teks dakwaan narkotika diproses secara runtut dari tahap input hingga tahap evaluasi paralel. Proses diawali dengan memuat dataset yang berisi 500 dokumen putusan melalui *widget corpus*. Dokumen tersebut kemudian dialirkan ke *widget select columns* untuk mengisolasi fitur secara fungsional, di mana kolom teks "Dakwaan_Primair" ditetapkan sebagai variabel *metas* dan kolom "Label " dikunci sebagai variabel *target* biner yang akan diprediksi oleh model.



Gambar 2. Arsitektur Workflow Implementasi Model Klasifikasi pada *Orange*

Setelah struktur data terkunci, korpus teks dialirkan menuju *widget preprocess text* untuk dibersihkan dari komponen kata administratif peradilan yang tidak informatif melalui sub-proses pengubahan huruf kecil (*lowercase*), tokenisasi reguler ekspresi (*w+*), serta penyaringan kata hentian (*custom stopwords*). Guna membatasi dimensi fitur, kapasitas kata dioptimasi menggunakan batas dokumen frekuensi relatif (*document frequency*) berkisar antara 0,10 hingga 0,95, disertai perluasan frasa kustom *n-grams* (unigram dan bigram) untuk menjaga keutuhan makna istilah hukum. Teks bersih yang dihasilkan kemudian ditransformasikan ke dalam representasi matriks bobot numerik untuk memberikan nilai bobot matematis yang proporsional berdasarkan tingkat kemunculan kata di seluruh dokumen [17] melalui *widget bag of words* menggunakan skema kalkulasi *term frequency - inverse document frequency* (TF-IDF) dengan parameter *count* dan *inverse document frequency* (IDF) aktif tanpa normalisasi tambahan (*none*). Matriks bobot inilah yang disuplai secara simultan menuju *widget algoritma naive bayes* dan *logistic regression* sebelum akhirnya bermuara pada *widget test and score* untuk dievaluasi kinerjanya.

3.3. Analisis Komparasi Performa Model Klasifikasi

Evaluasi kinerja komputasi untuk menguji keandalan prediksi algoritma dilakukan menggunakan skema *stratified 10-fold cross validation*. Parameter evaluasi yang digunakan meliputi metrik *area under ROC curve* (auc), *classification accuracy* (ca), *f1-score*, *precision*, dan *recall*. Berkedudukan sebagai instrumen pengujian utama, hasil perbandingan performa antara algoritma *logistic regression* dan *naive bayes* disajikan secara ringkas pada tabel 2.

Tabel 2. Perbandingan Model *Naive Bayes* dan *Logistic Regression*

Model	AUC	CA	F1	Prec	Recall	MCC
Naive Bayes	0.784	0.730	0.728	0.738	0.730	0.468
Logistic Regression	0.978	0.916	0.916	0.918	0.916	0.834

Berdasarkan hasil pengujian pada tabel 2, terlihat adanya perbedaan signifikan performa akurasi yang cukup signifikan antara kedua algoritma yang diuji. Model *logistic regression* menunjukkan dominasi dengan mencatatkan tingkat akurasi (CA), *f1-score*, dan *recall* yang sangat konsisten sebesar 0,916 (91,6%), serta nilai *precision* sebesar 0,918. Nilai akurasi yang tinggi ini didukung oleh capaian skor AUC (*area under the ROC curve*) sebesar 0,978, yang membuktikan bahwa model *logistic regression* memiliki kapabilitas prediksi yang masuk dalam kategori sangat baik.

Sebaliknya, model *naive bayes* menunjukkan performa yang jauh lebih rendah di seluruh metrik evaluasi. Algoritma probabilistik ini hanya mampu memperoleh nilai CA dan *recall* sebesar 0,730 (73,0%), nilai *precision* sebesar 0,738, nilai *f1-score* sebesar 0,728, serta nilai AUC sebesar 0,784. Kesenjangan akurasi sebesar 18,6% ini mengindikasikan adanya perbedaan mendasar pada cara kedua algoritma dalam memproses fitur matriks renggang (*sparse matrix*) yang dihasilkan oleh pembobotan TF-IDF [18]. Dalam konteks teks dakwaan primair, pembobotan TF-IDF menghasilkan matriks fitur yang didominasi oleh istilah-istilah hukum spesifik dengan frekuensi kemunculan yang tidak merata antar dokumen. Istilah-istilah seperti pasal rujukan, jenis narkoba, dan frasa dakwaan formal cenderung memiliki distribusi yang sangat berbeda antara kasus vonis ringan dan vonis berat, sehingga bobot TF-IDF yang dihasilkan menjadi sangat informatif namun sekaligus kompleks untuk diproses oleh algoritma yang bergantung pada asumsi independensi antar fitur. Kondisi inilah yang secara langsung memengaruhi kemampuan masing-masing algoritma dalam mengekstraksi pola diskriminatif dari matriks fitur tersebut [19].

Dari perbandingan kedua model yang digunakan, sesuai dengan penelitian yang dilakukan oleh R.E. Pambudi, dkk model *logistic regression* terbukti mampu memberikan keseimbangan antara mendeteksi kelas yang ada, dalam hal ini kelas vonis ringan (label 0) dan kelas vonis berat (label 1) [20]. Sedangkan model *naive bayes* belum begitu akurat dalam membedakan kedua kelas biner tersebut, terutama akibat tingginya tingkat kesalahan klasifikasi pada dokumen yang memuat karakteristik fitur kata yang kompleks. Model probabilistik ini rentan mengalami bias akibat tumpang tindihnya penggunaan kosakata administratif di dalam berkas dakwaan, sehingga batas probabilitas antara kelas vonis ringan dan vonis berat menjadi samar.

3.4. Analisis Confusion Matrix

Guna membedah sebaran prediksi benar dan salah, instrumen *confusion matrix* digunakan untuk mengaudit luaran dari 500 dokumen sampel yang diuji. Visualisasi matriks kesalahan untuk model *naive bayes* dan *logistic regression* masing-masing disajikan secara berdampingan pada tabel 3 dan tabel 4

Tabel 3. Confusion Matrix Model *Logistic Regression*

Kelas Aktual / Prediksi	0	1	Σ
0	237	13	250
1	29	221	250
Σ	266	234	500

Tabel 4. Confusion Matrix Model *Naive Bayes*

Kelas Aktual / Prediksi	0	1	Σ
0	206	44	250
1	91	159	250
Σ	297	203	500

Berdasarkan pengujian pada tabel 3, dari total 500 dokumen putusan hukum tindak pidana narkoba, model *logistic regression* berhasil mengklasifikasi secara akurat sebanyak 237 dokumen sebagai kelas vonis ringan (label 0) dan 221 dokumen sebagai kelas vonis berat (label 1). Model ini hanya menghasilkan total kesalahan klasifikasi (*misclassification*) sebanyak 42 dokumen yang terdiri dari 29 dokumen salah mengklasifikasikan label 1 sebagai label 0 dan 13 dokumen salah mengklasifikasikan label 0 sebagai label 1.

Sebaliknya, pada tabel 4 tingkat kesalahan prediksi pada model *naive bayes* cukup tinggi, mencapai 135 dokumen salah prediksi. Model *naive bayes* hanya mampu mendeteksi secara tepat 206 dokumen vonis ringan (label 0) dan 159 dokumen vonis berat (label 1), sementara 44 dokumen label 0 keliru diklasifikasi menjadi label 1 dan 91 dokumen label 1 salah diklasifikasi sebagai label 0. Tingginya angka salah klasifikasi pada *naive bayes* sangat berisiko dalam konteks analisis hukum karena gagal mengenali karakteristik dokumen kasus-kasus narkoba.

3.5. Perbandingan Kedua Model

Perbedaan performa yang terjadi dipengaruhi secara kuat oleh karakteristik internal dari masing-masing algoritma dalam menangani dimensi data teks hukum. Algoritma *logistic regression* mampu bekerja secara lebih efektif dalam menangani keterkaitan antar kata yang kompleks karena didukung oleh implementasi regularisasi *ridge* (12). Regularisasi ini berfungsi untuk mengendalikan nilai koefisien regresi dari ribuan fitur kata hasil ekstraksi TF-IDF, meminimalkan dampak keterikatan antar-kata, serta mengeklusi bobot kata yang redundan tanpa menghilangkan informasi penting. Pada dokumen dakwaan primair yang kaya akan frasa hukum terikat seperti 'tanpa hak atau melawan hukum', 'menawarkan untuk dijual', dan 'sebagaimana diatur dan diancam pidana', regularisasi *ridge* efektif karena mampu mereduksi bobot fitur kata administratif yang berulang dan tidak diskriminatif antar kelas, sekaligus mempertahankan bobot fitur kata yang benar-benar membedakan karakteristik vonis ringan dan vonis berat. Hal ini memungkinkan *logistic regression* untuk menangkap pola linier yang bermakna dari matriks TF-IDF meskipun dimensi fiturnya sangat tinggi dan banyak fitur yang saling berkorelasi [21]. Keseimbangan nilai metrik pada *logistic regression* mencerminkan kemampuan model yang lebih efektif dalam menangani data text dengan fitur yang saling terkait, baik untuk meminimalkan kesalahan *false positive* maupun *false negative* [22].

Di sisi lain, kelemahan mendasar dari algoritma *naive bayes* terletak pada asumsi independensi mutlak yang diembannya. Algoritma ini mengasumsikan bahwa setiap kata di dalam dokumen dakwaan berdiri sendiri dan tidak memiliki keterkaitan konteks dengan kata lainnya [22]. Pada realitas dokumen hukum, asumsi ini tidak akurat karena teks dakwaan disusun berdasarkan rangkaian frasa yang saling mengikat (misalnya, kata "tanpa" dan "hak" harus dibaca sebagai satu kesatuan konteks "tanpa hak"). Dalam pembobotan TF-IDF, frasa-frasa hukum terikat tersebut dipecah menjadi token kata tunggal yang masing-masing memiliki bobot independen, sehingga makna kontekstual dari rangkaian frasa tersebut hilang. *Naive bayes* yang kemudian memproses token-token ini secara terpisah, tidak mampu merekonstruksi makna frasa hukum secara utuh, sehingga estimasi probabilitas kelas yang dihasilkan menjadi kurang akurat dalam membedakan karakteristik dakwaan vonis ringan dan vonis berat [23]. Akibatnya, tingginya volatilitas kata yang redundan dan pengabaian struktur frasa membuat *naive bayes* mengalami misklasifikasi yang tinggi, di mana banyak dokumen vonis berat keliru diprediksi sebagai vonis ringan, dan sebaliknya.

Selain itu, performa *naive bayes* yang kurang optimal dalam mengklasifikasikan dokumen hukum ini juga dipengaruhi oleh faktor keterbatasan jumlah *dataset* dan data latih yang digunakan. Karakteristik probabilistik dari algoritma *naive bayes* memerlukan pasokan sampel data yang melimpah guna membangun estimasi peluang kemunculan kata yang stabil dan komprehensif pada matriks fitur yang lebar, kurangnya kuantitas data latih berbanding lurus dengan penurunan tingkat akurasi pada model *naive bayes* [24]. Keterbatasan volume data ini menyebabkan model rentan mengalami bias saat mempelajari pola kata baru yang jarang muncul di dokumen latihan, sehingga memicu rendahnya nilai akurasi akhir. Hal ini menegaskan bahwa untuk memperoleh tingkat akurasi yang lebih tinggi dan akurat, algoritma *naive bayes* sangat bergantung pada perluasan dan penambahan jumlah *dataset* yang jauh lebih masif pada penelitian selanjutnya.

4. KESIMPULAN

Berdasarkan hasil eksperimen dan analisis komparatif yang telah dilakukan, dapat disimpulkan bahwa implementasi pemrosesan teks berbasis pembobotan *term frequency - inverse document frequency* (TF-IDF) berhasil mengklasifikasikan berat ringannya vonis hukuman penjara kasus narkoba secara biner menggunakan dokumen dakwaan primair. Melalui skema pengujian validasi silang terhadap keseluruhan dokumen sampel, algoritma *logistic regression* terbukti memiliki performa yang jauh lebih unggul dan stabil dibandingkan dengan algoritma *naive bayes*. Model *logistic regression* mencatatkan nilai tingkat akurasi, *f1-score*, *precision*, dan *recall* yang sangat konsisten, serta nilai *auc* yang mengindikasikan kualitas prediksi dalam kategori sangat baik. Keunggulan ini divalidasi secara nyata oleh instrumen *confusion matrix*, di mana model ini mampu memprediksi sebagian besar dokumen vonis ringan dan vonis berat secara akurat dengan tingkat kesalahan klasifikasi yang sangat minim.

Di sisi lain, model *naive bayes* menghasilkan performa yang kurang optimal di seluruh metrik evaluasi serta mencatatkan angka kesalahan klasifikasi yang jauh lebih tinggi. Perbedaan performa yang signifikan ini dipengaruhi secara kuat oleh karakteristik internal dari masing-masing algoritma dalam menangani dimensi data teks hukum. Faktor keunggulan *logistic regression* didukung oleh implementasi *regularisasi ridge* yang terbukti sangat efektif dalam mengendalikan nilai koefisien regresi dari ribuan fitur kata serta mengurangi efek tumpang tindih informasi antar-kata pada data teks. Sementara itu, kelemahan mendasar *naive bayes* terletak pada asumsi independensi mutlak antar-kata yang patah ketika dihadapkan pada karakteristik dokumen hukum yang didominasi oleh rangkaian frasa. Dampak dari kelemahan *naive bayes* ini menghasilkan angka kesalahan dalam mendeteksi vonis berat yang tinggi, sehingga sangat berisiko dalam konteks analisis hukum yuridis karena gagal mengenali

karakteristik dokumen kasus-kasus narkoba berskala besar. Dengan demikian, model *logistic regression* berbasis pembobotan TF-IDF lebih direkomendasikan sebagai model klasifikasi untuk membantu melakukan prediksi atau klasifikasi biner terhadap dokumen putusan hukum tindak pidana narkoba.

Penelitian ini memiliki beberapa keterbatasan yang perlu diperhatikan dalam penginterpretasian hasil yang didapatkan. Pertama, dataset yang digunakan terbatas pada 500 dokumen putusan dari Pengadilan Negeri Medan dalam rentang tahun 2022 hingga 2026, sehingga generalisasi model terhadap dokumen hukum dari pengadilan lain di luar Pengadilan Negeri Medan belum dapat dijamin. Karakteristik teks dakwaan primair, pola penyusunan frasa hukum, serta distribusi vonis yang berlaku di setiap pengadilan dapat bervariasi secara signifikan bergantung pada yurisdiksi dan kebijakan masing-masing institusi peradilan. Kedua, jumlah dataset sebanyak 500 dokumen masih tergolong terbatas untuk membangun model klasifikasi yang sepenuhnya komprehensif, terutama bagi algoritma naive bayes yang sangat bergantung pada kelimpahan data latih untuk menghasilkan estimasi probabilitas yang stabil. Oleh karena itu, penerapan model yang dikembangkan dalam penelitian ini pada dokumen hukum di luar Pengadilan Negeri Medan perlu dilakukan dengan kehati-hatian dan memerlukan pengembangan lebih lanjut.

DAFTAR PUSTAKA

- [1] Badan Narkotika Nasional Republik Indonesia, "Indonesia Drug Report 2025," 2025. Diakses: 15 Mei 2026. [Daring]. Tersedia pada: <https://puslitdatin.bnn.go.id/konten/unggah/2025/06/IDR-2025.pdf>
- [2] Chely Aulia Misrun, E. Haerani, M. Fikry, dan E. Budianita, "Analisis sentimen komentar youtube terhadap Anies Baswedan sebagai bakal calon presiden 2024 menggunakan metode naive bayes classifier," *Jurnal CoSciTech (Computer Science and Information Technology)*, vol. 4, no. 1, hlm. 207–215, Apr 2023, doi: 10.37859/coscitech.v4i1.4790.
- [3] Ash Shiddicky dan Surya Agustian, "Analisis Sentimen Masyarakat Terhadap Kebijakan Vaksinasi Covid-19 pada Media Sosial Twitter menggunakan Metode Logistic Regression," *Jurnal CoSciTech (Computer Science and Information Technology)*, vol. 3, no. 2, hlm. 99–106, Agu 2022, doi: 10.37859/coscitech.v3i2.3836.
- [4] B. Ramadhani dan R. R. Suryono, "Komparasi Algoritma Naïve Bayes dan Logistic Regression Untuk Analisis Sentimen Metaverse," *Jurnal Media Informatika Budidarma*, vol. 8, no. 2, hlm. 714, Apr 2024, doi: 10.30865/mib.v8i2.7458.
- [5] M. Monica dan A. Purwanto, "Evaluasi Komparatif Kinerja Algoritma Naïve Bayes dan Logistic Regression dalam Klasifikasi Sentimen Komentar YouTube Berbahasa Indonesia," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 6, no. 2, hlm. 934–943, 2026, doi: 10.57152/malcom.v6i2.2639.
- [6] D. A. Trisliatanto, *Metodologi Penelitian Panduan Lengkap Penelitian Dengan Mudah*. Yogyakarta: CV Andi Offset, 2020.
- [7] Muhammad Rizki Syafapri, Elin Haerani, Iwan Iskandar, dan Liza Afriyanti, "Klasifikasi sentimen terhadap larangan pernikahan beda agama menggunakan metode Naive Bayes Classifier," *Jurnal CoSciTech (Computer Science and Information Technology)*, vol. 5, no. 1, hlm. 10–18, Apr 2024, doi: 10.37859/coscitech.v5i1.6889.
- [8] A. Oktian Permana dan Sudin Saepudin, "Perbandingan algoritma k-nearest neighbor dan naïve bayes pada aplikasi shopee," *Jurnal CoSciTech (Computer Science and Information Technology)*, vol. 4, no. 1, hlm. 25–32, Apr 2023, doi: 10.37859/coscitech.v4i1.4474.
- [9] M. Mubarak, L. Tanti, dan R. Rosnelly, "Perbandingan Algoritma Decision Tree dan Naive Bayes Pada Analisis Sentimen Masyarakat Terhadap Pejabat Pertamina Pasca Kasus Pertamina Oplosan," *Jurnal Minfo Polgan*, vol. 15, no. 1, hlm. 179–188, Mar 2026, doi: 10.33395/jmp.v15i1.15971.
- [10] M. F. Rozi, R. Siregar, dan N. I. Syahputri, "Penerapan Data Mining Menggunakan Metode Naive Bayes Untuk Klasifikasi Data Penentuan Hasil Penjualan Dalam Strategi Pemasaran," *Jurnal Komputer Teknologi Informasi Sistem Komputer*, vol. 2, hlm. 444–454, 2023.
- [11] D. Utami dan P. A. R. Devi, "Klasifikasi Kelayakan Penerima Bantuan Program Keluarga Harapan (PKH) Menggunakan Metode Weighted Naive Bayes Dengan Laplace Smoothing," *JIPi (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, vol. 07, no. 04, hlm. 1373–1384, 2022.
- [12] Z. Zulkifli dan R. Fajri, "Klasifikasi Tingkat Kematangan Buah Strawberry Menggunakan Algoritma Logistic Regression," *Data Sciences Indonesia (DSI)*, vol. 4, no. 2, hlm. 50–59, Des 2024, doi: 10.47709/dsi.v4i2.4850.
- [13] W. A. Ma'arif, Sarwindo, dan T. Tamrin, "Analisis Sentimen Terhadap Ulasan Game Mobile Legend di Playstore menggunakan Algoritma Logistic Regression," *JUKI : Jurnal Komputer dan Informatika*, vol. 8, no. 1, hlm. 43–49, 2026.
- [14] N. F. Sahamony, T. Terttiaavini, dan H. Rianto, "Analisis Perbandingan Kinerja Model Machine Learning untuk Memprediksi Risiko Stunting pada Pertumbuhan Anak," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 2, hlm. 413–422, Feb 2024, doi: 10.57152/malcom.v4i2.1210.
- [15] I. Attyyatullatifah dan M. Kamayani, "Perbandingan Algoritma Klasifikasi untuk Prediksi Kelulusan Mahasiswa Teknik Informatika dengan Orange Data Mining," *Indonesian Journal of Computer Science*, vol. 13, no. 2, hlm. 3127–3140, 2024.
- [16] A. S. D. P. Sinaga dan A. S. Aji, "Analisis Sentimen Publik Terhadap Mayor Teddy Indra Wijaya dengan Pendekatan Logistic Regression," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 5, no. 1, hlm. 222–231, Des 2024, doi: 10.57152/malcom.v5i1.1752.
- [17] A. R. Manoppo, D. Nurandani, D. Adrian, dan M. I. Ayuda, "Penerapan Naive Bayes untuk Memprediksi Status Keberhasilan Transaksi pada Sistem Top Up AY Pulsa Menggunakan Aplikasi Orange," 2026.
- [18] R. A. Husen, R. Astuti, L. Marlia, R. Rahmaddeni, dan L. Efrizoni, "Analisis Sentimen Opini Publik pada Twitter Terhadap Bank BSI Menggunakan Algoritma Machine Learning," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 3, no. 2, hlm. 211–218, Okt 2023, doi: 10.57152/malcom.v3i2.901.
- [19] R. Wati, S. Ernawati, dan H. Rachmi, "Pembobotan TF-IDF Menggunakan Naïve Bayes pada Sentimen Masyarakat Mengenai Isu Kenaikan BIPIH," *Jurnal Manajemen Informatika (JAMIKA)*, vol. 13, no. 1, hlm. 84–93, Apr 2023, doi: 10.34010/jamika.v13i1.9424.
- [20] R. E. Pambudi, H. Purnomo, dan R. Irawan, "Pemanfaat Data Mining Untuk Prediksi Prestasi Akademik Siswa Berdasarkan Pola Kehadiran, Aktivitas Belajar Menggunakan Naive Bayes Logistic Regression," *Jurnal Teknologi Informasi Mura*, vol. 16, hlm. 132, 2024.
- [21] Sutarman, R. Siringoringo, D. Arisandi, E. Kurniawan, dan E. B. Nababan, "Model Klasifikasi Dengan Logistic Regression Dan Recursive Feature Elimination Pada Data Tidak Seimbang," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 11, no. 4, hlm. 735–742, Agu 2024, doi: 10.25126/jtiik.1148198.
- [22] O. N. Cahyani dan F. Budiman, "Performa Logistic Regression dan Naive Bayes dalam Klasifikasi Berita Hoax di Indonesia," *Edumatic: Jurnal Pendidikan Informatika*, vol. 9, no. 1, hlm. 60–68, Apr 2025, doi: 10.29408/edumatic.v9i1.28987.
- [23] A. Gerliandeva, Y. Chrisnanto, dan H. Ashaury, "Optimasi Klasifikasi Sentimen pada Komentar Online menggunakan Multinomial Naïve Bayes dan Ekstraksi Fitur TF-IDF serta N-grams," *Jurnal Pekommas*, vol. 9, no. 2, hlm. 260–272, Des 2024, doi: 10.56873/jpkm.v9i2.5585.
- [24] Noviolan Jehovan Dieksa dan I. Pakereng, "Analisis sentimen masyarakat terhadap putusan Mahkamah Konstitusi tentang batasan usia calon Presiden dan Wakil Presiden di media sosial Twitter," *IT-Explore: Jurnal Penerapan Teknologi Informasi dan Komunikasi*, vol. 5, no. 1, hlm. 1–10, Feb 2026, doi: 10.24246/itexplore.v5i1.2026.pp1-10.