

BAB II

LANDASAN TEORI

2.1. Konsep Kepuasan Pelanggan

Kepuasan pelanggan merupakan salah satu aspek terpenting dalam manajemen bisnis modern. Dalam era persaingan yang semakin ketat, perusahaan dituntut untuk tidak hanya menawarkan produk atau layanan berkualitas, tetapi juga untuk memahami dan memenuhi harapan pelanggan. Kepuasan pelanggan tidak hanya berpengaruh pada loyalitas pelanggan, tetapi juga pada reputasi perusahaan dan kinerja keuangan. Dalam tulisan ini, kita akan membahas konsep kepuasan pelanggan, faktor-faktor yang mempengaruhinya, model-model yang ada, serta implikasi praktisnya dalam dunia bisnis.

Menurut Kotler dan Keller (2020), kepuasan pelanggan dapat didefinisikan sebagai "perasaan senang atau kecewa seseorang yang dihasilkan dari perbandingan antara kinerja produk yang dipersepsikan dan harapan pelanggan (Kotler dan Keller, 2020) Kepuasan pelanggan yang tinggi dapat meningkatkan loyalitas pelanggan, yang pada gilirannya dapat berdampak positif pada kinerja bisnis. Definisi ini menunjukkan bahwa kepuasan pelanggan bersifat subjektif dan dapat bervariasi dari satu individu ke individu lainnya. Kepuasan pelanggan sangat penting karena dapat mempengaruhi loyalitas pelanggan, reputasi merek, dan pada akhirnya, profitabilitas perusahaan.

Kepuasan pelanggan dapat didefinisikan sebagai evaluasi emosional yang muncul setelah pelanggan menggunakan produk atau layanan, di mana mereka membandingkan hasil yang diperoleh dengan harapan yang dimiliki sebelumnya.

2.2. Produk Mochi

Mochi memiliki sejarah yang panjang dan kaya, yang dapat ditelusuri kembali ke zaman kuno di Jepang. Menurut catatan sejarah, mochi telah ada sejak lebih dari 1.000 tahun yang lalu. Awalnya, mochi digunakan dalam upacara keagamaan dan perayaan, serta sebagai makanan yang memberikan kekuatan dan

keberuntungan. Pada zaman Heian (794-1185), mochi mulai menjadi makanan yang lebih umum dan digunakan dalam berbagai hidangan. Selama periode Edo (1603-1868), mochi semakin populer dan mulai diproduksi secara massal. Makanan ini juga menjadi simbol perayaan Tahun Baru Jepang, di mana mochi sering disajikan dalam bentuk "ozoni", sup yang berisi mochi dan sayuran.

Menurut Kato (2021), mochi awalnya dibuat sebagai simbol keberuntungan dan kemakmuran, dan sering kali disajikan dalam bentuk kue untuk menghormati dewa-dewa (Kato, 2021). Mochi adalah makanan tradisional Jepang yang terbuat dari beras ketan yang di haluskan dan dibentuk menjadi bulatan. Dalam beberapa tahun terakhir, mochi telah menjadi populer di berbagai negara, termasuk Indonesia. Mochi dapat diisi dengan berbagai bahan, seperti pasta kacang merah, es krim, atau buah-buahan, yang membuatnya menjadi pilihan dessert yang menarik. Penjualan mochi yang terlaris dapat memberikan wawasan tentang preferensi pelanggan dan faktor-faktor yang mempengaruhi kepuasan mereka.

Mochi di Indonesia merupakan hasil dari interaksi budaya yang kaya antara Jepang dan Indonesia. Dari makanan tradisional yang dibawa oleh imigran Jepang, mochi telah berkembang menjadi camilan yang populer dan diterima oleh masyarakat Indonesia. Dengan inovasi rasa dan adaptasi terhadap selera lokal, mochi kini menjadi bagian dari keragaman kuliner Indonesia yang terus berkembang. Mochi adalah salah satu kue yang berasal dari Jepang serta terbuat dari tepung ketan dicampur dengan bahan lain, setelah itu dikukus hingga matang. Mochi yang telah matang dibentuk bulatan serta ditaburi tepung sagu ataupun tepung maizena yang telah disangrai. Kandungan gizi yang terdapat pada mochi sebanyak 75-90% karbohidrat dan kandungan proteinnya sedikit sekali. Produk mochi dalam satu porsi mengandung protein 1,3 g, fiber 1,3 g, lemak 1,3 g dan karbohidrat 16 g.

Dalam beberapa tahun terakhir, mochi telah menjadi semakin populer di berbagai belahan dunia, termasuk Indonesia. Popularitas ini tidak hanya disebabkan oleh rasa dan tekstur uniknya, tetapi juga oleh inovasi dalam variasi rasa dan penyajian yang menarik. Dengan meningkatnya permintaan akan mochi,

banyak pelaku usaha yang mulai menjual produk ini, baik dalam bentuk tradisional maupun modern. Mochi terlaris di pasaran tidak hanya bergantung pada rasa dan kualitas produk, tetapi juga pada pengalaman pelanggan saat melakukan pembelian. Oleh karena itu, penting untuk menganalisis kepuasan pelanggan terhadap produk mochi terlaris agar pemilik usaha dapat mengidentifikasi area yang perlu diperbaiki dan strategi pemasaran yang lebih efektif.

2.3. Knowledge Discovery in Database (KDD)

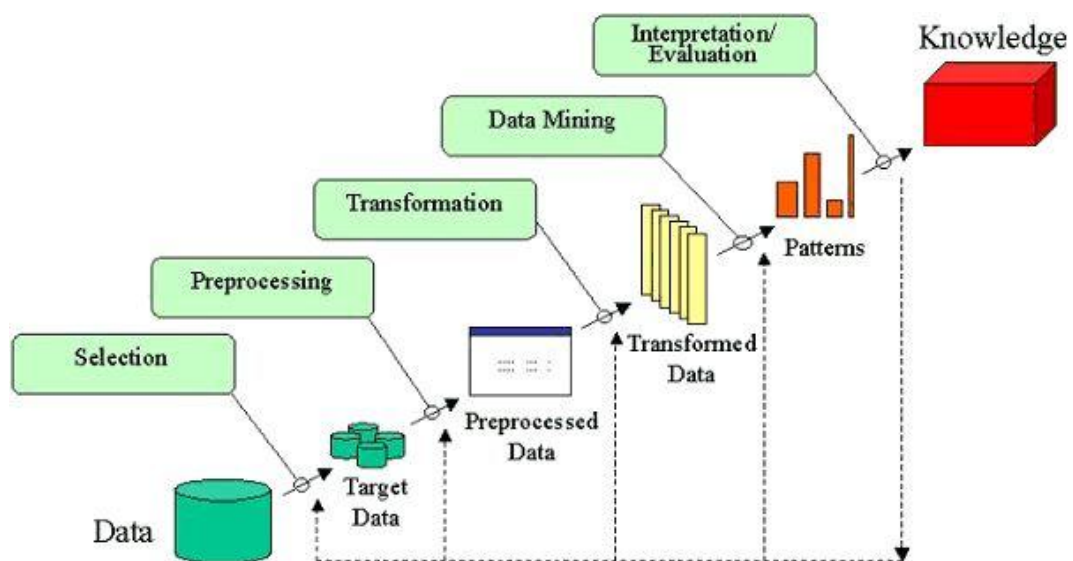
Knowledge Discovery in Database Process (KDD) adalah salah satu metode yang bisa digunakan dalam melakukan data mining. KDD sebagai proses dari menggunakan metode data mining untuk mencari informasi-informasi yang berharga, pola yang ada di dalam data, yang melibatkan algoritma untuk mengidentifikasi pola pada data. Proses KDD dari berbagai step, yaitu: seleksi data, pra-proses data, transformasi data, data mining, dan yang terakhir interpretasi dan evaluasi. Menurut Fayyad et al. (2020), KDD adalah "proses yang mencakup langkah-langkah dari pemilihan data, pra-pemrosesan, transformasi, data mining, evaluasi, hingga penyajian hasil" (Fayyad, Piatetsky-Shapiro, & Smyth, 2020).

Knowledge Discovery in Databases (KDD) adalah proses yang digunakan untuk menemukan pola, informasi, dan pengetahuan yang berguna dari kumpulan data besar. Dalam era big data saat ini, KDD menjadi semakin penting karena organisasi dan individu memiliki akses ke volume data yang sangat besar dan beragam. Proses KDD melibatkan beberapa langkah yang sistematis, mulai dari pengumpulan data hingga interpretasi hasil, dan sering digunakan dalam berbagai bidang, termasuk bisnis, kesehatan, ilmu sosial, dan banyak lagi.

Knowledge Discovery in Databases (KDD) mencakup berbagai teknik dan metode untuk mengekstrak pengetahuan dari data. Proses ini tidak hanya melibatkan teknik data mining, tetapi juga mencakup langkah-langkah sebelumnya dan sesudahnya, seperti pemilihan data, pra-pemrosesan, dan evaluasi

hasil. KDD bertujuan untuk mengubah data mentah menjadi informasi yang dapat digunakan untuk pengambilan keputusan yang lebih baik.

Berikut adalah ilustrasi serta penjelasan mengenai proses KDD secara detail:



Gambar 2.1 Knowledge Discovery in Database

1. Data Cleansing, Proses dimana data diolah lalu dipilih data yang dianggap bisa dipakai.
2. Data Integration, Proses menggabungkan data yang dianggap berulang akan digabungkan menjadi satu.
3. Selection, Proses seleksi atau pemilihan data yang dianggap relevan terhadap analisis.
4. Data Transformation, Proses transformasi data terpilih ke dalam bentuk mining procedure.
5. Data Mining, Proses dimana dilakukan beragam teknik untuk mengekstrak pola-pola potensial menghasilkan data yang berguna.
6. Pattern Evolution, Proses dimana pola-pola yang telah diidentifikasi berdasarkan measure yang diberikan.

7. Knowledge Presentation, Proses paling akhir dari proses KDD, Data-data yang sudah diproses divisualisasikan agar lebih mudah dipahami oleh pengguna dan diharapkan bisa diambil Tindakan berdasarkan analisis.

2.4. Data Mining

Data mining adalah suatu proses ekstraksi atau penggalian data yang belum diketahui sebelumnya, namun dapat dipahami dan berguna dari database yang besar serta digunakan untuk membuat suatu keputusan bisnis yang sangat penting. Data mining biasa juga disebut dengan “Data atau knowledge discovery” atau menemukan pola tersembunyi pada data. Data mining adalah proses dari menganalisa data dari prespektif yang berbeda dan menyimpulkannya ke dalam informasi yang berguna. Data mining didefinisikan sebagai proses mengekstrak atau menambang pengetahuan yang dibutuhkan dari sejumlah data besar. Menurut Han et al. (2020), data mining adalah "proses otomatisasi untuk menemukan pola dalam data besar yang dapat digunakan untuk pengambilan keputusan" (Han, Kamber, & Pei, 2020).

Pada prosesnya data mining akan mengekstrak informasi yang berharga dengan cara menganalisis adanya pola-pola ataupun hubungan keterkaitan tertentu dari data-data yang berukuran besar. Data mining berkaitan dengan bidang ilmu lain, seperti Database System, Data Warehousing, Statistic, Machine Learning, Information Retrieval, dan Komputasi Tingkat Tinggi. Selain itu data mining didukung oleh ilmu lain seperti Neural Network, Pengenalan Pola, Spatial Data Analysis, Image Database, Signal Processing.

Data mining adalah salah satu langkah dalam proses KDD secara keseluruhan. Secara umum, data mining digunakan oleh banyak peneliti sebagai sinonim dari proses KDD. Akhir-akhir ini, data mining dan knowledge discovery telah diusulkan sebagai nama yang paling memadai untuk keseluruhan proses KDD. Knowledge Discovery in Database berkaitan dengan proses penemuan pengetahuan yang diterapkan pada database. Hal ini juga didefinisikan sebagai

proses nontrivial untuk identifikasi data yang valid dan baru berpotensi bermanfaat, dan akhirnya memiliki pola yang dapat dimengerti.

Data mining memiliki aplikasi luas di berbagai sektor, termasuk bisnis, keuangan, kesehatan, ilmu pengetahuan, dan lainnya. Contoh penggunaannya meliputi deteksi kecurangan keuangan, analisis perilaku pelanggan, prediksi penjualan, pengembangan obat dalam bidang kesehatan, dan pemrosesan bahasa alami. Dengan meningkatnya volume data yang dihasilkan oleh berbagai sumber seperti sensor, media sosial, dan perangkat Internet of Things (IoT), peran data mining semakin krusial dalam mengolah informasi yang relevan dan bernilai dari tumpukan data yang besar dan kompleks.

Meskipun memberikan banyak manfaat, data mining juga melibatkan beberapa tantangan, seperti privasi data dan etika penggunaan informasi yang ditemukan. Oleh karena itu, pemahaman mendalam tentang metode, alat, dan proses yang terlibat dalam data mining sangat penting untuk memastikan bahwa analisis data dilakukan dengan benar dan menghasilkan hasil yang dapat diandalkan. Dengan terus berkembangnya teknologi dan inovasi, peran data mining diharapkan terus berkembang dan memberikan kontribusi besar pada perkembangan ilmu pengetahuan dan pengambilan keputusan.

2.5. Metode *Naïve Bayes*

Naïve Bayes adalah salah satu algoritma klasifikasi yang berbasis pada teorema Bayes dengan asumsi independensi antar fitur. Metode ini sering digunakan dalam analisis data untuk mengklasifikasikan data ke dalam kategori tertentu. Menurut Rahman et al. (2021), "*Naïve Bayes* adalah metode yang efisien dan sederhana untuk klasifikasi yang sering digunakan dalam pengenalan teks dan analisis sentimen" (Rahman, Hossain, & Rahman, 2021). Dalam konteks analisis kepuasan pelanggan, *Naïve Bayes* dapat digunakan untuk memprediksi tingkat kepuasan pelanggan berdasarkan fitur-fitur tertentu, seperti rasa, harga, dan kualitas produk. *Naïve Bayes* memiliki keunggulan dalam hal kecepatan dan efisiensi, terutama ketika digunakan pada dataset yang besar.

2.5.1. Teori *Naïve Bayes*

Teori *Naïve Bayes* beroperasi dengan menghitung probabilitas posterior dari suatu kelas berdasarkan fitur yang ada. Rumus dasar dari *Naïve Bayes* adalah:

$$[P(C|X) = \frac{P(X|C) \cdot P(C)}{P(X)}]$$

Di mana:

$P(C|X)$ adalah probabilitas kelas (C) diberikan fitur (X).

$P(X|C)$ adalah probabilitas fitur (X) diberikan kelas (C).

$P(C)$ adalah probabilitas prior dari kelas (C).

$P(X)$ adalah probabilitas fitur (X).

2.6. Metode *K-means clustering*

K-means clustering adalah metode pengelompokan yang digunakan untuk mengelompokkan data ke dalam beberapa kluster berdasarkan kesamaan fitur. Metode ini bekerja dengan cara mengidentifikasi pusat kluster dan mengelompokkan data berdasarkan jarak terdekat ke pusat tersebut. K-means sering digunakan dalam analisis segmentasi pasar dan dapat membantu dalam memahami karakteristik pelanggan yang berbeda. K-means clustering efektif dalam mengidentifikasi pola dalam data dan memberikan wawasan yang berharga untuk pengambilan keputusan. K-means adalah salah satu metode Clustering yang paling populer karena kesederhanaan dan efisiensinya. Menurut Jain (2020), "K-Means adalah salah satu algoritma clustering yang paling populer dan banyak digunakan karena kesederhanaannya dan efisiensinya dalam menangani data besar" (Jain, 2020).

K-means clustering merupakan salah satu algoritma klasterisasi yang banyak digunakan dalam analisis data dan machine learning. Algoritma ini bekerja dengan cara membagi sekumpulan data menjadi k kluster, di mana setiap kluster memiliki 18 pusat yang disebut centroid. Proses klasterisasi dilakukan dengan mengelompokkan data ke kluster terdekat berdasarkan jarak Euclidean antara data

dan centroid. Algoritma *K-means* berupaya mengoptimalkan penempatan centroid sedemikian rupa sehingga total jarak antara data dan centroid di dalam kluster menjadi minimal. Satu keunggulan utama *K-means* adalah kecepatan eksekusinya, membuatnya cocok untuk mengelompokkan data dalam skala besar.

Algoritma ini juga relatif sederhana untuk diimplementasikan dan diinterpretasikan. Meskipun demikian, *K-means* memiliki beberapa asumsi, termasuk asumsi terhadap bentuk kluster yang berbentuk bola dan ukuran kluster yang seimbang, yang dapat memengaruhi kinerjanya pada data dengan struktur yang kompleks atau tidak teratur. Untuk metode yang akan digunakan yaitu metode *K-means* dan untuk rumus yang akan penulis gunakan yaitu menggunakan rumus jarak *Euclidean Distance*. Untuk rumusnya yaitu sebagai berikut.

$$d = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2}$$

Keterangan:

“d” merupakan jarak *Euclidean Distance*

(x1 , y1) merupakan koordinat titik pertama

(x2, y2) merupakan koordinat titik kedua

2.6.1. Proses *K-means*

Keberhasilan algoritma sangat tergantung pada pemilihan nilai k (jumlah kluster) yang tepat dan inisialisasi awal centroid yang baik. *K-means* sering digunakan dalam berbagai aplikasi, termasuk segmentasi pelanggan, analisis pola geografis, dan kompresi gambar. Meskipun *K-means* memiliki kegunaan yang luas, perlu diingat bahwa hasilnya dapat bervariasi tergantung pada karakteristik data dan parameter yang dipilih. Proses *K-means* terdiri dari beberapa langkah:

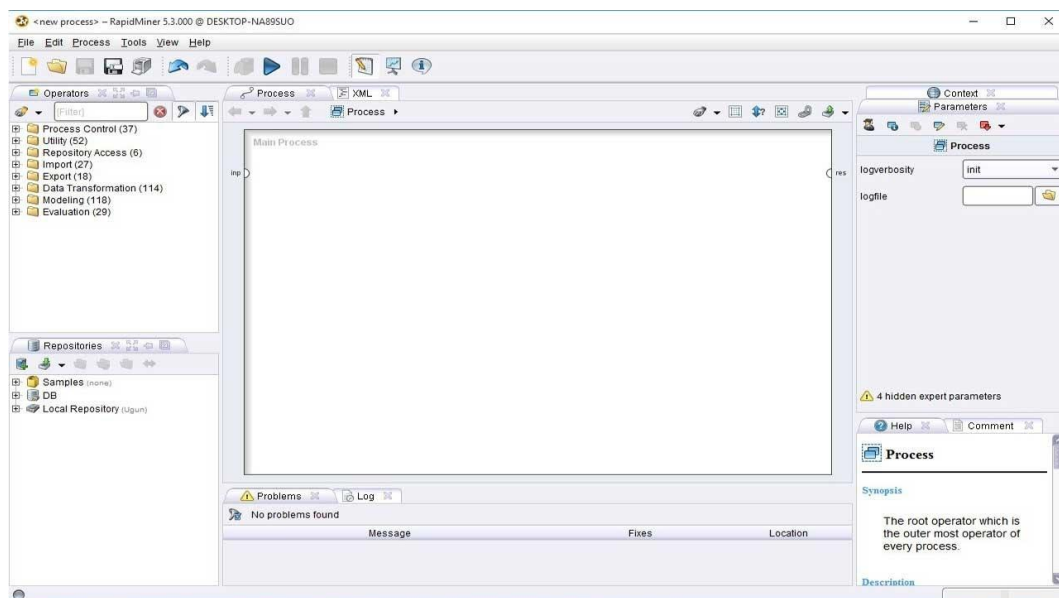
1. Inisialisasi: Memilih (k) centroid awal secara acak.

2. Penugasan: Mengelompokkan setiap data ke dalam *cluster* berdasarkan jarak terdekat ke centroid.
3. Pembaruan Centroid: Menghitung ulang posisi centroid berdasarkan rata-rata data dalam *cluster*.
4. Iterasi: Mengulangi langkah 2 dan 3 hingga tidak ada perubahan dalam penugasan *cluster*.

2.7. Penerapan *Naïve Bayes* dan *K-means* dalam Analisis Kepuasan Pelanggan

Kombinasi antara *Naïve Bayes* dan *K-means clustering* dapat memberikan analisis yang lebih komprehensif terhadap kepuasan pelanggan. *K-means* dapat digunakan untuk mengelompokkan pelanggan berdasarkan preferensi dan tingkat kepuasan mereka, sementara *Naïve Bayes* dapat digunakan untuk mengklasifikasikan tingkat kepuasan berdasarkan fitur-fitur yang relevan. Dengan pendekatan ini, pengusaha dapat mengidentifikasi segmen pelanggan yang berbeda dan memahami faktor-faktor yang mempengaruhi kepuasan mereka.

2.8. Aplikasi *RapidMiner*



Gambar 2.2 Tampilan Awal *RapidMiner*

RapidMiner adalah platform perangkat lunak sumber terbuka yang populer untuk analisis data dan pembelajaran mesin. Platform ini menyediakan berbagai alat dan fungsi yang memungkinkan pengguna untuk menjelajahi, mengolah, dan menganalisis data dengan mudah tanpa perlu pengetahuan pemrograman yang mendalam. Salah satu keunggulan *RapidMiner* adalah antarmuka pengguna grafis yang intuitif, memungkinkan pengguna dari berbagai tingkat keahlian untuk membangun dan mengevaluasi model analisis data secara visual. Menurut Kurgan dan Musilek (2020), "*RapidMiner* adalah alat yang kuat untuk analisis data yang mendukung berbagai teknik pembelajaran mesin dan visualisasi data" (Kurgan & Musilek, 2020).

Aplikasi *RapidMiner* mencakup berbagai industri dan bidang, termasuk bisnis, keuangan, kesehatan, dan ilmu pengetahuan. Dalam bisnis, *RapidMiner* digunakan untuk mendukung pengambilan keputusan, analisis pelanggan, dan prediksi tren pasar. Di sektor kesehatan, platform ini dapat membantu dalam analisis data klinis, prediksi penyakit, dan penelitian medis. Di bidang ilmu pengetahuan, *RapidMiner* digunakan untuk eksplorasi dan analisis data besar, mendukung penelitian dan pengembangan di berbagai disiplin ilmu.

RapidMiner juga mendukung integrasi dengan berbagai sumber data, termasuk database, spreadsheet, dan alat analisis data lainnya. Ini memungkinkan pengguna untuk mengakses dan menggabungkan data dari berbagai sumber untuk analisis yang lebih komprehensif. Dengan dukungan untuk berbagai teknik pembelajaran mesin, termasuk klasifikasi, regresi, klustering, dan lainnya, *RapidMiner* menjadi alat yang sangat fleksibel dan kuat untuk mengatasi tantangan analisis data dalam berbagai konteks dan industri.

2.9. Penelitian Terdahulu

Pada penelitian ini, penulis menuliskan beberapa referensi yang digunakan dalam penelitian ini sebagai pendukung dan landasan dalam pengerjaan proposal ini, yaitu sebagai berikut:

Tabel 2.1. Penelitian Terdahulu

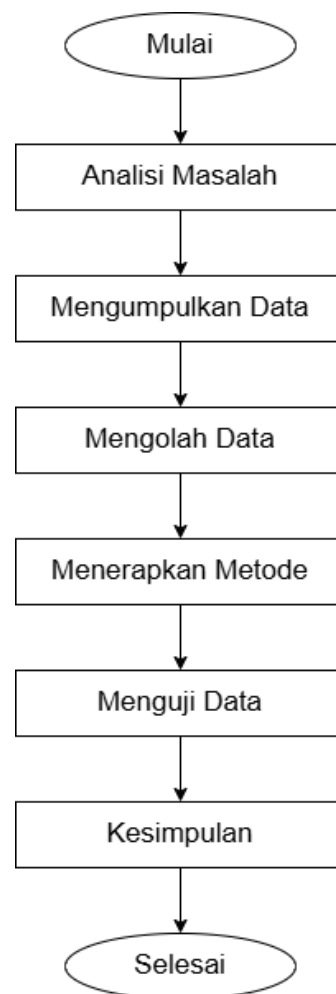
Referensi Penelitian	1
Judul	"Analisis Kepuasan Pelanggan Menggunakan Metode <i>K-means clustering</i> pada Produk Makanan"
Nama Penulis	Sari dan Prabowo
Tahun	2020
Hasil	Penelitian ini mengkaji penerapan <i>K-means clustering</i> untuk mengelompokkan pelanggan berdasarkan tingkat kepuasan mereka terhadap produk makanan. Hasil penelitian menunjukkan bahwa segmentasi pelanggan dapat membantu produsen memahami preferensi dan kebutuhan pelanggan, yang dapat meningkatkan strategi pemasaran. Meskipun tidak secara spesifik membahas mochi, temuan ini relevan untuk analisis kepuasan pelanggan dalam konteks produk makanan.
Referensi Penelitian	2
Judul	"Customer Satisfaction Analysis in Food Delivery Services Using Machine Learning Techniques"
Nama Penulis	Wang et al
Tahun	2021
Hasil	Penelitian ini mengeksplorasi penggunaan metode pembelajaran mesin, termasuk <i>Naïve Bayes</i> dan K-Means, untuk menganalisis kepuasan pelanggan dalam layanan pengantaran makanan. Penelitian ini menunjukkan bahwa kombinasi kedua metode dapat meningkatkan akurasi dalam memprediksi kepuasan pelanggan dan

	mengelompokkan pelanggan berdasarkan preferensi mereka. Temuan ini memberikan wawasan yang berharga untuk penelitian, terutama dalam konteks produk makanan seperti mochi.
Referensi Penelitian	3
Judul	"Application of <i>Naïve Bayes</i> and <i>K-means clustering</i> for Customer Satisfaction Analysis in E-commerce"
Nama Penulis	Alam dan Rahman
Tahun	2022
Hasil	Penelitian ini mengkaji penerapan metode <i>Naïve Bayes</i> dan <i>K-means</i> dalam analisis kepuasan pelanggan di platform e-commerce. Penelitian ini menemukan bahwa metode ini efektif dalam mengidentifikasi factor-faktor yang mempengaruhi kepuasan pelanggan dan dalam mengelompokkan pelanggan berdasarkan perilaku belanja mereka. Hasil penelitian ini relevan untuk analisis kepuasan pelanggan pada penjualan mochi, yang dapat dilakukan melalui platform e-commerce.
Referensi Penelitian	4
Judul	"Analyzing Customer Satisfaction in the Food Industry Using Machine Learning Approaches",
Nama Penulis	Khan et al
Tahun	2023
Hasil	Penelitian ini mengeksplorasi penggunaan berbagai teknik pembelajaran mesin, termasuk <i>Naïve Bayes</i> dan K-Means, untuk menganalisis kepuasan pelanggan di industri

	makanan. Penelitian ini menunjukkan bahwa metode ini dapat digunakan untuk mengidentifikasi preferensi pelanggan dan meningkatkan pengalaman pelanggan. Temuan ini sangat relevan untuk penelitian yang berfokus pada penjualan mochi.
--	---

Dari penelitian-penelitian terdahulu tersebut, dapat disimpulkan bahwa penggunaan metode *Naïve Bayes* dan *K-means clustering* dalam analisis kepuasan pelanggan telah terbukti efektif dalam meningkatkan akurasi prediksi dan memberikan wawasan yang lebih dalam tentang perilaku pelanggan. Penelitian ini memberikan dasar yang kuat untuk melanjutkan penelitian lebih lanjut dalam konteks penjualan mochi terlaris, dengan harapan dapat memberikan kontribusi yang signifikan terhadap pemahaman kepuasan pelanggan di industri makanan.

2.10. Langkah-Langkah Proses Penelitian (Kerangka Kerja)



Gambar 2.3. Kerangka Kerja Penelitian