



BAB II

TINJAUAN PUSTAKA

BAB II

TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Penelitian terdahulu dikutip dari beberapa referensi jurnal dan dapat dilihat pada tabel 2.I

Tabel 2.1. Penelitian Terdahulu

No	The Title Of Papers	Penulis	Publisher	Algorithm	Result
1	Klasterisasi kinerja karyawan menggunakan algoritma fuzzy cmeans	Martin dan Nataliani. 2020	AITI	Fuzzy C-Means	Dari pembahasan yang telah diberikan pada bagian sebelumnya dapat disimpulkan bahwa FCM dapat digunakan untuk mendapatkan kelompok karyawan berdasarkan kinerjanya. Kinerja tersebut diukur berdasarkan tiga kriteria yaitu presensi, kedisiplinan, dan waktu penyelesaian pekerjaan. Dengan menetapkan jumlah kelompok karyawan sebesar tiga kelompok, FCM bekerja dengan baik dalam melakukan klasterisasi atau pengelompokan kinerja karyawan. Hal ini terbukti dengan diperolehnya kelompok karyawan yang sama dengan hasil pengelompokan yang dilakukan oleh manajer. Kelompok pertama merupakan kelompok

					karyawan dengan kinerja tidak baik, kelompok kedua merupakan kelompok dengan kinerja sedang, dan kelompok ketiga merupakan kelompok dengan kinerja yang baik. Selain itu, dengan nilai eksponen fuzzy yang berbeda, dihasilkan kelompok yang berbeda juga. Nilai eksponen fuzzy 1.5 dan 2 menghasilkan kelompok karyawan yang sama dengan pengelompokan yang dilakukan oleh manajer.
2	Perbandingan Clustering Karyawan Berdasarkan Nilai Kinerja Dengan Algoritma K-Means Dan Fuzzy C-Means	Pramitasari dan Nataliani, 2021	JATISI	Fuzzy C-Means	Menurut dari hasil analisis yang sudah diperoleh dapat disimpulkan bahwa pengelompokan karyawan dari kepala bagian tidak sesuai dengan data penilaian kinerja karyawan yang berupa angka. Hasil pengelompokan dari data nilai kinerja karyawan menggunakan algoritma k-means dengan algoritma FCM berbeda. Dari hasil diagram cluster plot, pengelompokkan algoritma k-means dan algoritma FCM belum ideal. Jika dilihat dari hasil akurasi, algoritma FCM memperoleh tingkat akurasi yang lebih baik dengan hasil sebesar 76% dibandingkan dengan

					algoritma k-means dengan nilai akurasi sebesar 44%. Dari kedua algoritma tersebut yang lebih baik digunakan adalah algoritma FCM.
3	Analisis Kinerja Pegawai Menggunakan Algoritma KMeans Pada Dinas Pendidikan Dan Kebudayaan Kabupaten Bengkulu Tengah	Oktara, P., Yulianti, L., & Fredricka, J. (2021).	Jurnal Media Infotama	Fuzzy C-Means	Hasil pengelompokan terhadap data penilaian kinerja pegawai yang telah diinputkan, didapatkan bahwa jumlah pegawai yang berada pada cluster tinggi (Cluster I) 22% dan jumlah pegawai yang berada pada cluster rendah (Cluster II) 82% dari 27 Orang
4	Classification of Talent of Employee Using C4.5, KNN, SVM	Stephanie & Sarno (2019)	ICOIACT	C4.5	The conclusion of the research is that classification algorithms can help mapping talent employees. The methods that can used for classified talent are C4.5, KNN, and SVM. The highest accuracy results obtained using the SVM method, which is 94.62%. In the next study, mix algorithm can be propose to classify talent. Besides, it is necessary to check the correlation between independent variables and the dependent variable. So that variables are found that have a strong correlation with talent mapping. This can be compared with the variables that have been owned by the company and can be adjusted to the company.

					In addition, the classification that produces four guess labels can be tested using a multi-label classification.
5	A proposed Model for Predicting Employees' Performance Using Data Mining Techniques: Egyptian Case Study	Nasr, dkk, 2019	IJCSIS	C4.5	Applying the DM techniques in the different problem domains in the HRM field is considered as an important and urgent issue. Especially, at the public sector in Egypt. In addition, increasing the horizons of academic and practice research on DM in HR for reaching a government sector with a high performance.
6	Prediction of Employee Attendance Factors Using C4.5 Algorithm, Random Tree, Random Forest	Fahlapi et al (2020)	SEMESTA TEKNIKA	C4.5	The research can be concluded : from the results of the tests that have been conducted to produce, the value of accuracy for the factor of delay of workers at PT. Permata Karya Jasa uses the data mining classification algorithm, it can be proved by the results of accuracy and AUC values of each algorithm, for C.45 accuracy = 79.37% and AUC = 0.646, Random Forest accuracy = 78.58% and AUC = 0.807 while for the Random Tree algorithm accuracy = 76.26% and AUC = 0.610. As for the advice given by the author to refine the results of this study is

					the application of the algorithm C.45 and the Random Forest is expected to be able to provide solutions for workers placement PT. Permata Karya Services in order to be more disciplined in carrying out their duties as workers, especially in carrying out attendance at work in accordance with company regulations and work agreements. The following suggestion is to use another algorithms such as Particle Swarm Optimization (PSO) or another algorithm likes Genetic Algorithm (GA) as well as other algorithms to increase the level of accuracy especially in providing further analysis for more specific results in order to foster employee attendance.
7	Application of C4.5 Algorithm for Simulation of Prediction of Victory in Acceptance of Several Prospective Employees	Marpaung (2019)	Journal of Computer Networks, Architecture and High Performance Computing	C4.5	Data mining dan algoritma C4.5 sangat berguna dan memudahkan dalam pengambilan keputusan yang efektif dengan memproyeksikan data-data yang ada ke dalam bentuk phon keputusan, berdasarkan nilai entropy dan gain yang dimiliki masing-masing atribut data. Algoritma C4.5 membantu menentukan beberapa

					pegawai yang masuk dalam kriteria yang sudah di tentukan, Berdasarkan dari hasil penelitian
8	Hybrid of K-means clustering and naive Bayes classifier for predicting performance of an employee	Fadhil (2021)	Periodicals of Engineering and Natural Sciences	K-Mean Clustering	Clustering of K-means is used to assemble the information, and then the NB algorithm is used to classify it. Data was obtained from prior publications, which were employed as a test case in this study. When the dataset is put to the test, the previous label is ignored and the new data is labeled with k-means clustering. As a test tool, use to the original dataset using clustering techniques and the data should be broken down into the appropriate number of groups, such as dividing the data of 145 employees into three clusters based on the total number of labels. Next, make use of the NB algorithm for classification to anticipate the findings of employee performance data. Investigate how categorization might be improved through clustering and integration. Table 3 shows the findings of the employee performance data that was computed using the Weka. Using formulas 1, 2, and 3

					from the earlier section of this research, in terms of ACC, precision, and recall, the results are shockingly superior.
9	Analisis RapidMiner Dan Weka Dalam Memprediksi Kualitas Kinerja Karyawan Menggunakan Metode Algoritma C4.5	Hijrah (2022)	Jurnal Teknik Informatika dan Sistem Informasi	C4.5	Pengujian data kinerja karyawan pada penelitian ini dilakukan dengan metode C45, dari pengujian menggunakan aplikasi data mining Rapid Miner 7.6.0.0.1 dan Weka 3.8.1. maka diperoleh hasil aplikasi rapidminer lebih baik dari segi Accuracy sebesar 94,12 % dan Recall 95,6 %. Sedangkan aplikasi weka lebih unggul dari segi Precision 98,6 %.
10	Klasterisasi Perpustakaan Sekolah di Indonesia menggunakan C4.5 dan algoritma C-Means	Sjuchro et al (2021)	Jurnal Filosofi dan praktik perpustakaan	C4,5 dan C-Means	Pengelompokan dan pengklasifikasian algoritma memiliki akurasi 97,5% presisi 100%, dan recall 75%.

2.2. Landasan Teori

Landasan teori peneliti kutip dari beberapa teori yang berkaitan dengan penelitian ini yaitu:

2.2.1. Analisis

Analisis adalah kegiatan berpikir untuk menguraikan suatu keseluruhan menjadi komponen sehingga dapat mengenal tanda-tanda komponen, hubungannya satu sama lain dan fungsi masing-masing dalam satu keseluruhan yang terpadu. Pengertian analisis juga adalah memecahkan atau menguraikan

sesuatu unit menjadi unit terkecil. Dari pendapat diatas dapat ditarik kesimpulan bahwa analisis merupakan suatu kegiatan berfikir untuk menguraikan atau memecahkan suatu permasalahan dari unit menjadi unit terkecil. (Septiani et al., 2020).

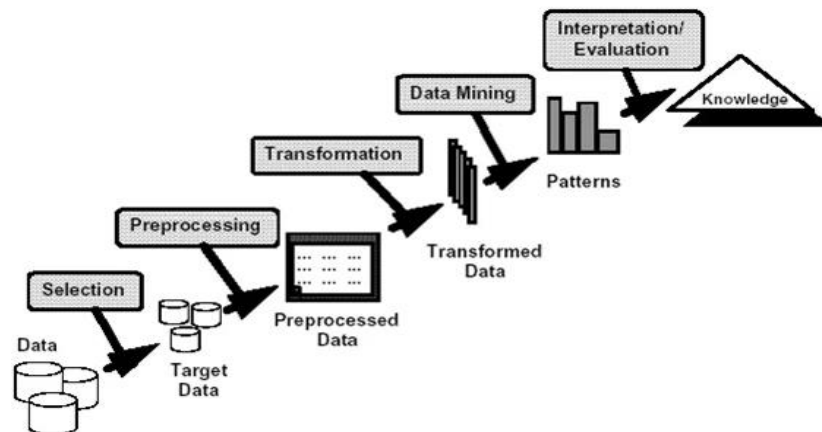
Analisis adalah sebuah kegiatan untuk mencari suatu pola selain itu analisis merupakan cara berpikir yang berkaitan dengan pengujian secara sistematis terhadap sesuatu untuk menentukan bagian, hubungan antar bagian dan hubungannya dengan keseluruhan. Analisis dilakukan terhdap data yang telah diperoleh dari hasil wawancara yang telah dilakukan. (Pasha et al., 2020).

2.3. Data Mining

Data mining dapat diartikan sebagai proses mengekstrak atau menggali knowledge yang ada pada sekumpulan data. Informasi dan knowledge yang didapat tersebut dapat digunakan pada banyak bidang, seperti manajemen bisnis, pendidikan, kesehatan dan sebagainya. (Martin & Nataliani, 2021).

2.3.1. *Knowledge Discovery In Database (KDD)*

Knowledge Discovery In Databases (KDD) adalah keseluruhan proses non-trivial untuk mencari dan mengidentifikasi pola (*pattern*) dalam data, dimana pola yang ditemukan bersifat sah, baru, dapat bermanfaat dan dapat dimengerti. KDD berhubungan dengan teknik integrasi dan penemuan ilmiah, interpretasi dan visualisasi dari pola-pola sejumlah kumpulan data. Berikut merupakan Tahapan dari KDD:



Gambar 2.1. Tahapan KDD

(Sumber: Fadhil, 2021)

Keterangan Gambar 2.1:

1. Data Cleaning

Proses pembersihan data ini merupakan proses menghilangkan noise dan data yang tidak konsisten atau data yang tidak relevan. Data-data yang harus dibersihkan seperti data yang memiliki isian yang tidak sempurna, seperti data yang hilang atau data yang tidak valid. Pembersihan data ini juga akan mempengaruhi hasil dari proses data mining yang dihasilkan karena data yang akan ditangani berkurang secara kuantitas dan kompleksitas pemrosesannya.

2. Data Integration

Proses ini merupakan proses penggabungan data ke dalam satu database dimana terkadang proses data mining yang akan dilakukan memerlukan data lebih dari satu database atau memerlukan pemrosesan menggunakan database lain. Integrasi data perlu dilakukan secara cermat karena dapat menghasilkan hasil yang menyimpang jika terjadi

kesalahan pada proses ini.

3. Data Selection

Proses penyeleksian data yang dilakukan karena tidak semua data yang ada di dalam database akan digunakan, hanya data yang sesuai untuk di analisis yang akan digunakan.

4. Preprocessing

Proses pembentukan data yang dapat digunakan.

5. Transformation

Proses penyeleksian data yang dilakukan karena tidak semua data yang ada di dalam database akan digunakan, hanya data yang sesuai untuk di analisis yang akan digunakan.

6. Data Mining

Proses utama yang dilakukan berdasarkan metode yang sudah dipilih untuk menghasilkan suatu pengetahuan.

7. Evaluation

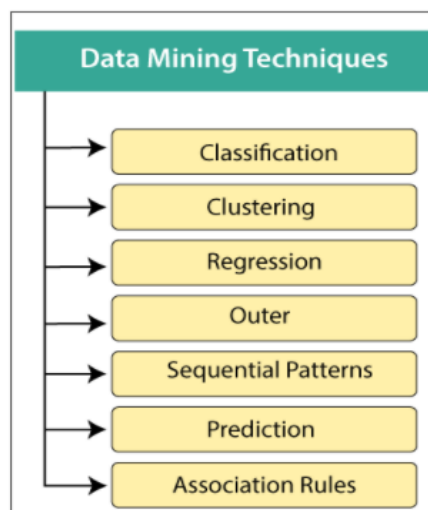
Proses ini diidentifikasikan pola-pola yang ada ke dalam knowledge based dimana hasil dari teknik data mining dievaluasi untuk dilakukan proses penilaian apakah pola hipotesa yang ada benar-benar tercapai, bila hasil pemrosesan tidak sesuai maka dapat dilakukan perbaikan pada proses data mining, mencoba metode pemrosesan yang lain atau menerima hasil ini sebagai bahan acuan untuk penelitian data mining dimasa akan datang.

8. Knowledge Persentation

Presentasi pengetahuan merupakan penyajian hasil dari pemrosesan yang ada. Pada tahap ini terfokus pada bagaimana penyajian atau visualisasi dari hasil pengetahuan yang sudah didapat sehingga dapat dipahami oleh pengguna atau yang memanfaatkan hasil dari pemrosesan tersebut. (Hijrah, 2022).

2.3.2. Teknik-Teknik Data Mining

Data mining adalah serangkaian proses untuk menggali nilai tambah dari suatu kumpulan data berupa pengetahuan yang selama ini tidak diketahui secara manual. Perlu diingat bahwa kata mining sendiri berarti usaha untuk mendapatkan sedikit data berharga dari sejumlah besar data dasar. Karena itu data mining sebenarnya memiliki akar yang panjang dari bidang ilmu seperti kecerdasan buatan (*artificial intelligent*), *machine learning*, statistik dan basis data.



Gambar 2.2. Teknik Data Mining
(Sumber: Marpaung, 2019)

2.3.3. Pengelompokan Data Mining

Data Mining dibagi menjadi beberapa kelompok berdasarkan tugas yang dapat dilakukan, berikut penjelasannya dapat dilihat dibawah ini:

1. Deskripsi

Terkadang secara sederhana ingin mencoba mencari cara untuk menggambarkan pola dan kecenderungan yang terdapat dalam data. Sebagai contoh, petugas pengumpulan suara mungkin tidak dapat menemukan keterangan atau fakta bahwa siapa yang tidak cukup profesional akan sedikit didukung dalam pemilihan presiden. Deskripsi dari pola dan kecenderungan sering memberikan kemungkinan penjelasan untuk suatu pola.

2. Estimasi

Estimasi hampir sama dengan klasifikasi, kecuali variabel target estimasi lebih ke arah numerik daripada ke arah kategori. Model dibangun menggunakan record lengkap yang menyediakan nilai dari variabel target sebagai nilai prediksi. Selanjutnya, pada penilaian berikutnya estimasi nilai dari variabel target dibuat berdasarkan nilai variabel prediksi. Sebagai contoh, akan dilakukan estimasi tekanan darah sistolik pada pasien rumah sakit berdasarkan umur pasien, jenis kelamin, indeks berat badan, dan level sodium darah. Hubungan antara tekanan darah sistolik dan nilai variabel prediksi dalam proses pembelajaran akan menghasilkan model estimasi. Model estimasi yang dihasilkan dapat digunakan untuk kasus baru lainnya.

3. Prediksi

Prediksi hampir sama dengan klasifikasi dan estimasi, kecuali dalam prediksi nilai dari hasil akan ada di masa mendatang. Contoh prediksi dalam bisnis dan penelitian adalah prediksi harga beras dalam tiga bulan mendatang dan prediksi presentasi kenaikan kecelakaan lalu lintas tahun depan jika batas bawah kecepatan dinaikkan.

4. Klasifikasi

Klasifikasi terdapat target variabel kategori. Sebagai contoh penggolongan pendapatan dapat dipisahkan dalam tiga kategori, yaitu pendapatan tinggi, pendapatan sedang, dan pendapatan rendah. Kemudian untuk menentukan pendapatan seorang pegawai, dipakai cara klasifikasi dalam data mining.

5. Pengklusteran

Pengklusteran merupakan pengelompokan record, pengamatan atau memperhatikan, dan membentuk kelas objek-objek yang mempunyai kemiripan. Kluster adalah kumpulan record yang memiliki kemiripan satu dengan yang lainnya dan memiliki ketidakmiripan dengan record-record dalam kluster lain. Contoh pengklusteran dalam bisnis dan penelitian, antara lain: (a) Mendapatkan kelompok-kelompok konsumen untuk target pemasaran dari suatu produk bagi perusahaan yang tidak memiliki dana pemasaran yang besar. (b) Untuk tujuan audit akuntansi, yaitu melakukan pemisahan terhadap perilaku nansial. (c) Melakukan pengklusteran terhadap ekspresi dari gen untuk

mendapatkan kemiripan perilaku dari gen dalam jumlah besar.

6. Asosiasi Tugas

Asosiasi dalam data mining adalah menemukan atribut yang muncul dalam satu waktu. Dalam dunia bisnis lebih umum disebut analisis keranjang belanja. Asosiasi mencari kombinasi jenis barang yang akan terjual untuk bulan depan. Contoh asosiasi dalam bisnis dan penelitian, antara lain: (a) Meneliti jumlah pelanggan dari perusahaan telekomunikasi seluler yang diharapkan untuk memberikan respon positif terhadap penawaran upgrade layanan yang diberikan. (b) Menentukan barang dalam supermarket yang dibeli secara bersamaan dan yang tidak pernah dibeli secara bersamaan. (Nasr et al., 2019).

2.4. Fuzzy C-Means

Fuzzy C-Mean (FCM) adalah suatu teknik pengelompokan data yang mana keberadaan tiap-tiap data dalam suatu cluster ditentukan oleh nilai keanggotaan. (Saputra, 2021).

Fuzzy logic sebagai cabang dari sistem kecerdasan buatan dapat diterapkan untuk melakukan klasifikasi buah kopi. Salah satu *fuzzy logic* yang dapat diterapkan sebagai kecerdasan buatan ke dalam sistem klasifikasi adalah logika *Fuzzy C-Means*. *Fuzzy C-Means* merupakan algoritma iteratif, yang menerapkan iterasi pada proses clustering data. Tujuan dari *Fuzzy C-Means* adalah untuk mendapatkan pusat *cluster* yang nantinya akan digunakan untuk mengetahui data yang masuk ke dalam sebuah *cluster*. *Fuzzy C-Means* menerapkan pengelompokan *fuzzy*, dimana setiap data dapat menjadi anggota dari beberapa

cluster dengan derajat keanggotaan yang berbeda-beda pada setiap *cluster*. (Pramitasari & Nataliani, 2021).

Fuzzy clustering adalah salah satu Teknik *clustering* dengan menggunakan konsep logika *Fuzzy*. *Fuzzy clustering* sangat berguna bagi pemodelan *Fuzzy* terutama dalam mengidentifikasi aturan-aturan *Fuzzy*. Salah satu metode dalam *Fuzzy clustering* yaitu *Fuzzy C-Means*. *Fuzzy C-Means* memiliki kelebihan, yaitu tingkat akurasi yang tinggi. Beberapa kelemahan dari *Fuzzy C-Means*, antara lain membutuhkan banyaknya kelompok dan matriks keanggotaan kelompok yang ditetapkan sebelumnya. Biasanya matriks keanggotaan kelompok awal ditetapkan secara acak yang menyebabkan metode *Fuzzy C-Means* memiliki masalah inkonsistensi. Untuk mengatasi kelemahan *Fuzzy C-Means*, Metode gabungan *Fuzzy C-Means* dengan *Subtractive Clustering* yang disebut *Subtractive Fuzzy C-Means*. *Subtractive Fuzzy C-Means* berhasil menetapkan jumlah *cluster*, matriks keanggotaan awal dan mengatasi masalah inkonsistensi dari *Fuzzy C-Means*.

Pemetaan atau *clustering* dengan menggunakan Algoritma *Fuzzy C-Means* dengan tahapan sebagai berikut:

1. Menetapkan matriks partisi awal U berupa matriks berukuran $n \times m$ (n adalah jumlah sampel data, yaitu = 45, dan m adalah parameter atau atribut setiap data, yaitu = 3). X_{ij} = data sampel ke- i ($i = 1, 2, \dots, n$), atribut ke- j ($j = 1, 2, \dots, m$).
2. Menentukan nilai parameter awal
 - a. Jumlah cluster (c)
 - b. Pangkat (w)

- c. Maksimum iterasi (MaxIter)
 - d. Error terkecil yang diharapkan (ϵ)
 - e. Fungsi objektif awal (P0)
 - f. Iterasi awal (t)
3. Membangkitkan bilangan random μ_{ik} , $i = 1, 2, \dots, n$; $k = 1, 2, \dots, c$ sebagai elemen – elemen matriks partisi awal (U).
 4. Menentukan pusat cluster (V) Pada iterasi pertama dengan menggunakan persamaan:

$$V_{kj} = \frac{\sum_{i=1}^n ((\mu_{ik})^2 * X_{ij})}{\sum_{i=1}^n (\mu_{ik})^2}$$

5. Menghitung nilai Fungsi Objektif (Pt) Fungsi objektif pada iterasi pertama
Pt dihitung dengan menggunakan persamaan:

$$P1 = \sum_{i=1}^n \sum_{j=1}^n \left[\sum_{k=1}^n (X_{ij} - V_{kj})^2 (\mu_{ik})^2 \right]$$

6. Mengecek Kondisi Berhenti. (Oktara et al., 2021).

2.5. Algoritma C4.5

Algoritma C4.5 merupakan salah satu teknik klasifikasi pada machine learning yang digunakan pada proses data mining dengan membentuk sebuah pohon keputusan (decision tree) yang direpresentasikan dalam bentuk aturan. Berikut adalah langkah-langkah perhitungan:

1. Data dikelompokkan berdasarkan atribut beserta nilai di dalamnya;
2. Menghitung jumlah data pada setiap nilai atribut yang ada;

3. Mengklasifikasi data yang sudah dihitung menjadi dua kelompok berdasarkan terget tujuan, yaitu laris (L) dan tidak laris (TL);
4. Menghitung entropy total dari 20 data penjualan barang;
5. Menghitung entropy dari masing-masing nilai atribut;
6. Menghitung gain dari tiap atribut;
7. Menghitung gain ratio dari tiap atribut;
8. Mencari atribut dengan gain ratio tertinggi untuk dijadikan root;
9. Menentukan nilai atribut yang akan dijadikan cabang;
10. Menentukan node selanjutnya dari atribut yang terpilih berdasarkan nilai gain ratio tertinggi. (Stephanie dan Samo, 2019).

Algoritma C4.5 adalah algoritma klasifikasi data dengan teknik pohon keputusan yang memiliki kelebihan-kelebihan. Kelebihan ini misalnya dapat mengolah data numerik (kontinyu) dan diskret, dapat menangani nilai atribut yang hilang, menghasilkan aturan-aturan yang mudah diinterpretasikan dan tercepat diantara algoritma-algoritma yang lain. Ada beberapa tahap untuk membuat pohon keputusan dengan algoritma C.45 yaitu:

1. Menyiapkan data training. Data training di ambil dari data atau dokument yang tersimpan dan tersusun dalam kelas-kelas tertentu.
2. Menentukan akar pohon. Akar dipilih dari atribut yang dipilih, dengan cara menghitung nilai gain dari masing-masing atribut, nilai tertinggi akan menjadi akar. Sebelum menghitung gain dari atribut, hitung dahulu nilai entropy. Untuk menghitung nilai entropy dengan rumus:

$$Entropy(S) = \sum_{i=1}^n - P_i * \log_2 P_i$$

Keterangan:

- S = himpunan kasus
- N = Jumlah partisi
- P_i = Proporsi S_i terhadap S

3. Menghitung nilai gain menggunakan rumus.

$$Gain(S, A) = Entropy(S) - \sum_{i=1}^n - \frac{|S_i|}{|S|} * Entropy(S_i)$$

Keterangan:

S = Himpunan kasus

A = Fitur

N = Jumlah partisi atribut A

$|S_i|$ = proporsi S_i terhadap S

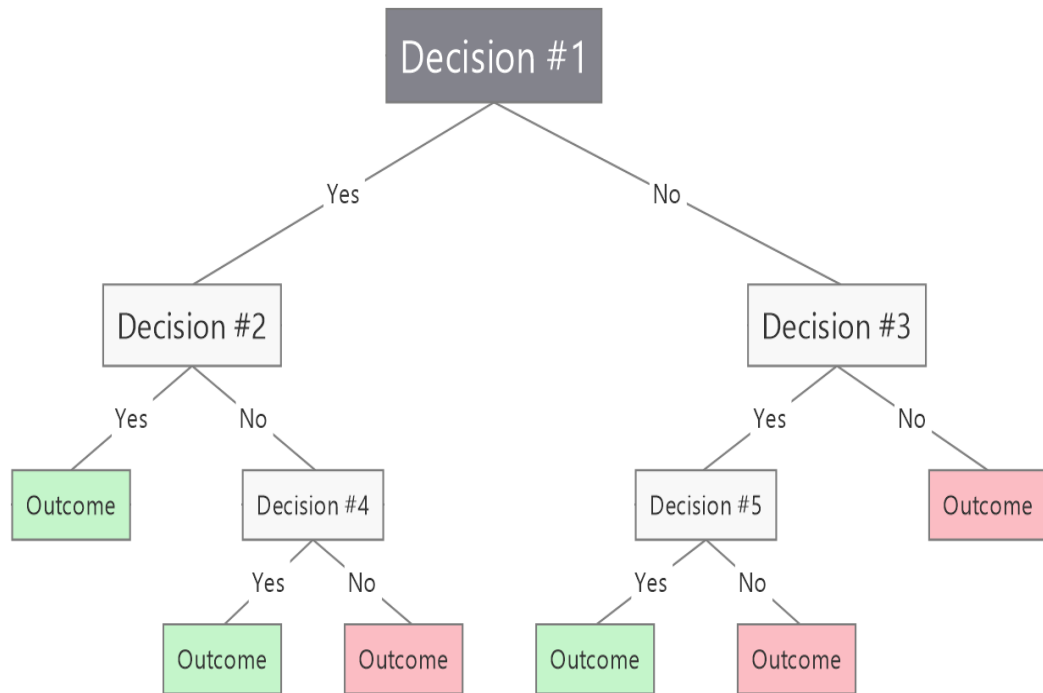
$|S|$ = Jumlah kasus dalam S

4. Ulangi langkah ke dua hingga semua record terpartisi

5. Proses partisi berhenti jika:

- a. Semua record dalam simpul N mendapat nilai kelas yang sama.
- b. Tidak ada atribut dalam record yang dipartisi lagi.
- c. Tidak ada record di dalam cabang yang kosong. (Nasr et al., 2019).

Decision Tree atau pohon keputusan adalah model prediksi menggunakan struktur pohon atau hirarki dengan mengubah data menjadi pohon keputusan dan aturan-aturan keputusan.



Gambar 2.3. *Decision Tree*
(Sumber: Nasr et al., 2019)

2.6. Klasifikasi

Klasifikasi merupakan bagian dari prediksi, dimana nilai yang diprediksi berupa label. Klasifikasi menentukan *class* atau grup untuk tiap contoh data, *input* dari model klasifikasi adalah atribut dari contoh data (*data sample*) dan outputnya adalah *class* dari data *samples* itu sendiri, dalam *machine learning* untuk membangun model klasifikasi digunakan metode *supervised learning*. Klasifikasi dapat diartikan proses menemukan kumpulan pola atau fungsi-fungsi yang mendeskripsikan dan memisahkan kelas data satu dengan lainnya, untuk dapat digunakan untuk memprediksi data yang belum memiliki kelas data tertentu. Jadi secara singkat, klasifikasi adalah proses untuk membedakan atau memisahkan kelas. (Fadhil, 2021).

2.7. Kinerja Pegawai

Kinerja pada umumnya dapat diartikan sebagai kesuksesan seseorang dalam melaksanakan suatu pekerjaan, kinerja yang baik adalah kinerja yang mengikuti tata cara atau prosedur sesuai standar yang telah ditetapkan. Kinerja merupakan hasil kerja seseorang, dimana keseluruhan hasil tersebut dapat dibuktikan secara konkrit dan dapat diukur. Kinerja yang dapat dinilai dan diukur secara objektif akan meningkatkan motivasi karyawan untuk dapat bekerja lebih baik, tetapi apabila kinerja dinilai secara subjektif dan tidak ada pengukuran yang jelas akan menyebabkan karyawan terdemotivasi dan membuat ketidakpuasan dalam bekerja. (Tyas & Sunuharyo, 2018).

2.7.1. Indikator Kinerja

Terdapat beberapa indikator kinerja yaitu:

1. Kualitas Kerja.
2. Kuantitas Kerja.
3. Tanggung Jawab.
4. Kerjasama.
5. Inisiatif. (Tyas & Sunuharyo, 2018).

2.7.2. Faktor-Faktor Yang Mempengaruhi Kinerja

Faktor-faktor yang mempengaruhi kinerja karyawan dalam melaksanakan tugasnya dapat bersumber dari faktor individu dan faktor lingkungan organisasi. Perusahaan dapat meningkatkan kinerja karyawannya apabila di dalam diri karyawannya memiliki kesadaran untuk dapat patuh dan mentaati aturan di perusahaan. Faktor lingkungan organisasi juga dapat mempengaruhi kinerja karyawan, lingkungan kerja yang baik akan mempermudah karyawan dalam menjalankan pekerjaannya karena karyawan akan merasa terbantu dengan lingkungan kerja fisiknya dan akan merasa nyaman bekerja dengan lingkungan non fisiknya, dengan memperhatikan disiplin kerja dan lingkungan kerja karyawan akan memiliki kinerja yang baik sehingga tujuan perusahaan dapat tercapai. (Tyas & Sunuharyo, 2018).

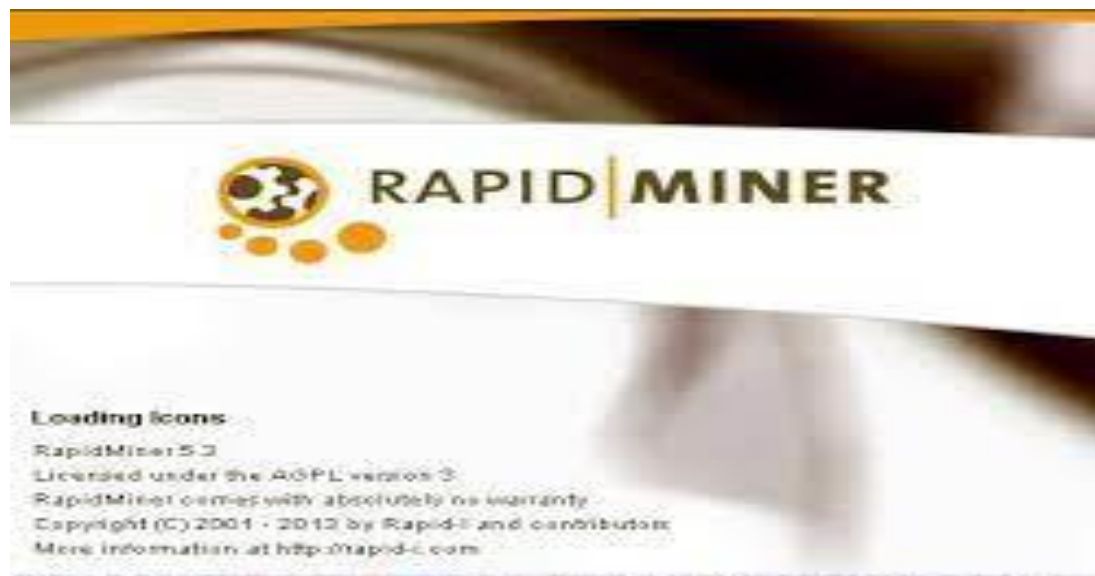
2.7.3. Kriteria Utama Dalam Pengukuran Kinerja

Kriteria utama dalam pengukuran kinerja dapat dilakukan dengan cara sebagai berikut:

1. Kualitas, yaitu berkaitan dengan baik tidaknya mutu yang dihasilkan. Pengukuran kualitatif keluaran mencerminkan pengukuran “tingkat kepuasan” yaitu seberapa baik penyelesaiannya.
2. Kuantitas, yaitu jumlah yang harus diselesaikan atau dicapai. Pengukuran kuantitatif melibatkan perhitungan keluaran dari proses atau pelaksanaan kegiatan. Ini berkaitan dengan jumlah keluaran yang dihasilkan.
3. Ketepatan waktu, yaitu sesuai tidaknya dengan waktu yang direncanakan. Ketepatan waktu merupakan jenis khusus dari pengukuran kuantitatif yang menentukan ketepatan waktu penyelesaian suatu kegiatan. (Tyas & Sunuharyo, 2018).

2.8. *Rapid Miner*

Rapid Miner merupakan perangkat lunak yang dapat diakses oleh siapa saja dan bersifat terbuka (*open source*). *Rapid Miner* ini dijadikan sebuah solusi untuk menganalisa terhadap data *processing*. Pada *Rapid Miner* ini digunakan berbagai teknik seperti teknik deskriptif dan prediksi. *Rapid Miner* ini menggunakan bahasa *Java* untuk pengoperasiannya. Berkat adanya kecanggihan teknologi algoritma komputasi dan analisis data berbasis komputer, datamining dapat diolah menggunakan *software Rapid Miner*. (Sari et al., 2020).



Gambar 2.4. *Rapid Miner*
(Sumber: Sari et al., 2020)