

BAB IV

IMPLEMENTASI DAN PEMBAHASAN

4.1. Implementasi Sistem

Implementasi sistem merupakan tahap penting dalam proses pembangunan sistem informasi, karena pada tahap inilah seluruh rancangan dan perencanaan yang telah dibuat sebelumnya diterapkan dalam bentuk nyata. Secara umum, implementasi adalah pelaksanaan atau penerapan dari suatu konsep atau rencana untuk mencapai tujuan tertentu. Dalam konteks penelitian ini, implementasi merujuk pada penerapan sistem *data mining* yang dirancang untuk menganalisis data pelanggan guna mengetahui tingkat kepuasan mereka terhadap produk yang ditawarkan. Sistem ini menggabungkan dua metode analisis yang populer, yaitu *Naïve Bayes* dan *K-Means Clustering*. Metode *Naïve Bayes* digunakan untuk melakukan klasifikasi terhadap data pelanggan berdasarkan atribut-atribut tertentu seperti rasa, kemasan, pelayanan, dan kepuasan, sedangkan *K-Means Clustering* digunakan untuk mengelompokkan pelanggan ke dalam beberapa *cluster* berdasarkan kesamaan karakteristik. Dengan kombinasi kedua metode ini, sistem dapat menghasilkan informasi yang lebih lengkap dan akurat mengenai perilaku dan tingkat kepuasan pelanggan.

Tujuan utama dari implementasi sistem ini adalah untuk membantu pihak **Sweetdessert** dalam mengelola dan memahami data pelanggan secara lebih efisien dan mendalam. Melalui analisis data yang terstruktur, manajemen dapat mengetahui faktor-faktor yang paling memengaruhi kepuasan pelanggan serta mengidentifikasi kelompok pelanggan dengan preferensi yang serupa. Informasi ini sangat berharga untuk dijadikan dasar dalam pengambilan keputusan strategis, seperti peningkatan kualitas produk, perbaikan layanan, dan penyusunan strategi pemasaran yang lebih tepat sasaran. Dengan sistem ini, perusahaan tidak hanya dapat meningkatkan kepuasan pelanggan secara signifikan, tetapi juga memperkuat loyalitas pelanggan terhadap merek Sweetdessert. Implementasi sistem ini diharapkan menjadi langkah awal menuju pengelolaan bisnis yang lebih

berbasis data (*data-driven decision making*), sehingga perusahaan dapat bersaing lebih baik di tengah pasar yang semakin kompetitif.

4.2. Hasil Pengujian Menggunakan Rapid Miner

4.2.1. Implementasi Metode *Naive Bayes* dengan Aplikasi *Rapidminer*

Pada tahap ini pemodelan data, metode yang digunakan pada penelitian ini adalah probabilitas (prediksi) dengan menggunakan metode *Naive Bayes* Data yang telah dikumpulkan ditransformasi akan dikelola menggunakan probabilitas dan perhitungan jarak. Metode ini dapat digunakan dalam memprediksi peluang yang akan terjadi dimasa depan berdasarkan pengalaman dimasa sebelumnya sebagai bahan perbandingan.

Data yang akan diujikan dibagi menjadi dua yaitu data *training* dan data *testing* kemudian dianalisis dengan menggunakan *Software RapidMiner*. Data Kepuasan pelanggan mochi terlaris memiliki 80 *record*, untuk data *training* 64 *record* dan data *testing* memiliki 16 *record*.

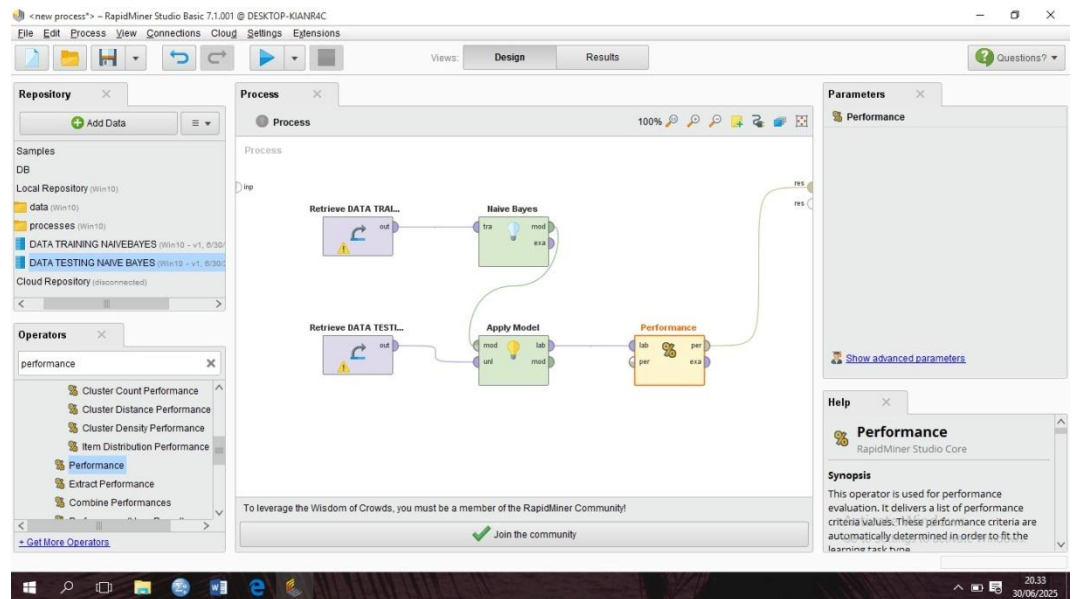
Pada tahap ini penulis mengambil data training 80% dari jumlah data keseluruhan dan 20% data yang digunakan untuk data testing. Maka diketahui jika data keseluruhan tersebut berjumlah 80 data.

$$\text{Data training} = 80\% \times 80 = 64$$

$$\text{Data testing} = 20\% \times 80 = 16$$

Atribut yang digunakan sebagai label adalah kepuasan pelanggan, sedangkan data *testing* yang digunakan diambil secara acak atau random dari 64 data yang telah ada. Selanjutnya lakukan importing data terlebih dahulu yaitu *data training* dan *data testing* yang tersedia di pc atau computer.

Jika kedua data atau file sudah selesai di *import* maka langkah selanjutnya *drop* kedua file ke dalam lembar kerja dan *drop operator* *Naive Bayes*, *Apply Model* dan *Performane* lalu hubungkan dengan data *training* dan data *testing* seperti gambar 4.1 dibawah ini



Gambar 4.1 Menghubungkan Operator

Setelah semua operator terhubung, selanjutnya klik icon *Run* pada *toolbar*. Selanjutnya *RapidMiner* akan menunjukkan hasil dari perhitungan yang telah dilakukan.

Criterion	accuracy
precision	
recall	
AUC (optimistic)	
AUC	
AUC (pessimistic)	

	true Puas	true Tidak Puas	class precision
pred. Puas	14	0	100.00%
pred. Tidak Puas	1	1	50.00%
class recall	93.33%	100.00%	

Gambar 4.2 Hasil dari perhitungan

Pada gambar 4.2 diatas tingkat *accuracy* dari *performace vector* adalah 93,75 %, *class precision* Puas 100,00% dan *class precision* Tidak Puas 50,00%, sedangkan untuk *class recall true* Puas 93,33% dan *class recall true* Tidak Puas 100,00%.

Secara umum precision, recall, dan accuracy dirumuskan sebagai berikut:

$$\begin{aligned} \text{Accuracy} &= \frac{TP+TN}{TP+FP+FN+TN} \\ &= \frac{14+1}{14+0+1+1} = \frac{15}{16} = 94 * 100\% = 94\% \end{aligned}$$

$$\begin{aligned} \text{Precision} &= \frac{TP}{TP+FP} \\ &= \frac{14}{14+0} = \frac{14}{14} = 1 * 100\% = 100\% \end{aligned}$$

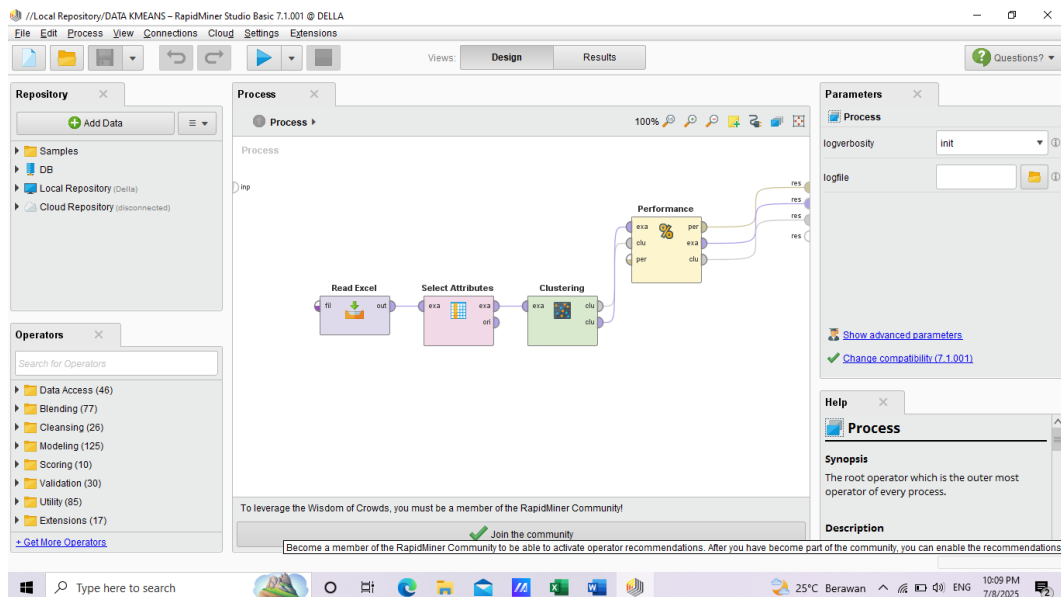
$$\begin{aligned} \text{Recall} &= \frac{TP}{TP+FN} \\ &= \frac{14}{14+1} = \frac{14}{15} = 0,93 * 100\% = 93\% \end{aligned}$$

$$\begin{aligned} \text{F1-Score} &= 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \\ &= 2 \times \frac{1 \times 93}{1 + 93} = 2 \times \frac{93}{94} = 98\% \end{aligned}$$

Setelah mengetahui hasil dari uji performa diatas, maka klasifikasi kepuasan pelanggan pada mochi terlaris dengan menggunakan metode *Naïve Bayes* dinyatakan Puas. Hal ini dikarenakan memiliki nilai akurasi yang tinggi yaitu mencapai 94%.

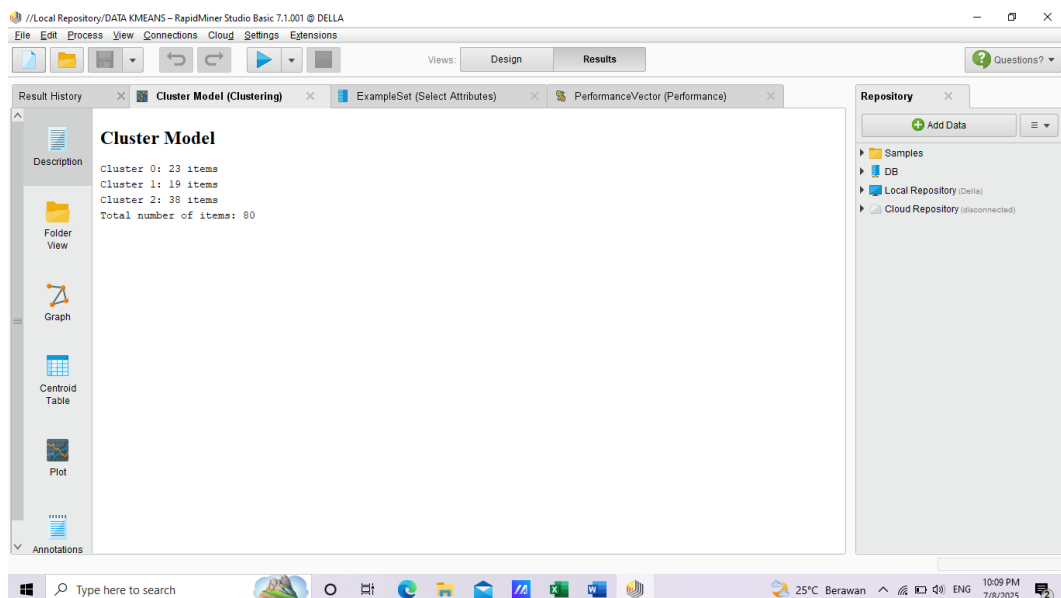
4.2.2. Implementasi Metode *K-Means Clustering* dengan Aplikasi *Rapidminer*

Implementasikan metode Metode *K-means clustering* pada data tersebut dengan menggunakan aplikasi *rapidminer*. Setelah semua operator yang kita butuhkan sudah berada dihalaman main proses selanjutnya hubungkan *read excel* dengan Dataset Kepuasan pelanggan, kemudian hubungkan ke *select attributes* lalu ke *clustering*, selanjutnya ke *Performance*, dan dapat dilihat pada gambar berikut.



Gambar 4.3 Susunan Antar Operator dan Data *K-means Clustering*

Setelah semua operator terhubung, selanjutnya klik icon *Run* pada *toolbar*. Selanjutnya *RapidMiner* akan menunjukkan hasil dari perhitungan yang telah dilakukan.



Gambar 4.4. Hasil Prediksi Data

Pada gambar diatas merupakan hasil perhitungan yang dilakukan dengan menggunakan *RapidMiner* dengan mencantumkan hasil prediksi pada data kepuasan pelanggan . Hasil dari metode *K-means clustering* menunjukkan bahwa diperoleh tiga kelompok pelanggan dengan jumlah anggota sebagai berikut:

1. Klaster 0 : 23 pelanggan (R3, R6, R9, R12, R15, R18, R24, R27, R30, R33, R36,R39,R42,R45,R48,R51,R57,R60,R63,R66,R69,R72, dan R78).
2. Klaster 1 : 19 pelanggan (R4, R8, R11, R21, R23, R26, R32, R35, R43, R49, R50, R54, R56, R67, R70, R73, R75, R77, dan R80).
3. Klaster 2 : 38 pelanggan (R1, R2, R5, R7, R10, R13, R14, R16, R17, R19, R20, R22, R25, R28, R29, R31, R34, R37, R38, R40, R41, R44, R46, R47, R52, R53, R55, R58, R59, R61, R62, R64, R65,R68, R71, R74, R76, dan R79).

Sebagian besar pelanggan masuk dalam Klaster 2 (paling puas), yang menunjukkan bahwa layanan atau produk saat ini secara umum mendapat tanggapan positif. Masih ada pelanggan dalam Klaster 0 dan 1, yang berarti ada peluang untuk meningkatkan kepuasan mereka