

OPTIMASI MANAJEMEN PERSEDIAAN DENGAN KLASIFIKASI NAIVE BAYES DI TOKO BAJA MANDIRI

Prilanisa Zhahiran Herlambang¹, Nining Rahaningsih², Irfan Ali³

^{1,2} Komputerisasi Akuntansi, STMIK IKMI Cirebon

³ Program Studi Rekayasa Perangkat Lunak, STMIK IKMI Cirebon

Jl. Perjuangan No. 10 B, Majasem, Cirebon Jawa Barat Indonesia

prilanisazhahiran@gmail.com

ABSTRAK

Pengelolaan persediaan merupakan aspek penting dari operasional perusahaan, khususnya di Toko Baja Mandiri yang menjual perlengkapan bahan bangunan berbasis baja. Pengelolaan yang tidak efektif dapat menyebabkan kekurangan pasokan atau penumpukan barang, sehingga berpotensi mengakibatkan kerugian finansial. Penelitian ini bertujuan untuk mengembangkan sistem klasifikasi status persediaan menggunakan algoritma Naive Bayes dengan data jumlah persediaan, harga satuan, dan total harga barang selama periode satu bulan. Pendekatan *Knowledge Discovery in Databases* (KDD) diterapkan, dimulai dari seleksi data yang relevan, diikuti oleh tahap praproses untuk membersihkan dan mentransformasi data agar siap digunakan. Setelah melalui tahap transformasi untuk menetapkan atribut yang berguna, data dibagi menjadi set pelatihan dan pengujian. Algoritma Naive Bayes diterapkan untuk mengklasifikasikan barang menjadi kategori “STOCK” (perlu *restock*) dan “READY” (*stock* aman). Evaluasi menunjukkan akurasi model sebesar 92.59%, dengan *recall* 100% untuk kategori “READY” dan *precision* 27.27%. Ketidakseimbangan data target, dengan 96.8% dalam kategori “STOCK” dan 3.2% kategori “READY”, memengaruhi performa model, terutama dalam *precision*. Sistem ini terbukti membantu pemilik toko mengoptimalkan pengelolaan stok, meminimalkan risiko kekosongan, dan mendukung pengambilan keputusan strategis secara lebih efektif dengan proses yang terstruktur dan andal.

Kata kunci : Klasifikasi, *Knowledge Discovery in Databases* (KDD), Manajemen Persediaan, Naive Bayes, Persediaan, RapidMiner.

1. PENDAHULUAN

Pada masa sekarang ini perkembangan teknologi dan komunikasi dari waktu ke waktu dirasakan semakin meningkat pesat, yang mendorong penggunaan dan pemanfaatan perkembangan teknologi di berbagai bidang dan aspek kehidupan[1].

Salah satu pengaruhnya yang signifikan terlihat dalam aspek bisnis, terutama dalam hal pengelolaan persediaan barang. Sistem pengelolaan persediaan barang menjadi bagian penting dari operasional perusahaan, karena jika gagal akan mempunyai pengaruh yang sangat besar terkait masa depan keberlangsungan perusahaan.

Toko Baja Mandiri merupakan sebuah usaha yang bergerak dibidang penjualan bahan bangunan khususnya baja. Dalam operasionalnya, Toko Baja Mandiri menghadapi tantangan besar dalam mengelola persediaan barang. Persediaan barang dikelola dengan memperhatikan keseimbangan antara kekurangan dan kelebihan *stock* selama masa perencanaan[2].

Jika jumlah *stock* tidak dikelola dengan tidak benar, ada risiko terjadinya kekurangan barang ketika pelanggan membutuhkannya. Di sisi lain, *stock* yang berlebihan justru dapat menambah biaya penyimpanan dan meningkatkan risiko kerusakan barang selama penyimpanan.

Tantangan lain yang sering dihadapi toko adalah pencatatan *stock* secara manual. Proses manual membutuhkan waktu cukup lama untuk menemukan informasi yang diperlukan, yang pada akhirnya

menurunkan efisiensi. Selain itu, risiko kesalahan manusia dalam pencatatan atau pengolahan data lebih tinggi dalam proses manual ini, yang bisa berdampak pada kelancaran distribusi barang[3].

Untuk menghadapi tantangan ini, Toko Baja Mandiri perlu mengadopsi cara-cara baru dalam mengelola persediaan barang. Salah satu solusi yang dapat diterapkan adalah pemanfaatan metode Naive Bayes dalam data *mining*.

Naive Bayes adalah metode klasifikasi probabilistik sederhana yang bekerja dengan menghitung sejumlah probabilitas berdasarkan frekuensi dan kombinasi nilai dari dataset yang tersedia. Algoritma ini memanfaatkan teorema Bayes sebagai dasar perhitungannya dan mengasumsikan bahwa seluruh atribut bersifat independen atau tidak saling bergantung, dengan syarat nilai pada variabel kelas telah ditentukan[4].

Penelitian sebelumnya yang dilakukan oleh Gitty Putri Asri dengan judul Klasifikasi Produk Persediaan pada PT HMS Kompresindo Sukses Menggunakan Algoritma Naive Bayes memaparkan tingkat akurasi dengan klasifikasi menggunakan metode Naive Bayes memperoleh nilai *Accuracy* 88.89%, *Precision* 81.82%, dan *Recall* 81.82%. Hasil tersebut menunjukkan kecenderungan klasifikasi sehingga dengan metode yang digunakan mampu memberikan dalam mengambil keputusan pada perusahaan[5].

Sementara itu, penelitian yang dilakukan oleh Putra dengan judul Sistem Predeksi Persediaan

Barang di Mini Market Mars Menggunakan Algoritma Naive Bayes Classifier dengan Pemrograman Python melaporkan akurasi sebesar 75,7%, dengan precision 64%, recall 75,7%, dan F1 score 69,3%. Hasil ini menunjukkan bahwa model Naive Bayes dapat digunakan untuk memprediksi stock barang di masa mendatang[6].

Naive Bayes dinilai efektif karena strukturnya sederhana dan memanfaatkan atribut-atribut independen untuk menyelesaikan masalah[7].

Berdasarkan penjelasan tersebut, penelitian ini menerapkan penggunaan Naive Bayes dengan tujuan agar Toko Baja Mandiri dapat mengklasifikasikan status persediaan barang dengan lebih akurat, sehingga dapat mengoptimalkan manajemen stock, mengurangi risiko kekurangan atau kelebihan barang, serta meningkatkan efisiensi dalam pengambilan keputusan terkait persediaan barang.

2. TINJAUAN PUSTAKA

2.1. Data Mining

Data Mining adalah proses mengumpulkan data berukuran besar, kemudian mengekstraksi data tersebut menjadi informasi yang nantinya dapat digunakan. Data mining merupakan teknik menggabungkan teknik analisis data dan menemukan pola-pola yang penting pada data. Data mining tersebut akan menjadi tolak ukur ataupun acuan untuk mengambil keputusan.

Tidak semua proses pencarian informasi dapat dikategorikan sebagai data mining. Data mining dimulai dengan menyeleksi data dari sumber hingga menjadi data target. Selanjutnya, tahap preprocessing dilakukan untuk meningkatkan kualitas data, diikuti dengan tahap transformasi, penerapan teknik data mining, serta interpretasi dan evaluasi. Proses ini nantinya akan menghasilkan pengetahuan baru yang diharapkan dapat memberikan manfaat yang lebih optimal[8].

2.2. Algoritma Naive Bayes

Algoritma Naive Bayes merupakan algoritma klasifikasi yang berdasarkan teorema Bayes dengan dugaan bahwa setiap fitur yang digunakan untuk klasifikasi yaitu independen. Secara sederhana Algoritma Naive Bayes adalah algoritma yang menggunakan pendekatan statistik untuk mengambil sebuah keputusan, dapat dijelaskan kembali Algoritma Naive Bayes adalah metode klasifikasi yang terstruktur pada Teorema Bayes. Proses klasifikasi dengan menerapkan metode probabilitas dan statistik yang dipaparkan oleh ilmuwan Inggris yaitu Thomas Bayes. Meskipun konsep ini bertentangan dengan kehidupan nyata perlu diketahui atribut tidak sama penting atau independen, tetapi Naive Bayes akan memperlihatkan performa yang mampu bersaing dengan algoritma klasifikasi yang terkenal, decision tree dan neural network[9].

Teorema Bayes memiliki bentuk umum sebagai berikut :

$$P(H|X) = \frac{P(X|H) \cdot P(H)}{P(X)} \dots\dots\dots (1)$$

Di mana:

X : Data dengan class yang belum diketahui

H : Hipotesis data yang merupakan suatu class spesifik

P(H|X) : Probabilitas hipotesis H berdasarkan kondisi X (posteriori probabilitas)

P(H) : Probabilitas hipotesis H (prior probabilitas)

P(X|H) : Probabilitas X berdasarkan kondisi pada hipotesis H

P(X) : Probabilitas X[10]

2.3. Persediaan

Secara umum istilah persediaan barang dipakai untuk menunjukkan barang-barang yang dimiliki untuk dijual kembali atau digunakan untuk memproduksi barang-barang yang akan dijual.

Pengertian menurut Ikatan Akuntan Indonesia, Persediaan merupakan suatu aset :

- a. Tersedia untuk dijual dalam kegiatan usaha normal
- b. Dalam proses produksi ataupun dalam perjalanan
- c. Dalam bentuk bahan atau perlengkapan (supplies) untuk digunakan dalam proses produksi atau pemberian jasa[11]

Persediaan merupakan elemen terpenting dalam kegiatan produksi yang dilakukan oleh seluruh perusahaan di dunia. Selain itu persediaan menjadi salah satu masalah yang perlu diperhatikan dalam kaitannya dengan kegiatan proses produksi, biaya serta distribusi barang-barang baik itu bahan baku, barang dalam proses atau barang setengah jadi, ataupun barang jadi[12].

2.4. Klasifikasi

Klasifikasi adalah metode untuk mengembangkan model atau fungsi yang menerapkan model matematika pada data atau konsep untuk membuat inferensi statistik untuk data yang tidak diketahui[13].

Klasifikasi bisa dibuat dengan berbagai metode, baik dengan cara manual oleh manusia maupun sama bantuan teknologi, seperti menggunakan algoritma machine learning. Secara manual, klasifikasi dilakukan dengan cara mengidentifikasi karakteristik atau pola tertentu dalam data dan mengambil keputusan berdasarkan pengalaman atau pengetahuan yang dimiliki. Sementara itu, dengan bantuan teknologi, klasifikasi sering kali menggunakan pendekatan yang lebih terstruktur dan otomatis. Ini dapat mencakup penggunaan algoritma seperti Naive Bayes, Random Forests, dan lainnya. Algoritma-algoritma ini mampu belajar dari data yang ada dan menggunakan informasi tersebut untuk membuat klasifikasi atau prediksi terhadap data baru[14]

2.5. Rapid miner

Aplikasi Rapid miner digunakan untuk melakukan penelitian, pendidikan, pelatihan, pembuatan *prototype* dengan cepat, mendukung proses pembelajaran mesin termasuk persiapan data, visualisasi hasil, validasi dan pengoptimalan[15].

3. METODE PENELITIAN

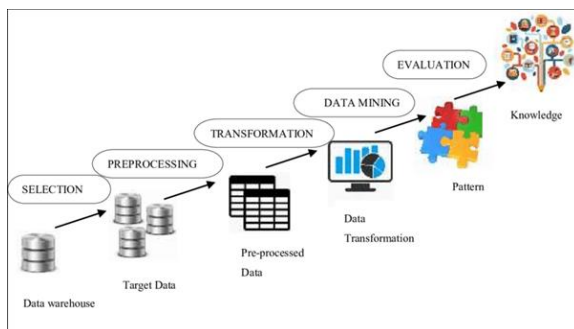
Dalam penelitian ini melakukan beberapa tahapan dalam proses pengumpulan data dan pengolahan data mining klasifikasi dengan menggunakan metode Naive Bayes.

3.1. Metode pengumpulan data

- Proses wawancara dengan menanyakan langsung kepada responden atau narasumber untuk mendapatkan sebuah informasi data.
- Metode Observasi dengan mengunjungi langsung ke lokasi atau tempat penelitian, guna mengumpulkan dan mengetahui informasi data yang didapat lebih akurat. Observasi yang dilakukan pada penelitian ini mendatangi secara langsung ke Toko Baja Mandiri.

3.2. Tahapan penelitian

Tahapan penelitian ini menggunakan proses tahapan *Knowledge Discovery Database* (KDD).



Gambar 1. Tahapan Penelitian KDD

Tahapannya *Knowledge Discovery in Databases* (KDD) sebagai berikut :

- Data**
Data merupakan sekumpulan data operasional yang diperlukan sebelum dilakukan tahap penggalian informasi dalam *Knowledge Discovery Database* (KDD) dimulai.
- Data Cleaning**
Proses data *cleaning* adalah proses dengan melakukan pembersihan data yang bertujuan untuk

menghilangkan data yang tidak memiliki nilai (*null*), data yang salah input, data yang tidak relevan, duplikat data dan data yang tidak konsisten.

c. Data transformation

Data *transformation* dilakukan dengan memberikan inisialisasi terhadap data yang memiliki nilai nominal menjadi bernilai numerik

d. Data Mining

Pada fase ini menerapkan algoritma atau metode pencarian pengetahuan. Algoritma Naïve Bayes diterapkan untuk melakukan klasifikasi data. Prosesnya meliputi:

- Menghitung peluang awal (Probabilitas *Prior*) untuk setiap kategorinya berdasarkan jumlah data yang tersedia.
- Menghitung peluang setiap fitur (*likelihood*) dalam masing-masing kategori.
- Menggunakan rumus Naïve Bayes untuk menghitung kemungkinan suatu data termasuk dalam kategori tertentu.
- Menentukan hasil klasifikasi dengan memilih kategori yang memiliki probabilitas tertinggi.

e. Evaluation

Pada tahap evaluasi, akan diketahui apakah hasil daripada tahap data mining dapat menjawab tujuan yang telah ditetapkan. Untuk itu akan dilakukan profilisasi pada setiap cluster yang telah terbentuk, untuk diketahui karakteristik pada kelompok tersebut. Disamping itu untuk diketahui kesesuaian dengan jalur perminatan akan dilakukan analisis lebih lanjut untuk dihubungkan dengan atribut perminatan, Sehingga diharapkan mendapatkan informasi atau pola yang berguna sebagai acuan pemutakhiran data.

f. Knowledge

Tahap terakhir dari proses data mining adalah bagaimana memformulasikan keputusan atau aksi dari hasil analisis yang didapat[16].

4. HASIL DAN PEMBAHASAN

4.1. Sumber Data

Sumber data ini merupakan data persediaan barang Toko Baja Mandiri periode bulan Juni 2024. Memiliki 8 atribut yaitu terdiri dari NO, ID BARANG, NAMA BARANG, SATUAN, STOCK, TOTAL HARGA, HARGA SATUAN, dan STATUS. Dengan jumlah data sebanyak 361.

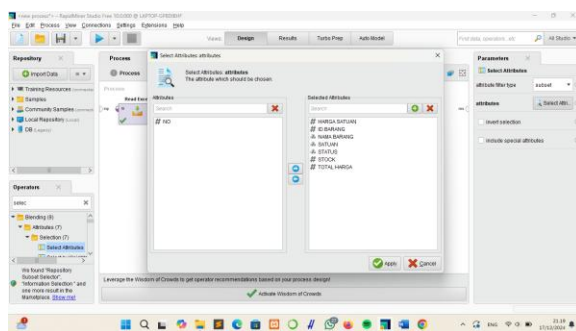
Tabel 1. Tabel Data Persediaan Barang Toko Baja Mandiri sebelum *Pre-processing*

No	Id Barang	Nama Barang	Stock	...	Status
1	201911005	A. Pinggir (035) WH DC	1		stock
2	201911015	Pipa Bulat 8 mm CA MF	3		ready
3	201912004	Holo Galvanis 4 x 4 0,5	5		stock
4	202001016	List Shadowline	16		stock
5.	202002011	Z 3" (0414) Ckt Forta	4		stock
	...	...	...		...
361	201911002	T. Sliding (9055) WH FORTA	9		stock

Tabel 1. adalah data persediaan Toko Baja Mandiri yang digunakan.

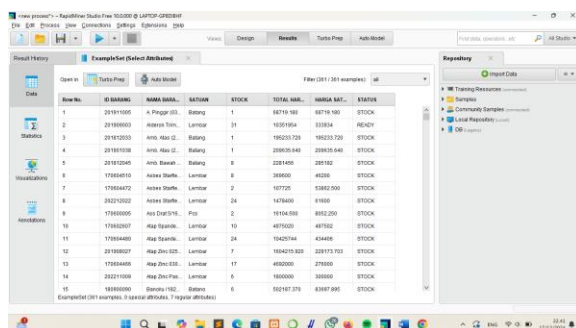
4.2. Pembersihan Data

Pembersihan Data (*Data Cleaning*) merupakan langkah penting untuk memastikan bahwa data yang digunakan dalam analisis benar-benar berkualitas. Proses ini melibatkan pengecekan dan perbaikan data yang tidak lengkap, duplikat, atau tidak relevan. Pembersihan data melibatkan identifikasi dan penghapusan data yang tidak diperlukan, mengatasi nilai yang hilang atau tidak valid, serta memastikan bahwa format dan struktur data seragam. Dengan pembersihan data yang tepat, data yang digunakan untuk analisis menjadi lebih siap, sehingga meningkatkan keakuratan hasil dan keandalan model yang dihasilkan.



Gambar 2. Tahapan Select Attribute dan Select Subset

Pada Gambar 2. Merupakan proses *select attribute* dan *select subset*. Untuk *attribute filter type* diganti dengan menggunakan *subset*, sedangkan pada *select subset* adalah dengan menghilangkan atribut yang tidak relevan yaitu *NO*, yang tidak memberikan informasi penting untuk analisis persediaan barang. Penelitian ini berfokus pada persediaan barang di Toko Baja Mandiri sehingga atribut-atirbut yang digunakan dalam analisis meliputi HARGA SATUAN, ID BARANG, NAMA BARANG, SATUAN, STATUS, STOCK, dan TOTAL HARGA.

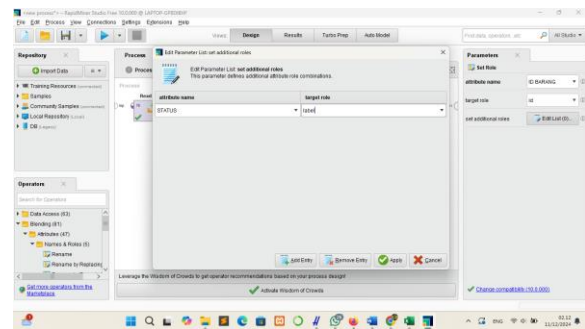


Gambar 3. Exampleset Select Attribute

Gambar 3. menunjukkan hasil dari pembersihan data yang telah dilakukan.

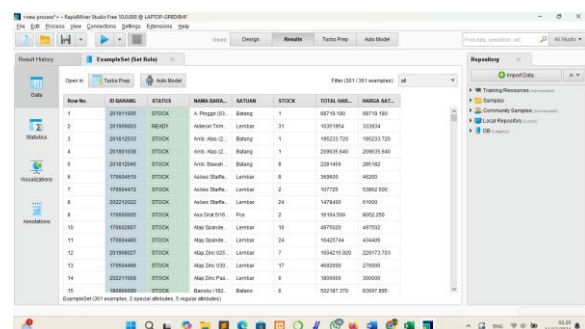
4.3. Transformasi Data

Transformasi data merupakan proses mengubah data mentah menjadi format yang lebih sesuai untuk analisis atau pemodelan. Pada tahap ini, data tidak mengalami perubahan nilai tetapi diubah perannya agar sesuai dengan kebutuhan algoritma. Perubahan ini dilakukan melalui *operator Set Roles*, yang mengatur atribut-atribut menjadi *label* dan *id*.



Gambar 4. Penentuan Peran Atribut (*Set Roles*)

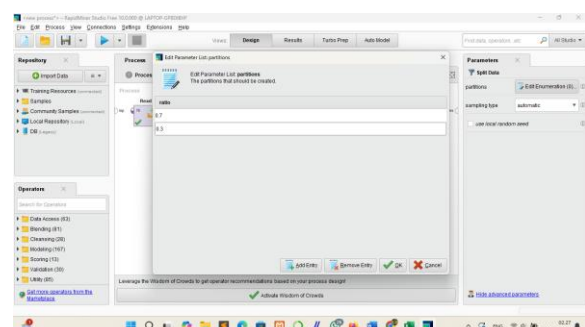
Pada Gambar 4. merupakan penentuan peran atribut (*Set Roles*) dengan menetapkan atribut STATUS sebagai label (*target*) dan ID BARANG sebagai *id*, sementara atribut lainnya berperan sebagai *regular*. Maka hasilnya akan menjadi sebagai berikut.



Gambar 5. Hasil ExampleSet (*Set Role*)

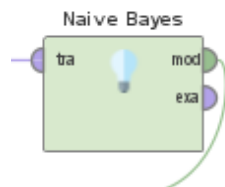
4.4. Data Mining

Pada tahapan data mining ini merupakan inti dari proses KDD, Dimana algoritma Naive Bayes digunakan untuk membangun model klasifikasi. Dataset dibagi menjadi dua bagian menggunakan operator *Split Data* dengan rasio 70:30, dimana 70% data digunakan sebagai *training set* dan 30% sebagai *test set*.



Gambar 6. Melakukan Split Data

Data *training* yang digunakan sebanyak 252 digunakan untuk melatih model agar dapat mengenali pola hubungan antara atribut ID BARANG dan atribut lainnya terhadap STATUS sebagai label atau targetnya. Sedangkan sebanyak 109 data uji digunakan untuk mengukur performa model dalam memprediksi STATUS secara akurat berdasarkan pola yang sudah dipelajari. Pengujian ini bertujuan untuk memastikan bahwa model mampu mengklasifikasikan data baru dengan tingkat akurasi yang memadai.

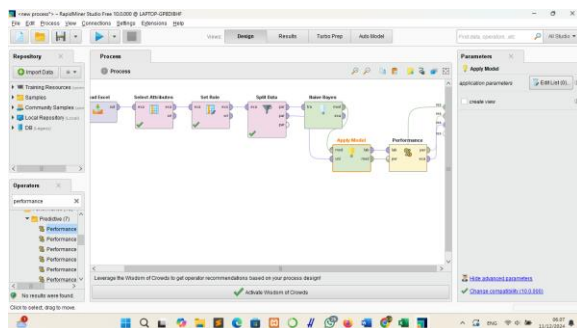


Gambar 7. Operator Naive Bayes

Kemudian pada Gambar 7. Operator Naive Bayes diterapkan pada data *training* untuk membangun model klasifikasi berdasarkan probabilitas kondisi dari masing-masing atribut terhadap STATUS. Setelah model terbentuk, operator.

4.5. Evaluasi

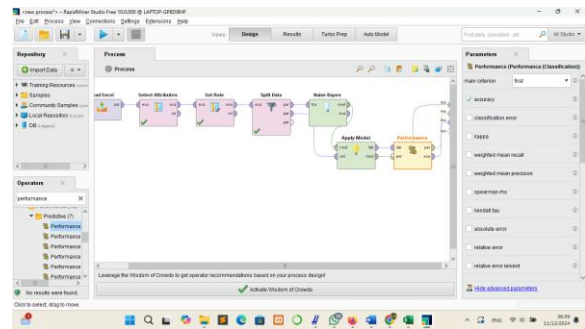
Tahap dilakukan untuk mengukur akurasi model yang dihasilkan. Operator *Apply Model* diterapkan pada data pengujian untuk menghasilkan prediksi, yang kemudian dievaluasi menggunakan operator *Performance*.



Gambar 8. Apply model

Gambar 8. merupakan tahapan penerapan operator *Apply Model* guna menerapkan model yang telah dilatih pada dataset baru. Operator ini menghasilkan prediksi atau klasifikasi berdasarkan pola yang dipelajari oleh model sebelumnya.

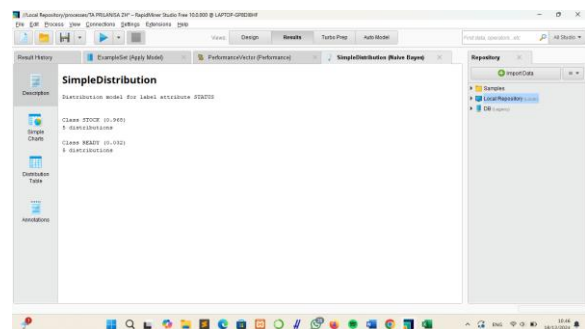
Pada gambar 9 ini menunjukkan tahapan terakhir yaitu *Performance* yang digunakan untuk mengevaluasi kinerja dari model dengan menghitung metrik seperti akurasi, presisi, *recall*, dan *F1-score*, sehingga memudahkan penilaian efektivitas model dalam analisis data.



Gambar 9. Penggunaan Operator performance

4.6. Distribusi Data

Distribusi data dalam model Naive Bayes menggambarkan sebaran nilai dari tabel atau atribut target yang digunakan untuk klasifikasi. Distribusi ini sangat penting karena mempengaruhi kinerja model, terutama jika terjadi ketidakseimbangan data (*imbalanced data*).



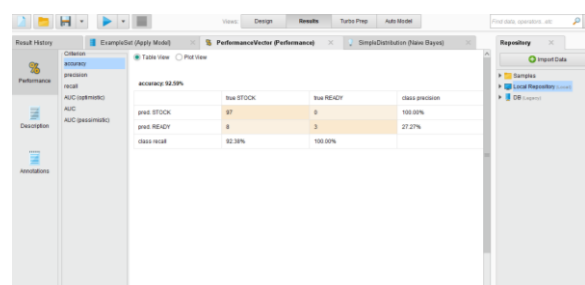
Gambar 10. Hasil Simple Distribution Naive Bayes

Berdasarkan Gambar 10. hasil distribusi data pada model klasifikasi menggunakan algoritma Naive Bayes menunjukkan adanya ketidakseimbangan data pada atribut target STATUS. Kelas STOCK memiliki proporsi sebesar 96.8%, sedangkan kelas READY hanya sebesar 3.2%. Ketidakseimbangan ini mengindikasikan bahwa data didominasi oleh kelas STOCK, sementara kelas READY memiliki jumlah data yang sangat sedikit.

4.7. Hasil Performa Vector

4.7.1. Hasil Accuracy

Hasil *Accuracy* didefinisikan sebagai tingkat kedekatan antara nilai prediksi.



Gambar 11. Hasil Accuracy

$$accuracy = \frac{TP + TN}{TP + TN + FP + FN} \times 100\%$$

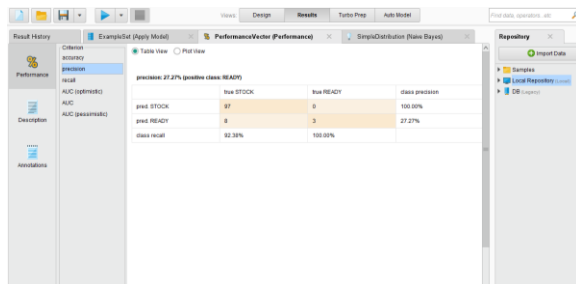
$$accuracy = \frac{97 + 3}{97 + 3 + 8 + 0} \times 100\%$$

$$accuracy = 0,9259 = 92.59 \%$$

Dalam perhitungan tersebut menunjukkan hasil pengujian prediksi menggunakan Algoritma Naive Bayes dengan nilai akurasi sebesar 92.59%.

4.7.2. Hasil Akurasi Precision

Akurasi *precision* merupakan tingkatan ketepatan antara data yang diminta oleh pengguna dengan jawaban yang diberikan oleh sistem.



	prec STOCK	prec READY	class precision
prec STOCK	0.2727	0	100.00%
prec READY	0	0.2727	27.27%
class recall	100.00%	100.00%	

Gambar 12. Hasil Precision

$$precision = \frac{TP}{TP + FP} \times 100\%$$

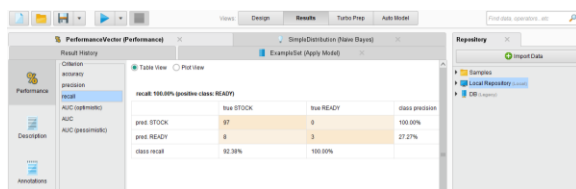
$$precision = \frac{3}{3 + 8} \times 100\%$$

$$precision = 0,2727 = 27.27\%$$

Dalam perhitungan tersebut dapat diketahui nilai akurasi *precision* untuk kelas *READY* adalah 27.27%, yang berarti hanya 27.27% dari data yang diprediksi sebagai *READY* benar-benar termasuk dalam kategori tersebut.

4.7.3. Akurasi Recall

Akurasi *Recall* merupakan tingkatan keberhasilan sistem dalam mendapatkan kembali data. Dengan jumlah data dapat diketahui nilai Akurasi *Recall* dengan hasil prediksinya yaitu 100%.



	prec STOCK	prec READY	class precision
prec STOCK	0.2727	0	100.00%
prec READY	0	0.2727	27.27%
class recall	100.00%	100.00%	

Gambar 13. Hasil Akurasi Recall

$$recall = \frac{TP}{TP + FN} \times 100\%$$

$$recall = \frac{3}{3 + 0} \times 100\%$$

$$recall = 100 = 100\%$$

Gambar 13. menunjukkan hasil akurasi *Recall* sebesar 100%. Artinya, model mampu mengenali semua data yang termasuk dalam kelas "*READY*" dengan benar tanpa ada kesalahan klasifikasi.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian diatas diperoleh kesimpulan bahwa klasifikasi *stock* barang di Toko Baja Mandiri dapat diterapkan menggunakan data mining dengan melakukan tata cara algoritma naive bayes. Melalui tahapan seperti pembersihan data terlebih dahulu kemudian transformasi atribut, pembagian *dataset*, hingga evaluasi performa model.

Hasil klasifikasi *stock* barang menggunakan Naïve Bayes menunjukkan akurasi sebesar 92.59%, yang memperlihatkan bahwa metode ini cukup andal dalam menentukan status *stock* barang di Toko Baja Mandiri; namun, masih terdapat ketidakseimbangan dalam *precision* dan *recall*, terutama pada kategori "*READY*", dengan *precision* hanya 27.27%, yang berarti masih ada kesalahan dalam mengklasifikasikan barang sebagai *READY* padahal sebenarnya bukan, sementara *recall* untuk kelas "*READY*" adalah 100%, yang berarti semua barang *ready* yang ada dalam dataset berhasil dikenali oleh model tanpa kesalahan. Dari 109 data uji, sekitar 101 barang diklasifikasikan dengan benar sebagai "*STOCK*" dan 3 barang sebagai "*READY*". Untuk meningkatkan akurasi, penelitian selanjutnya disarankan menambah jumlah data latih, menggunakan fitur tambahan seperti tren permintaan barang, mencoba model lain seperti *Random Forest* atau *Neural Network*, serta mengoptimalkan teknik *oversampling* guna menyeimbangkan data. Selain itu, implementasi model dalam sistem manajemen *stock* berbasis aplikasi dapat membantu pengambilan keputusan yang lebih cepat dan akurat, sehingga mengurangi risiko kelebihan atau kekurangan *stock* yang dapat berdampak pada efisiensi operasional toko.

DAFTAR PUSTAKA

- [1] A. Arisusanto, N. Suarna, and G. Dwilestari, "Analisa Klasifikasi Data Harga Handphone Menggunakan Algoritma Random Forest Dengan Optimize Parameter Grid," *J. Teknol. Ilmu Komput.*, vol. 1, no. 2, pp. 43–47, 2023, doi: 10.56854/jtik.v1i2.51.
- [2] M. Sabrina, A. F. Novendri, R. B. Utami, D. Felisha, A. Ekawati, and N. R. Faizah, "Analisis Prosedur Pengendalian Persediaan Barang Di Toko Alfamart," *J. Bisnis Corp.*, vol. 9, no. 1, pp. 61–66, 2024, doi: 10.46576/jbc.v9i1.4560.
- [3] A. al afif fadhil Aqilah, S. Bustamin, and S. Sultan sahrir, "Sistem Informasi Manajemen Persediaan Berbasis Web di CV. Makmur Sejahtera Palopo," *J. Process.*, vol. 18, no. 2, 2023, doi: 10.33998/processor.2023.18.2.1385.
- [4] E. Martantoh and N. Yanhi, "Implementasi Metode Naïve Bayes Untuk Klasifikasi Karakteristik Kepribadian Siswa Di Sekolah MTS Darussa'adah Menggunakan Php Mysql," *J. Teknol. Sist. Inf.*, vol. 3, no. 2, pp. 166–175, 2022, doi: 10.35957/jtsi.v3i2.2896.
- [5] G. P. Asri and F. R. Doni, "Klasifikasi Produk Persediaan pada PT HMS Kompresindo Sukses Menggunakan Algoritma Naive Bayes," vol. 16,

- no. 2, pp. 1–9, 2024.
- [6] P. W. P. Putra, D. Irawan, A. T. Martadinata, and A. Heryati, “Sistem Prediksi Persediaan Barang Di Mini Market Mars Menggunakan Algoritma Naïve Bayes Classifier Dengan Pemrograman Phyton,” *JUTIM (Jurnal Tek. Inform. Musirawas)*, vol. 7, no. 2, pp. 148–159, 2022, doi: 10.32767/jutim.v7i2.2101.
- [7] P. Ayu, W. Purnama, and T. A. Putra, “Klasifikasi Penjualan Produk Menggunakan Algoritma Naive Bayes pada Konter HP Bayu Cell,” *Remik Ris. dan E-Jurnal Manaj. Inform. Komput.*, vol. 8, no. 1, pp. 286–292, 2024, [Online]. Available: <http://doi.org/10.33395/remik.v8i1.13207>
- [8] R. Aztin and K. A. Musodo, “Penerapan Text Mining Dengan Algoritma Naïve Bayes Untuk Mengklasifikasikan Sentimen Rakyat Terhadap Minyak Goreng Subsidi Pemerintah,” *Semin. Nas. Mhs. Fak. ...*, no. September, pp. 645–652, 2022, [Online]. Available: <http://senafiti.budiluhur.ac.id/index.php/senafiti/article/view/347%0Ahttp://senafiti.budiluhur.ac.id/index.php/senafiti/article/download/347/76>
- [9] E. Andrian and A. R. Isnain, “Analisis Sentimen Masyarakat Terhadap Tiktok Shop di Twitter Menggunakan Metode Naive Bayes Classifier,” *J. Media Inform. Budidarma*, vol. 8, no. 2, p. 788, 2024, doi: 10.30865/mib.v8i2.7530.
- [10] R. Rachman and R. N. Handayani, “Klasifikasi Algoritma Naive Bayes Dalam Memprediksi Tingkat Kelancaran Pembayaran Sewa Teras UMKM,” *J. Inform.*, vol. 8, no. 2, pp. 111–122, 2021, doi: 10.31294/ji.v8i2.10494.
- [11] S. Ayu, Anggi; Nugroho, Nur; Muhammad, “Penerapan Data Mining Untuk Menentukan Persediaan Barang Pada PT. Deli Food Menggunakan Metode K-Means Anggi Ayu Ningtiyas *, Nurcahyo Budi Nugroho**, Muhammad Syaifuddin*** * Program Studi Sistem Informasi, STMIK Triguna Dharma ** Program Studi Sistem In,” vol. 4, no. 7, 2021, [Online]. Available: <https://ojs.trigunadharma.ac.id/>
- [12] F. Azhima, Refiza, and Z. Siregar, “Bahan Baku Menggunakan Klasifikasi Abc Dan Metode,” vol. 04, no. 02, 2023, doi: 10.54123/vorteks.v4i2.309.
- [13] A. Julianto and S. Andayani, “Penerapan Data Mining Untuk Klasifikasi Produk Terlaris Menggunakan Algoritma Naive Bayes Pada Bengkel Motor,” *JuSiTik J. Sist. dan Teknol. Inf. Komun.*, vol. 7, no. 2, pp. 50–58, 2024, doi: 10.32524/jusitik.v7i2.1148.
- [14] T. Ponda and D. Kusumaningsih, “IMPLEMENTASI DATA MINING UNTUK KLASIFIKASI KELULUSAN MENGGUNAKAN METODE NAIVE BAYES BERBASIS WEB PADA SISWA IMPLEMENTATION OF DATA MINING FOR GRADUATION CLASSIFICATION USING THE WEB-BASED NAÏVE BAYES METHOD IN,” vol. 3, no. September, pp. 856–864, 2024.
- [15] Y. Mulyanto, F. Idifitriani, A. Wati, U. T. Sumbawa, D. Mining, and K. P. Tano, “Vol 7 No 2 , September 2024 KLASIFIKASI DATA MINING UNTUK PENENTUAN STUNTING,” vol. 7, no. 2, pp. 119–125, 2024.
- [16] Heliyanti Susana, “Penerapan Model Klasifikasi Metode Naive Bayes Terhadap Penggunaan Akses Internet,” *J. Ris. Sist. Inf. dan Teknol. Inf.*, vol. 4, no. 1, pp. 1–8, 2022, doi: 10.52005/jursistekni.v4i1.96.