

Implementasi *Naïve Bayes* untuk Evaluasi dan Klasifikasi Beban Kerja Pegawai di Badan Kepegawaian Kabupaten Barru

Muhammad Alif Tenriadjeng^a, Ihwana As'ad^b, Amaliah Faradibah^c

Universitas Muslim Indonesia, Makassar, Indonesia

^a13020200275@umi.ac.id; ^bihwana.asad@umi.ac.id; ^camaliah.faradibah@umi.ac.id;

Received: 10-02-2025 | Revised: 02-08-2025 | Accepted: 12-09-2025 | Published: 29-09-2025

Abstrak

Penelitian ini mengkaji penerapan algoritma *Naïve Bayes* dalam mengklasifikasikan beban kerja pegawai di Badan Kepegawaian Kabupaten Barru guna meningkatkan efektivitas manajemen sumber daya manusia berbasis data. Tantangan utama dalam pengelolaan pegawai meliputi ketidakseimbangan distribusi tugas, kurangnya transparansi penilaian, serta keterbatasan metode konvensional, yang dapat diatasi melalui analisis data historis dengan pendekatan probabilistik. Algoritma *Naïve Bayes* digunakan untuk mengelompokkan beban kerja pegawai berdasarkan parameter jumlah jam kerja, kompleksitas tugas, dan pencapaian target, yang menghasilkan distribusi beban kerja rendah (55,7%), sedang (31,4%), dan tinggi (12,8%). Model ini menunjukkan akurasi tinggi dalam klasifikasi dan mampu mengidentifikasi pegawai dengan kelebihan atau kekurangan beban kerja secara objektif. Studi kasus pada PT. Rikasa Dinar Djaya juga membuktikan keandalan metode ini dalam mengevaluasi kinerja karyawan dengan akurasi 99,44%. Hasil penelitian ini diharapkan dapat meningkatkan transparansi dan keadilan dalam evaluasi kinerja pegawai, meminimalkan bias subjektif, serta menjadi referensi bagi instansi pemerintah lainnya dalam mengadopsi teknologi machine learning untuk pengelolaan SDM yang lebih modern, efisien, dan berkelanjutan.

Kata kunci: *Naïve Bayes*, Klasifikasi Beban Kerja, Evaluasi Kinerja Pegawai.

Pendahuluan

Di era transformasi digital, pengelolaan sumber daya manusia (SDM) menjadi faktor kunci dalam meningkatkan efektivitas dan produktivitas organisasi, khususnya di sektor pemerintahan [1]. Badan Kepegawaian Kabupaten Barru, sebagai instansi yang bertanggung jawab dalam manajemen pegawai, menghadapi tantangan dalam menilai serta mengklasifikasikan beban kerja secara objektif. Ketidakseimbangan dalam distribusi tugas, kurangnya transparansi dalam penilaian, serta keterbatasan metode konvensional kerap menjadi kendala dalam mengoptimalkan kinerja pegawai. Oleh karena itu, diperlukan pendekatan berbasis data yang dapat memberikan solusi sistematis untuk mengatasi permasalahan tersebut [2].

Salah satu solusi yang dapat diterapkan adalah pemanfaatan algoritma *Naïve Bayes*, sebuah metode klasifikasi berbasis probabilitas yang dikenal efisien dan akurat dalam mengolah data kategorikal [3]. Algoritma ini dianggap sesuai untuk mengelompokkan beban kerja pegawai berdasarkan berbagai parameter, seperti jumlah jam kerja, tingkat kompleksitas tugas, pencapaian target, serta indikator kinerja lainnya [4]. Dengan menganalisis data historis dan pola kerja yang terukur, *Naïve Bayes* dapat membantu mengidentifikasi pegawai yang mengalami kelebihan atau kekurangan beban kerja. Hal ini memungkinkan distribusi tugas yang lebih seimbang dan efektif dalam meningkatkan kinerja organisasi [5].

Penelitian ini bertujuan untuk menguji efektivitas algoritma *Naïve Bayes* dalam mengevaluasi dan mengklasifikasikan beban kerja pegawai di lingkungan Badan Kepegawaian Kabupaten Barru. Data yang digunakan mencakup rekam jejak kinerja pegawai yang kemudian diproses melalui tahapan pra-pemrosesan data, pelatihan model, dan validasi hasil klasifikasi [6]. Pendekatan ini diharapkan mampu menghasilkan rekomendasi berbasis data untuk meningkatkan kualitas pengambilan keputusan manajerial, khususnya dalam penempatan pegawai dan penyusunan kebijakan pengembangan SDM.

Penelitian [7] dalam melakukan evaluasi kinerja pada PT. Rikasa Dinar Djaya dengan menerapkan metode *Naïve Bayes* Berdasarkan hasil analisis menggunakan RapidMiner dengan algoritma *Naïve Bayes*, karyawan dengan jabatan *Maintenance*, *Admin*, *Leader*, *Scheduler*, dan *Team Sales* tercatat sebagai kelompok yang paling sering mengalami keterlambatan dibandingkan dengan jabatan lainnya. Model *Naïve Bayes* menunjukkan tingkat akurasi sebesar 99,44%, yang menegaskan keandalan metode ini dalam evaluasi kinerja karyawan. Oleh karena itu, penerapan algoritma *Naïve Bayes* dapat menjadi alternatif dalam proses

pengambilan keputusan di PT. Riksa Dinar Djaya. Dengan akurasi yang tinggi, metode ini diharapkan dapat membantu perusahaan dalam menentukan kebijakan yang mendorong perbaikan kinerja bagi karyawan yang belum mencapai standar optimal.

Penelitian ini memiliki signifikansi dalam meningkatkan transparansi dan keadilan dalam evaluasi beban kerja pegawai. Dengan menghilangkan bias subjektif, algoritma *Naïve Bayes* dapat berfungsi sebagai alat bantu keputusan yang andal bagi pimpinan institusi. Selain itu, penerapan metode ini berpotensi meningkatkan motivasi pegawai melalui distribusi tugas yang lebih proporsional, sekaligus meminimalkan risiko kelelahan kerja (burnout) yang dapat berdampak negatif terhadap produktivitas.

Hasil penelitian ini diharapkan tidak hanya bermanfaat bagi Badan Kepegawaian Kabupaten Barru, tetapi juga menjadi referensi bagi instansi pemerintah lainnya dalam mengadopsi teknologi machine learning untuk pengelolaan SDM. Dengan demikian, inovasi ini dapat mendorong terciptanya tata kelola kepegawaian yang lebih modern, responsif, dan berkelanjutan di tingkat daerah maupun nasional.

Metode

Klasifikasi Bayesian merupakan metode pengklasifikasian statistik yang digunakan untuk memprediksi probabilitas suatu objek termasuk dalam kelas tertentu. Metode ini didasarkan pada Teorema Bayes dan memiliki kemampuan klasifikasi yang sebanding dengan decision tree serta neural network. Klasifikasi Bayesian terbukti efektif dalam hal akurasi dan kecepatan, terutama saat diterapkan pada database berukuran besar. Salah satu algoritma dalam metode ini adalah *Naïve Bayes*, yang meskipun sederhana, tetap mampu memberikan hasil klasifikasi dengan tingkat akurasi yang tinggi [8].

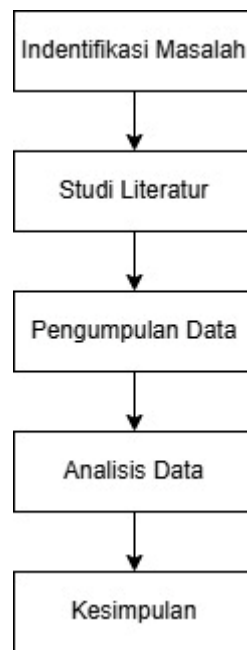
$$P(C_i | X) = \frac{P(X|C_i) \cdot P(C_i)}{P(X)} \quad (1)$$

Naïve Bayes merupakan metode klasifikasi yang sederhana dan banyak digunakan karena kemudahan penerapannya serta kemampuannya menghasilkan prediksi yang baik dalam berbagai kasus. Namun, kelemahan utama dari metode ini terletak pada asumsi bahwa setiap fitur bersifat independen dalam menentukan kelas, yang dalam kenyataannya tidak selalu berlaku. Dalam praktiknya, sering kali terdapat keterkaitan antara variabel, sehingga dapat memengaruhi akurasi model [9]. Alasan lain dalam memilih metode *Naïve Bayes* adalah karena proses klasifikasinya yang efektif serta mampu menghasilkan tingkat akurasi yang tinggi. Meskipun menggunakan algoritma yang sederhana, *Naïve Bayes* tetap dapat memberikan hasil prediksi yang akurat. Selain itu, algoritma ini memiliki beberapa keunggulan seperti mudah di implementasi dan memberikan hasil yang baik dalam beberapa kasus [10].

Confusion Matrix adalah alat dan metode yang digunakan untuk evaluasi visual dalam konsep Machine Learning. Kolom pada Confusion Matrix mewakili hasil prediksi kelas, sementara baris menunjukkan kelas yang sebenarnya. Confusion Matrix mencakup semua kemungkinan kasus dalam masalah klasifikasi, dengan mengambil contoh klasifikasi biner yang memiliki dimensi 2×2 . Berbagai ukuran kinerja algoritma dapat dihitung menggunakan Confusion Matrix, seperti tingkat deteksi negatif dan tingkat recall untuk kelas positif. Langkah-langkah ini umumnya berlaku untuk semua algoritma klasifikasi [11].

Dalam Confusion Matrix, setiap kolom mewakili instance dari kelas yang diprediksi, sementara setiap baris mewakili instance dari kelas yang sebenarnya. Matriks ini menggambarkan bagaimana model klasifikasi menunjukkan kebingungannya dalam membuat prediksi. Confusion Matrix tidak hanya memberikan informasi mengenai kesalahan yang dilakukan oleh model, tetapi juga jenis kesalahan yang terjadi. Pada masalah klasifikasi biner, Confusion Matrix berbentuk tabel 2×2 . Dengan asumsi satu kelas sebagai "positif" dan yang lainnya sebagai "negatif", terdapat empat kombinasi kemungkinan antara kelas yang diprediksi dan kelas aktual, yang disebut: True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN) [12].

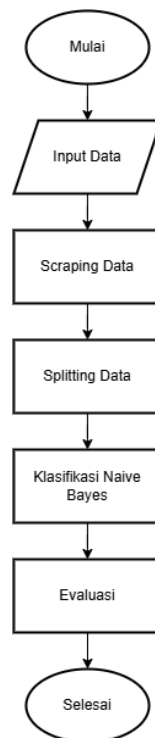
Penulis memberikan gambaran terhadap alur penelitian yang dilakukan. Tahapan penelitian yang dibangun ditunjukkan pada Gambar 1



Gambar 1. Tahap Penelitian

Penelitian ini dimulai dengan mengidentifikasi masalah pada kinerja pegawai pada kantor badan kepegawaian Kabupaten Barru. Tahap selanjutnya adalah studi literatur dengan memahami penelitian serupa yang telah dilakukan oleh peneliti sebelumnya dengan penerapan teori yang sama mengenai klasifikasi tingkat kinerja pegawai. Pengumpulan data dilakukan berdasarkan rekam jejak kinerja pegawai selama dalam waktu kerja. Selanjutnya dilakukan analisis data dengan menerapkan klasifikasi Naive Bayes untuk menghasilkan nilai akurasi. Library yang digunakan pada Google colab seperti stopwords dan stemming, serta nilai akurasi dari metode tersebut

Perancangan



Gambar 2. Desain Penelitian

A. Input Data

Langkah pertama dalam penelitian ini adalah pengumpulan data terkait beban kerja pegawai di Badan Kepegawaian Kabupaten Barru. Data diperoleh melalui rekaman kinerja pegawai sesuai dengan beban pekerjaan. Data ini kemudian dimasukkan ke dalam sistem untuk diproses lebih lanjut.

B. Preprocessing

Tahap *preprocessing* ini bertujuan untuk mengolah data menta menjadi data yang siap digunakan sebelum digunakan sebagai bahan analisis. Sehingga, proses ini memastikan bahwa data yang telah di *processing* telah bebas dari gangguan yang dapat mempengaruhi hasil. Proses ini meliputi:

a. Case Folding

Mengubah semua teks menjadi huruf kecil untuk menghilangkan perbedaan antar kapitalisasi.

b. Tokenizing

Memisahkan teks menjadi kata-kata individual atau token.

c. Filtering

Menghapus kata-kata yang tidak relevan, seperti tanda baca, angka, atau simbol.

d. Stemming

Mengubah kata-kata ke bentuk pada dasarnya untuk menyederhanakan analisis.

C. Klasifikasi Naive Bayes

Pada tahap ini, data yang telah diproses dianalisis menggunakan algoritma *Naïve Bayes*. Algoritma ini bekerja dengan prinsip probabilitas untuk mengklasifikasikan beban kerja pegawai dalam beberapa kategori. Setiap kategori ditentukan berdasarkan pola data yang ada. *Naïve Bayes* dipilih karena kemampuannya yang efisien dalam menangani data berbasis teks dan memberikan hasil yang akurat.

D. Evaluasi

Langkah terakhir adalah mengevaluasi performa model klasifikasi yang telah dibuat. Dalam penelitian ini, digunakan *Accuracy Score* untuk menilai tingkat akurasi model. Evaluasi dilakukan dengan cara membandingkan hasil prediksi model dengan data aktual. Nilai akurasi digunakan untuk menentukan seberapa baik algoritma *Naïve Bayes* dalam mengklasifikasi beban kerja pegawai.

Dalam penelitian ini, proses persiapan data berkaitan dengan data yang diperoleh dari media sosial, yang merupakan sumber informasi paling dinamis dalam mencerminkan perilaku manusia dan kondisi sosial terkini. Namun, data dari media sosial tidak akan memiliki nilai guna jika tidak melalui tahap pengolahan yang tepat. Oleh karena itu, diperlukan teknik khusus untuk mengekstraksi, membersihkan, dan mengolah data agar dapat diinterpretasikan secara akurat [13]. Proses pengolahan data, khususnya ekstraksi informasi dari web atau media sosial untuk mengidentifikasi pola tersembunyi di dalamnya, dikenal sebagai *scraping*, yang menjadi langkah awal dalam analisis lebih lanjut [14].

Tahap *preprocessing* terdiri dari empat langkah utama. Pertama, data cleaning dan case folding, yang bertujuan membersihkan data dari kesalahan, termasuk noise, serta mengubah seluruh teks menjadi huruf kecil. Kedua, filtering, yang berfungsi menghapus kata atau karakter yang tidak relevan dan sering muncul, termasuk slang dan kata-kata pendek. Ketiga, stemming, yaitu proses mengambil kata dasar atau akar kata untuk mendukung klasifikasi teks. Keempat, tokenizing, yang memecah kalimat menjadi unit makna lebih spesifik dengan menggunakan konsep bigram, yakni pasangan dua kata [15].

Pemodelan

A. Scraping Data

Hasil *scraping* data yang ditampilkan dalam tabel di Google Colab menunjukkan dataset pegawai di Badan Kepegawaian Kabupaten Barru yang telah diproses untuk keperluan klasifikasi beban kerja menggunakan *Naïve Bayes*. Dataset ini berisi 100 baris data pegawai, dengan informasi utama seperti nama, NIP, jabatan, golongan, tanggal mulai jabatan, jumlah jam kerja, kompleksitas tugas, pencapaian target, dan kategori beban kerja. Data yang diperoleh menunjukkan variasi dalam beban kerja pegawai berdasarkan jumlah jam kerja dan kompleksitas tugas. Pegawai dengan jam kerja yang lebih tinggi serta tugas lebih kompleks cenderung masuk dalam kategori beban kerja tinggi, sementara pegawai dengan jam kerja lebih sedikit

atau tugas yang lebih ringan diklasifikasikan ke dalam beban kerja rendah atau sedang. Misalnya, seorang pegawai dengan jumlah jam kerja 182 dan kompleksitas tugas 5 dikategorikan memiliki beban kerja tinggi, sedangkan pegawai dengan jam kerja 154 dan kompleksitas 2 dikategorikan memiliki beban kerja rendah.

	NO	NAMA	NIP	JABATAN	TMT JABATAN	GOL	TMT GOL	Jumlah Jam Kerja	Kompleksitas Tugas	Pencapaian Target (%)	Kategori Beban Kerja
0	1	SYAMSIR, S.IP, M.Si.	197001011990031012	KEPALA BADAN KEPEGAWAIAN DAN PENGEMBANGAN SUMB...	2022-04-28 00:00:00	IV/c	2019-04-01 00:00:00	178	5	68	Sedang
1	2	SUPARMAN, S.Sos	196811111993031014	SEKRETARIS BADAN KEPEGAWAIAN DAN PENGEMBANGAN ...	2020-01-07 00:00:00	IV/b	2021-04-01 00:00:00	191	1	85	Rendah
2	3	MUHAMMAD HARIS RAHIM, S.AP	198505262003121002	KEPALA BIDANG MUTASI, PROMOSI DAN KOMPETENSI	28/04/2022	III/d	2023-04-01 00:00:00	168	4	61	Sedang
3	4	NUR AMDAN, S.Sos	1983111112008031001	KEPALA BIDANG PENGADAAN, PEMBERHENTIAN DAN INF...	2023-05-11 00:00:00	IV/a	2023-10-01 00:00:00	154	2	79	Rendah
4	5	HASANUDDIN RASYID, S.Pd	197311222009011002	KABID PENDIDIKAN DAN LATIHAN, PENILAIAN KINERJ...	2023-05-11 00:00:00	III/d	2021-04-01 00:00:00	182	5	87	Tinggi
...	...	...	...	...	...	...	...	...	...	...	...
95	96	Herlina	1970960119900310096	Staf BKPSDM	2023-01-01	III/c	2022-01-01	193	3	99	Sedang
96	97	Yusril Ananda	1970970119900310097	Staf BKPSDM	2023-01-01	IV/a	2022-01-01	162	1	92	Rendah
97	98	Andi Faizal	1970980119900310098	Staf BKPSDM	2023-01-01	III/d	2022-01-01	176	5	68	Sedang
98	99	Taufik Hidayat	1970990119900310099	Staf BKPSDM	2023-01-01	IV/a	2022-01-01	171	2	98	Rendah
99	100	Andi Wahyudi	197010001199003100100	Staf BKPSDM	2023-01-01	III/c	2022-01-01	172	2	88	Rendah

Gambar 1. Hasil Scraping Data

B. Split Data

Hasil split data menunjukkan bahwa dari 100 data pegawai, sebanyak 70 data digunakan untuk pelatihan model (training set) dan 30 data digunakan untuk pengujian model (testing set). Pembagian ini dilakukan dengan proporsi 70% untuk training dan 30% untuk testing, yang merupakan praktik umum dalam machine learning untuk memastikan model memiliki cukup data untuk belajar serta diuji dengan data yang belum pernah dilihat sebelumnya. Pembagian ini bertujuan untuk melatih model *Naïve Bayes* dengan data training agar dapat mempelajari pola hubungan antara fitur seperti jumlah jam kerja, kompleksitas tugas, dan pencapaian target, kemudian menguji performanya menggunakan data testing untuk melihat seberapa baik model dapat mengklasifikasikan kategori beban kerja pegawai.

C. Klasifikasi *Naïve Bayes*

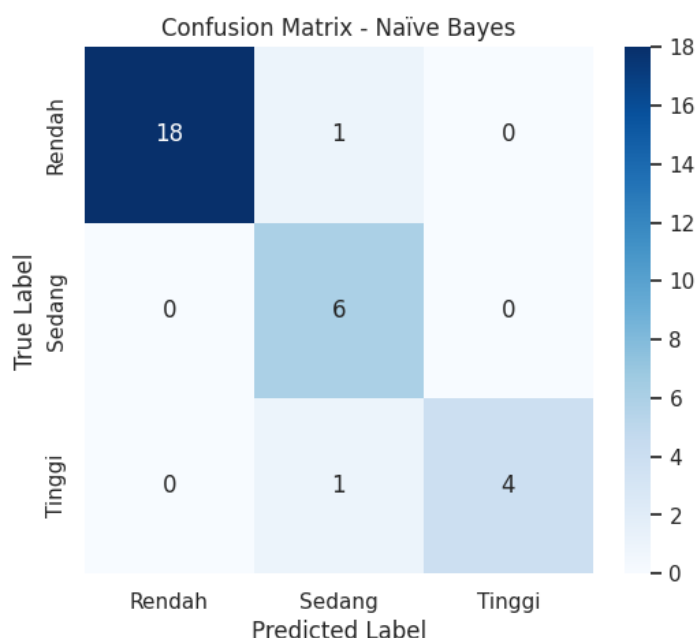
Pemodelan *Naïve Bayes* pada gambar yang diberikan menunjukkan bagaimana model mempelajari hubungan antara fitur-fitur yang digunakan untuk mengklasifikasikan beban kerja pegawai. Model ini bekerja berdasarkan prinsip probabilitas bersyarat, di mana ia menghitung kemungkinan seorang pegawai termasuk dalam kategori rendah, sedang, atau tinggi berdasarkan nilai jumlah jam kerja, kompleksitas tugas, dan pencapaian target. Dalam proses pemodelan, model pertama-tama menentukan prior probabilities, yang menunjukkan distribusi awal kelas sebelum mempertimbangkan fitur. Hasilnya, mayoritas pegawai berada dalam kategori beban kerja rendah, diikuti oleh kategori sedang, sementara kategori beban kerja tinggi memiliki jumlah yang lebih sedikit. Setelah itu, model menghitung mean (rata-rata) dan variance (variansi) untuk setiap fitur dalam masing-masing kategori. Rata-rata fitur menunjukkan bahwa pegawai dengan beban kerja tinggi memiliki nilai yang lebih besar untuk jumlah jam kerja dan kompleksitas tugas, sedangkan kategori rendah cenderung memiliki nilai lebih kecil. Variansi fitur menggambarkan seberapa tersebar nilai fitur dalam masing-masing kategori, di mana kategori sedang memiliki variasi yang lebih besar, sedangkan kategori tinggi memiliki variasi yang lebih kecil, menunjukkan data yang lebih seragam.

D. Klasifikasi *Naïve Bayes*

Hasil evaluasi model menggunakan confusion matrix menunjukkan performa klasifikasi model *Naïve Bayes* dalam mengelompokkan beban kerja pegawai ke dalam tiga kategori: rendah, sedang, dan tinggi. Pada confusion matrix, terdapat tiga kelas yang diprediksi oleh model: Rendah (0), Sedang (1), dan Tinggi (2). Dari total 30 data uji, model berhasil mengklasifikasikan sebagian besar data dengan benar, menghasilkan akurasi keseluruhan sebesar 93%. Model mengklasifikasikan 19 sampel sebagai beban kerja rendah, di mana 18 benar diklasifikasikan sebagai rendah (True Positive), dan 1 salah diklasifikasikan sebagai sedang (False Negative). Untuk kategori beban kerja sedang, model berhasil mengklasifikasikan 6 sampel dengan benar, tanpa kesalahan klasifikasi ke kategori lain. Sedangkan untuk beban kerja tinggi, model mengklasifikasikan 5 sampel, di mana 4 benar diklasifikasikan sebagai

tinggi, dan 1 salah diklasifikasikan sebagai sedang. Dari hasil precision, recall, dan F1-score, terlihat bahwa model memiliki precision tinggi (100%) untuk kelas "rendah" dan "tinggi", artinya semua sampel yang diprediksi dalam kategori tersebut mayoritas benar. Namun, recall untuk kelas "tinggi" adalah 80%, yang berarti ada beberapa data dari kelas tinggi yang salah diklasifikasikan ke kelas lain. Secara keseluruhan, model *Naïve Bayes* bekerja dengan baik, dengan sedikit kesalahan klasifikasi, terutama dalam membedakan kategori "rendah" dan "sedang". Kesalahan utama terjadi ketika beberapa sampel dengan beban kerja tinggi diklasifikasikan sebagai sedang, yang mungkin disebabkan oleh kemiripan karakteristik antara dua kelas tersebut dalam dataset.

Laporan Klasifikasi:				
	precision	recall	f1-score	support
0	1.00	0.95	0.97	19
1	0.75	1.00	0.86	6
2	1.00	0.80	0.89	5
accuracy			0.93	30
macro avg	0.92	0.92	0.91	30
weighted avg	0.95	0.93	0.94	30



Gambar 2. Confusion Matrix

Kesimpulan

Hasil klasifikasi menggunakan metode *Naïve Bayes* pada gambar menunjukkan bagaimana model mempelajari distribusi data berdasarkan tiga kategori beban kerja pegawai: rendah, sedang, dan tinggi. Dari hasil prior probabilities, terlihat bahwa kelas dengan jumlah terbesar dalam dataset adalah beban kerja rendah dengan probabilitas 55.7%, diikuti oleh beban kerja sedang sebesar 31.4%, dan beban kerja tinggi yang hanya memiliki probabilitas 12.8%. Hal ini menunjukkan bahwa mayoritas pegawai dalam dataset memiliki beban kerja yang relatif lebih ringan dibandingkan dengan kategori lainnya. Rata-rata fitur dalam setiap kelas menggambarkan bagaimana karakteristik jumlah jam kerja, kompleksitas tugas, dan pencapaian target berbeda di antara kategori beban kerja. Pegawai dengan beban kerja tinggi memiliki rata-rata nilai fitur yang lebih besar dibandingkan kategori lainnya, yang berarti mereka cenderung memiliki jam kerja yang lebih panjang, tugas yang lebih kompleks, serta pencapaian target yang lebih tinggi. Sebaliknya, pegawai dengan beban kerja rendah memiliki nilai rata-rata yang lebih kecil, yang menandakan bahwa setelah standarisasi,

mereka cenderung memiliki jam kerja lebih sedikit, tugas yang lebih sederhana, dan pencapaian target yang lebih rendah. Variansi fitur dalam setiap kelas menunjukkan sejauh mana penyebaran data dalam kategori tersebut. Kategori beban kerja sedang memiliki variansi yang lebih tinggi dibandingkan dengan kategori lainnya, yang menunjukkan bahwa pegawai dalam kategori ini memiliki distribusi nilai yang lebih beragam dalam jumlah jam kerja, kompleksitas tugas, dan pencapaian target. Sebaliknya, kategori beban kerja tinggi memiliki variansi yang lebih kecil, yang menunjukkan bahwa pegawai dalam kelompok ini memiliki karakteristik yang lebih seragam. Secara keseluruhan, model *Naïve Bayes* telah mampu menangkap pola hubungan antara fitur-fitur yang digunakan dalam klasifikasi beban kerja pegawai. Model ini menggunakan distribusi probabilitas untuk menentukan kemungkinan suatu pegawai masuk ke dalam kategori tertentu berdasarkan data yang telah diberikan. Hasil ini juga memberikan wawasan bahwa dalam dataset yang digunakan, mayoritas pegawai berada dalam kategori beban kerja rendah, sementara pegawai dengan beban kerja tinggi jumlahnya lebih sedikit tetapi memiliki pola kerja yang lebih terstruktur dan seragam.

Daftar Pustaka

- [1] D. Fajriyani, A. Fauzi, M. Devi Kurniawati, A. Yudo Prakoso Dewo, A. Fahri Baihaqi, and Z. Nasution, "Tantangan Kompetensi SDM dalam Menghadapi Era Digital (Literatur Review)," *Jurnal Ekonomi Manajemen Sistem Informasi*, vol. 4, no. 6, pp. 1004–1013, 2023, doi: 10.31933/jemsi.v4i6.1631.
- [2] M. Xanderina, A. Aditya Nafil, and F. Jatmiko, "Analisis Manajemen Sumber Daya Manusia Instansi Negeri Era Digitalisasi Dengan Kecerdasan Buatan," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 4, pp. 4451–4456, 2024, doi: 10.36040/jati.v8i4.9952.
- [3] Safitri, "Penerapan Algoritma Naïve Bayes Untuk Penentuan Calon Penerimaan Beasiswa Pada Sd Negeri 6 Ketapang," *Jurnal Informatika dan Sistem Informasi*, vol. 06, no. 01, pp. 43–52, 2020.
- [4] T. A. Q. Putri, A. Triayudi, and R. T. Aldisa, "Implementasi Algoritma Decision Tree dan Naïve Bayes Untuk Klasifikasi Sentimen Terhadap Kepuasan Pelanggan Starbucks," *Journal of Information System Research (JOSH)*, vol. 4, no. 2, pp. 641–649, 2023, doi: 10.47065/josh.v4i2.2949.
- [5] H. A. R. Harpizon, R. Kurniawan, Iwan Iskandar, R. Salambue, E. Budianita, and F. Syafria, "Analisis Sentimen Komentar Di YouTube Tentang Ceramah Ustadz Abdul Somad Menggunakan Algoritma Naïve Bayes," *JNKTI (Jurnal Nasional Komputasi dan Teknologi Informasi)*, vol. 5, no. 1, pp. 131–140, 2022.
- [6] T. Al Kautsar, "Klasifikasi Multi Label Pada Jawaban Esai Menggunakan Algoritma Multi Label Knearest Neighbor (Mlkn)," *Repository.Uinjkt.Ac.Id*, 2023.
- [7] E. Poerwandono and J. Perwitosari, "Penerapan Data Mining Untuk Penilaian Kinerja Karyawan Di PT. Riksa Dinar Djaya Menggunakan Metode Naïve Bayes Classification (Edhy Poerwandono 1 , Faizal Joko Perwitosari 2) Penerapan Data Mining Untuk Penilaian Kinerja Karya Di PT Riksa Dinar Djaya Menggunakan Metode Naive Bayes Classification," *Jurnal Sains dan Teknologi*, vol. 5, no. 1, p. |pp, 2023.
- [8] H. F. Putro, R. T. Vlandari, and W. L. Y. Saptomo, "Penerapan Metode Naive Bayes Untuk Klasifikasi Pelanggan," *Jurnal Teknologi Informasi dan Komunikasi (TIKoSIN)*, vol. 8, no. 2, 2020, doi: 10.30646/tikomsin.v8i2.500.
- [9] A. A. N. Risal, N. I. Yusuf, A. B. Kaswar, and F. Adiba, "Penerapan data mining dalam mengklasifikasikan tingkat kasus Covid-19 di Sulawesi Selatan menggunakan Algoritma Naive Bayes," *Indonesian Journal of Fundamental Sciences*, vol. 7, no. 1, pp. 18–28, 2021.
- [10] Ernianti Hasibuan and Elmo Allistair Heriyanto, "Analisis Sentimen Pada Ulasan Aplikasi Amazon Shopping Di Google Play Store Menggunakan Naive Bayes Classifier," *Jurnal Teknik dan Science*, vol. 1, no. 3, pp. 13–24, 2022, doi: 10.56127/jts.v1i3.434.
- [11] J. Xu, Y. Zhang, and D. Miao, "Three-way confusion matrix for classification: A measure driven view," *Information Sciences*, vol. 507, pp. 772–794, 2020.
- [12] C. A. Pamungkas and W. W. Widiyanto, "Klasifikasi Indeks Pembangunan Manusia Di Indonesia Tahun 2022 Dengan Support Vector Machine," *Jurnal ilmiah Sistem Informasi dan Ilmu Komputer*, vol. 2, no. 3, pp. 139–145, 2023, doi: 10.55606/juisik.v3i1.407.
- [13] E. Yuniar, D. S. Utsalinah, and D. Wahyuningsih, "Implementasi Scrapping Data Untuk Sentiment Analysis Pengguna Dompot Digital dengan Menggunakan Algoritma Machine Learning," *Jurnal Janitra Informatika dan Sistem Informasi*, vol. 2, no. 1, pp. 35–42, 2022, doi: 10.25008/janitra.v2i1.145.

- [14] K. Anwar, “Analisa sentimen Pengguna Instagram Di Indonesia Pada Review Smartphone Menggunakan Naive Bayes,” *KLIK: Kajian Ilmiah Informatika dan Komputer*, vol. 2, no. 4, pp. 148–155, 2022, doi: 10.30865/klik.v2i4.315.
- [15] A. S. A. Herdianti Darwis, Nugraha Wanaspati, “Support Vector Machine Untuk Analisis Sentimen Masyarakat Terhadap Penggunaan Antibiotik Di Indonesia,” *Indonesian Journal Of Computer Science*, vol. 12, pp. 2196–2206, 2023.