

BAB II

TINJAUAN PUSTAKA

2.1 Penelitian Terdahulu

Penelitian yang sudah dilakukan sebagai solusi untuk menyelesaikan masalah data *imbalanced* dengan menggunakan beberapa metode serta teknik optimasi sehingga memperoleh performa *accuracy* yang tepat dan akurat. Salah satu metode yang mampu menyelesaikan masalah data *imbalanced* yaitu *cost sensitive* yang di implementasikan pada metode *decision tree*. Metode tersebut bekerja dengan meminimalkan kesalahan *decision tree*. Penelitian yang telah dilakukan dengan menggunakan metode tersebut antara lain:

Tabel 2.1 Penelitian Neelesh Haldankar Akash

<i>Author, Year</i>	Neelesh Haldankar Akash, 2016
<i>Problem</i>	Memprediksi resiko pengajuan nasabah baru untuk menghindari pinjaman macet pada bank
<i>Scope of problem</i>	Data penelitian diperoleh di <i>UCI Machine Learning Repository</i> dengan menggunakan 1000 data, 20 atribut dan 2 class <i>imbalanced</i> .
<i>Method</i>	• Fitur seleksi <i>atribut</i> menggunakan metode <i>boruta</i> • Algoritma <i>decision tree</i> menggunakan <i>random forest</i> • <i>Cost sensitive</i> menerapkan metode <i>thresholding</i> untuk menghitung minimum <i>cost</i>.
<i>Result</i>	Accuracy yang diperoleh sebesar 76%

Tabel 2.2 Penelitian Atefeh Daraei

<i>Author, Year</i>	Atefeh Daraei, 2016
<i>Problem</i>	Memprediksi pasien yang terkena penyakit <i>Myocardial infarction</i> (MI) dengan jumlah data <i>imbalanced</i> .
<i>Scope of problem</i>	Data penelitian diambil dari rumah sakit di Shahid Madani Specialized, Iran. Data yang diperoleh sebanyak 750 data dengan 295 pasien didiagnosa mengidap penyakit MI dan 455 di diagnosa sehat. Terdiri dari 92 atribut dan 2 class.
<i>Method</i>	• Fitur seleksi <i>atribut</i> menggunakan metode <i>genetic algoritma</i> • Algoritma <i>decision tree</i> menggunakan J4.5 • <i>Cost sensitive</i> menggunakan algoritma <i>metacost</i> untuk menghitung minimum <i>cost</i>.
<i>Result</i>	Hasil accuracy terbaik diperoleh dengan menggunakan fitur seleksi GA sebesar 82.67%.

Tabel 2.3 Penelitian Neelesh Haldankar Akash

Author, Year	Hong Zhao, et al, 2015																																																						
Problem	Melakukan pengujian klasifikasi <i>dataset</i> , tanpa adanya <i>preprocessing</i> / seleksi fitur terlebih dahulu untuk menghasilkan data baru yang akan diproses.																																																						
Scope of problem	Dataset yang digunakan diperoleh dari <i>UCI Machine Learning Repository</i> sebanyak 8 dataset yaitu: <table><tr><th>No.</th><th>Name</th><th>Domain</th><th>U</th><th>C</th><th>$\bar{d}$</th></tr><tr><td>1</td><td><i>Biodeg</i></td><td>Chemical</td><td>1055</td><td>41</td><td>Type</td></tr><tr><td>2</td><td><i>Breast</i></td><td>Clinic</td><td>699</td><td>9</td><td>Class</td></tr><tr><td>3</td><td><i>Clean</i></td><td>Physics</td><td>476</td><td>166</td><td>Class</td></tr><tr><td>4</td><td><i>Credit</i></td><td>Commerce</td><td>1000</td><td>20</td><td>Class</td></tr><tr><td>5</td><td><i>Diabetes</i></td><td>Clinic</td><td>768</td><td>8</td><td>Class</td></tr><tr><td>6</td><td><i>Eeg</i></td><td>Life</td><td>1923</td><td>14</td><td>Class</td></tr><tr><td>7</td><td><i>Wdbc</i></td><td>Clinic</td><td>569</td><td>30</td><td>Diagnosis</td></tr><tr><td>8</td><td><i>Wilt</i></td><td>Clinic</td><td>4839</td><td>5</td><td>Class</td></tr></table>	No.	Name	Domain	$ U $	$ C $	$\bar{d}$	1	<i>Biodeg</i>	Chemical	1055	41	Type	2	<i>Breast</i>	Clinic	699	9	Class	3	<i>Clean</i>	Physics	476	166	Class	4	<i>Credit</i>	Commerce	1000	20	Class	5	<i>Diabetes</i>	Clinic	768	8	Class	6	<i>Eeg</i>	Life	1923	14	Class	7	<i>Wdbc</i>	Clinic	569	30	Diagnosis	8	<i>Wilt</i>	Clinic	4839	5	Class
No.	Name	Domain	$ U $	$ C $	$\bar{d}$																																																		
1	<i>Biodeg</i>	Chemical	1055	41	Type																																																		
2	<i>Breast</i>	Clinic	699	9	Class																																																		
3	<i>Clean</i>	Physics	476	166	Class																																																		
4	<i>Credit</i>	Commerce	1000	20	Class																																																		
5	<i>Diabetes</i>	Clinic	768	8	Class																																																		
6	<i>Eeg</i>	Life	1923	14	Class																																																		
7	<i>Wdbc</i>	Clinic	569	30	Diagnosis																																																		
8	<i>Wilt</i>	Clinic	4839	5	Class																																																		
Method	• Algoritma <i>decision tree</i> menggunakan C4.5 • <i>Cost sensitive</i> menggunakan teknik <i>probabilistic pruning mechanism</i> untuk menghitung minimum <i>cost</i>.																																																						
Result	• Metode <i>Cost sensitive decision tree</i> berhasil menurunkan <i>cost</i> dibawah 10%.																																																						

Tabel 2.4 Penelitian Shichao Zhang

Author, Year	Shichao Zhang, 2005																								
Problem	Dataset yang terdapat <i>missing value</i> dapat menjadi permasalahan karena jika dihapus dapat mengurangi kandungan informasi dari data.																								
Scope of problem	Dataset yang digunakan diperoleh dari <i>UCI Machine Learning Repository</i> sebanyak 5 dataset imbalanced yaitu: <table><tr><td></td><td>No. of Attributes</td><td>No. of Examples</td><td>Class distribution (N/P)</td></tr><tr><td>Ecoli</td><td>6</td><td>332</td><td>230/102</td></tr><tr><td>Breast</td><td>9</td><td>683</td><td>444/239</td></tr><tr><td>Heart</td><td>8</td><td>161</td><td>98/163</td></tr><tr><td>Thyroid</td><td>24</td><td>2000</td><td>1762/238</td></tr><tr><td>Australia</td><td>15</td><td>653</td><td>296/357</td></tr></table>		No. of Attributes	No. of Examples	Class distribution (N/P)	Ecoli	6	332	230/102	Breast	9	683	444/239	Heart	8	161	98/163	Thyroid	24	2000	1762/238	Australia	15	653	296/357
	No. of Attributes	No. of Examples	Class distribution (N/P)																						
Ecoli	6	332	230/102																						
Breast	9	683	444/239																						
Heart	8	161	98/163																						
Thyroid	24	2000	1762/238																						
Australia	15	653	296/357																						
Method	Menggunakan metode <i>cost sensitive</i> dengan menggunakan 4 <i>strategy misclassified cost</i> yaitu: <i>The Null Strategy, The Internal Node Strategy, The C4.5 Strategy</i>																								
Result	Hasil dari 4 <i>strategy</i> yang diujikan <i>internal node strategy</i> menjadi metode terbaik dalam menyelesaikan permasalahan data <i>missing value</i> .																								

Tabel 2.5 Penelitian Polat

<i>Author, Year</i>	Polat dan Gunes, 2009
<i>Problem</i>	Menyelesaikan masalah klasifikasi <i>imbalanced multiclass</i>
<i>Scope of problem</i>	<i>dataset</i> dari UCI <i>machine learning repository</i> , diantaranya <i>dataset Dermatology, Image Segmentation, dan Lymphography</i> .
<i>Method</i>	- Allgoritma pengklasifikasi C4.5 dengan pendekatan <i>One-Against-All</i> - Pengujian menggunakan teknik <i>10-fold cross validation</i>
<i>Result</i>	Hasil accuracy yang lebih tinggi yaitu 96,71%, 95,18%, dan 87,95% pada masing-masing <i>dataset</i> .

Tabel 2.6 Penelitian Aitkenhead

<i>Author, Year</i>	Aitkenhead, 2008
<i>Problem</i>	Mengatasi masalah data <i>noise</i> dan masalah kombinasi data kuantitatif dan kualitatif pada <i>dataset</i> yang dapat menyulitkan proses kategorisasi kelas
<i>Scope of problem</i>	<i>Dataset</i> yang digunakan adalah <i>Glass Chemistry</i> dan <i>Car Costing</i>
<i>Method</i>	mengembangkan pendekatan <i>evolusioner</i> pada algoritma <i>Decision Tree C4.5</i>
<i>Result</i>	Hasil accuracy yang lebih tinggi yaitu 82,40% dan 89,20% pada masing-masing <i>dataset</i> .

Tabel 2.7 Penelitian Balamurugan

<i>Author, Year</i>	Balamurugan dan Rajaram, 2009
<i>Problem</i>	Mengatasi masalah kegagalan dalam pemilihan <i>atribut</i> yang akan di- <i>split</i> yang dapat menyebabkan label kelas dipilih secara acak.
<i>Scope of problem</i>	<i>Dataset</i> yang digunakan dalam penelitian tersebut diantaranya <i>Blood Transfusion, Teaching Assistant Evaluation, SPECT Heart, Haberman's Survival, Contraceptive Method Choice, Hayes Roth, Concrete, Forest-fires, Solarflare 1, dan Solarflare 2</i> .
<i>Method</i>	Algoritma <i>Decision Tree C4.5</i> dengan menambahkan prosedur penghitungan <i>Maximum Influence Factor (MIF)</i>
<i>Result</i>	Hasil penelitian tersebut didapatkan nilai accuracy lebih tinggi yaitu 85,16 %, 77,78 %, 71,70 %, 78,79 %, 77,50 %, 76,74 %, 76,74 %, 75,68 %, 77,09% pada masing-masing <i>dataset</i> .

Tabel 2.8 Penelitian Xiaoyong Chai

Author, Year	Xiaoyong Chai, Lin Deng and Qiang Yang, 2004																							
Problem	Data <i>imbalanced</i> terdapat <i>Missing value</i> pada <i>atribut</i> dapat menurunkan performa <i>classifier</i> .																							
Scope of problem	<i>Dataset</i> yang digunakan diperoleh dari <i>UCI Machine Learning Repository</i> sebanyak 8 <i>dataset imbalanced</i> yaitu: <table><tr><th>Name of datasets</th><th>No. of attributes</th><th>Name of datasets</th><th>No. of attributes</th></tr><tr><td>Ecoli</td><td>6</td><td>Breast</td><td>9</td></tr><tr><td>Heart</td><td>8</td><td>Thyroid</td><td>24</td></tr><tr><td>Australia</td><td>15</td><td>Cars</td><td>6</td></tr><tr><td>Voting</td><td>16</td><td>Mushroom</td><td>22</td></tr></table>				Name of datasets	No. of attributes	Name of datasets	No. of attributes	Ecoli	6	Breast	9	Heart	8	Thyroid	24	Australia	15	Cars	6	Voting	16	Mushroom	22
Name of datasets	No. of attributes	Name of datasets	No. of attributes																					
Ecoli	6	Breast	9																					
Heart	8	Thyroid	24																					
Australia	15	Cars	6																					
Voting	16	Mushroom	22																					
Method	<i>Cost sensitive naïve bayes</i> . Menggunakan metode klasifikasi <i>naïve bayes</i> dan <i>cost sensitive</i> digunakan untuk mencari jumlah minimal <i>cost</i> dari jumlah <i>cost</i> rata-rata <i>misclassification cost</i> dan <i>test cost</i> .																							
Result	<i>Cost sensitive naïve bayes</i> memiliki nilai total <i>cost</i> rata-rata <i>misclassification cost</i> dan <i>test cost</i> lebih baik dibandingkan dengan <i>cost sensitive decision tree</i> .																							

Tabel 2.9 Penelitian Khor Kok-Chin

<i>Author, Year</i>	Khor Kok-Chin dan Ng Keng-Hoong, 2016			
<i>Problem</i>	Data <i>imbalanced dua class</i> untuk menyeleksi customer bank yang ingin menggunakan produk bank			
<i>Scope of problem</i>	<i>Dataset</i> yang digunakan diperoleh dari <i>UCI Machine Learning Repository</i> . Data dikumpulkan 45.211 dari 14 organisasi marketing oleh bank Portugal. Data memiliki atribut sebanyak 16 dan 2 class dan memiliki sifat data numeric dan nominal			
<i>Method</i>	Penelitian ini menggunakan 3 metode pengujian yaitu <i>cost sensitive naïve bayes</i> , <i>cost sensitive decision tree C4.5</i> dan <i>cost sensitive naïve bayes tree</i> .			
<i>Result</i>	<i>Cost sensitive naïve bayes</i> memiliki nilai nilai <i>accuracy</i> terbaik sebesar 96.30%, <i>Cost sensitive decision tree C4.5</i> memiliki nilai nilai <i>accuracy</i> terbaik sebesar 96.50% dan <i>Cost sensitive naïve bayes tree</i> memiliki nilai nilai <i>accuracy</i> terbaik sebesar 96.70%,			

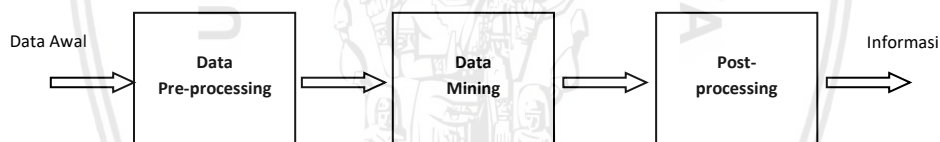
Pada tabel 2.1, 2.2 dan 2.3 merupakan penelitian yang menggunakan seleksi atribut dan metode *cost sensitive decision tree C4.5*. Hasil dari metode tersebut mampu meningkatkan performa *accuracy* pada data *imbalanced*. Pada tabel 2.4, 2.5, 2.6 dan 2.7 penelitian tersebut

menggabungkan teknik optimasi dengan metode *cost sensitive decision tree* C4.5. Teknik tersebut mampu meningkatkan performa *accuracy classifier*. Tabel 2.8 menggunakan metode *cost sensitive* pada *naïve bayes* dan *decision tree C4.5* hasilnya *cost sensitive naïve bayes* memiliki performa yang lebih baik dari pada menggunakan *decision tree C4.5*.

2.2 Data Mining

Data mining merupakan bidang keahlian yang menggabungkan teknik dari *machine learning* (mesin pembelajaran), *pattem recognition* (pengenalan pola), analisis, statistik dan visualisasi sebagai solusi permasalahan untuk mendapatkan informasi dari *database*. Pola-pola informasi tersebut dilakukan analisa untuk dipelajari sebagai perilaku yang terstruktur guna memberikan wawasan dan pengetahuan untuk mengambil keputusan (Larose, 2005).

Data mining merupakan bagian dari *knowledge discovery in database* (KDD). Sebelum mengambil keputusan, langkah awal dilakukan beberapa tahapan proses antara lain pengumpulan data sebagai masukan, pemrosesan data, hingga keluaran berupa informasi. Secara umum tahapan pengolahan data hingga memperoleh informasi digambarkan seperti gambar 2.1.



Gambar 2.1. Skema Proses KDD

Data awal menjadi masukan dari proses KDD yang didalamnya mengandung informasi. Proses pengolahan data awal (*preprocessing*) diperlukan sebelum masuk pada proses *data mining*. Pada proses ini, data diproses menjadi keluaran yang diinginkan dan dibutuhkan dalam proses *data mining*. Proses ini juga menggabungkan data dari berbagai sumber menjadi satu, menghilangkan *noise*, *menghapus* data duplikat, dan pemilihan *atribut* yang relevan.

Pre-processing merupakan proses yang memerlukan banyak sumber daya, karena proses ini merupakan tahapan penting dari keseluruhan proses KDD. Hal ini disebabkan karena kualitas data masukan memberikan pengaruh terhadap hasil keluaran proses. Tahap *data mining* dilakukan setelah tahap *pre-processing* data. Informasi diperoleh dengan melakukan

observasi pola dan hubungan dalam data menggunakan algoritma *machine learning*. Tahap terakhir adalah *post-processing*, dimana dalam tahap ini hasil penggalian data akan diintegrasikan dengan sistem pengambilan keputusan atau divisualisasikan untuk memudahkan pengguna dalam menganalisa informasi. *Data mining* secara umum terbagi menjadi dua jenis yaitu sebagai berikut:

1. **Prediktif.** Teknik ini bertujuan untuk memprediksi target *class* berdasarkan nilai dari karakteristik *atribut* lainnya. Teknik jenis prediktif antara lain Klasifikasi dan *regresi*.
2. **Deskriptif.** Teknik ini bertujuan untuk mendapatkan pola yang menghasilkan simpulan hubungan dari data. Metode ini digunakan untuk menjelaskan pola dan hubungan data. Teknik jenis deskriptif antara lain *Clustering* dan *asosiasi*.

2.3 Pre-processing

preprocessing adalah suatu proses yang dilakukan untuk membuat data awal menjadi data yang berkualitas (*input* yang baik untuk *data mining*). Teknik menghasilkan data yang berkualitas melalui beberapa proses tahapan antara lain pembersihan data, pengurangan data dan *transformasi* data. Proses pembersihan data seperti pengisian nilai kosong (*missing value*), menghapus nilai duplikasi, *noise* dan *outlier*. Proses *transformasi* data dengan menggunakan metode *normalisasi* atau *agregasi*.

2.3.1 Normalisasi Min-Max

Normalisasi merupakan suatu proses *transformasi* data. Salah satu teknik yang digunakan dalam proses *normalisasi* yaitu *min-max*. Teknik *min-max* merupakan metode *normalisasi* dengan cara pengskalaan nilai data antara 0 sampai 1. Tujuan teknik tersebut menyeimbangkan data dalam menentukan nilai informasi yang terkandung dalam sebuah data. Persamaan *min-max* ditunjukkan dibawah ini. (Nasution, 2019)

$$X_n = \frac{X_0 - X_{\min}}{X_{\max} - X_{\min}} \dots\dots\dots (2.1)$$

Keterangan:

X_n = Nilai data *normalisasi* (0 - 1)

X_0 = Nilai data *aktual*

$X_{\min}$ = Nilai *minimum* data aktual keseluruhan

$X_{\max}$ = Nilai *maksimum* data aktual keseluruhan

2.3.2 Seleksi Atribut PSO

Pemilihan *atribut* menjadi faktor penting untuk menghasilkan data yang berkualitas dalam proses *data mining*. Oleh sebab itu, perlu adanya seleksi atribut untuk memilih atribut relevan dan menghapus atribut yang tidak relevan sehingga menghasilkan pengetahuan yang *informatif* dan akurat. Salah satu metode yang baik untuk menyeleksi *atribut* adalah *particle swarm optimization* (PSO). PSO merupakan sebuah metode optimasi yang dapat digunakan untuk menyeleksi *atribut* berdasarkan nilai bobot informasi (w).

Metode PSO terinspirasi dari perilaku sekawanan hewan seperti ikan atau burung dalam mencari makan. PSO mengansumsikan atribut sebagai sekumpulan partikel dimana sekumpulan partikel mencari posisi dan kecepatan optimal dalam ruang fitur. Kelebihan metode optimasi *Particle Swarm Optimization* memiliki konsep yang sederhana, mudah diimplementasikan, dan efisien dalam perhitungan jika dibandingkan dengan algoritma matematika dan teknik *optimisasi heuristik* lainnya. Kemudian proses seleksinya sebagai berikut (Shih Wei, KuoChing, Shih-Chieh, & Zne-Jung, 2008).

1. Mengansumsikan jumlah kawanan partikel adalah N . Kecepatan dan posisi awal pada setiap partikel dalam N dimensi ditentukan secara *random* (acak).
2. Menghitung kecepatan semua partikel. Semua partikel bergerak menuju titik optimal yaitu 1. kecepatan awal dari semua partikel diasumsikan sama dengan nol.
3. Membandingkan hasil dari nilai *fitness* setiap partikel dengan nilai sebelumnya. Jika nilai *fitness* setiap partikel pada posisi saat ini lebih baik dari Pbest, maka Pbest diatur untuk posisi saat ini.
4. Membandingkan nilai *fitness* partikel dengan Gbest. Jika Gbest yang terbaik maka Gbest yang di *update*.
5. Persamaan (2.1) dan (2.2) ditunjukkan di bawah ini untuk memperbaharui (*update*) kecepatan (*velocity*) dan posisi (*position*) setiap partikel.

$$V_{id}^{k+1} = w \times V_{id}^k + c_1 \times rand_1 \times (P_{id} - X_{id}) + c_2 \times rand_2 \times (G_{id} - X_{id}) \dots\dots (2.2)$$

$$X_{id}^{k+1} = X_{id}^k + V_{id}^{k+1} \dots\dots\dots (2.3)$$

Keterangan:

V_{id} = Komponen kecepatan individu ke i pada d dimensi

X_{id} = Posisi individu i pada d dimensi

W = Parameter inertia *weight*

$c_1 c_2$ = Konstanta akselerasi (*learning rate*), nilainya antara 0 sampai 1

$rand_{1,2}$ = Parameter random antara 0 sampai 1

P_{id} = $Pbest$ (*local best*) individu i pada d dimensi

G_{id} = $Gbest$ (*global best*) pada d dimensi

2.4 Klasifikasi

Klasifikasi merupakan salah satu teknik dalam *data mining* dan termasuk ke dalam jenis *supervised methods*. *Supervised methods* merupakan pelabelan *class* yang telah didefinisikan berdasarkan nilai karakteristik atribut. Pada *dataset*, label *class* telah didefinisikan sehingga algoritma pembelajar dapat mempelajari pola atau hubungan antar *atribut* dengan label *class*, sehingga klasifikasi disebut dengan *supervised methods*. Klasifikasi dilakukan dua tahap proses, yaitu proses pembelajaran, dimana proses untuk menemukan pola data dan proses klasifikasi, dimana pola klasifikasi digunakan untuk memprediksi label *class*.

Pada tahap pertama yaitu tahap pembelajaran, algoritma *machine learning* digunakan untuk menganalisis pola atau hubungan dari *atribut* data dalam menentukan label *class* atau target *atribut*. Pola hasil pembelajaran kemudian dilakukan pengujian untuk dinilai performanya sebagai *classifier*. Jika hasil klasifikasi dinilai memiliki performa yang baik, maka pola klasifikasi dapat diterapkan pada kasus nyata. Klasifikasi memiliki beberapa metode antara lain *k-NN* (*k-Nearest Neighbor*), *SVM* (*support vector machine*), *NB* (*Naïve Bayes*) dan *DT* (*decision tree*).

2.4.1 *K-Nearest Neighbor*

Algoritma *K-Nearest Neighbor* (*K-NN*) merupakan metode klasifikasi objek berdasarkan data pembelajaran jarak yang paling dekat dengan objek tersebut. Data pembelajaran diproyeksikan ke ruang berdimensi banyak, dimana masing – masing dimensi merepresentasikan fitur dari data.

Algoritma KNN merupakan algoritma *supervised* yang bertujuan untuk menemukan pola baru dalam data dengan menghubungkan pola data yang sudah ada dengan data yang baru. Tujuan dari algoritma KNN adalah untuk mengklasifikasi objek baru berdasarkan *atribut* dan *training samples*. Dimana hasil dari *sampel* uji yang baru diklasifikasikan berdasarkan mayoritas dari kategori pada KNN. Pada proses pengklasifikasian, algoritma ini tidak menggunakan model apapun untuk dicocokkan dan hanya berdasarkan pada memori.

Algoritma KNN mengklasifikasikan berdasarkan ketetanggaan sebagai nilai prediksi dari contoh data uji yang baru. Jarak yang digunakan adalah jarak *Euclidean Distance*. Jarak *Euclidean* adalah jarak yang paling umum digunakan pada data *numeric*. Algoritma KNN merupakan algoritma yang menentukan nilai jarak pada data pengujian (*testing*) dengan data pelatihan (*training*) berdasarkan nilai terkecil dari nilai ketetanggaan terdekat. Algoritma *k*-NN menggunakan klasifikasi ketetanggaan sebagai nilai prediksi dari sampel uji yang baru. Jarak *Euclidean* adalah jarak yang paling umum digunakan pada data numerik. *Euclidean distance* didefinisikan sebagai berikut. (Krisandi, 2013)

$$d(x_i, x_j) = \sqrt{\sum_r^n (a_r(x_i) - a_r(x_j))^2} \dots\dots\dots (2.4)$$

Keterangan :

$d(x_i, x_j)$ = Jarak *Euclidean* (*Eucledean Distance*)

(x_i) = *Record* ke-i

(x_j) = *Record* ke-j

(a_r) = Data ke-r

i, j = 1,2,3.....n

Algoritma *k*-NN mencari nilai jarak pada data *testing* dengan data *training* berdasarkan nilai terkecil dari nilai ketetanggaan terdekat. Didefinisikan sebagai berikut (Krisandi, 2013).

$$D_{nn}(C_1, C_2) = \min_{1 \leq i \leq r, 1 \leq j \leq s} d(y_i, z_j) \dots\dots\dots (2.5)$$

Tabel 2.10 menunjukkan mengenai kelebihan dan kekurangan metode *K-NN* pada saat proses klasifikasi.

Tabel 2.10 Kelebihan dan kekurangan metode K-NN

Kelebihan	Kekurangan
• Mampu mengklasifikasikan data yang bersifat <i>non-linier</i> • Mudah dipahami dan diimplementasikan • <i>Efektif</i> meskipun data <i>training</i>-nya berjumlah besar • Mampu memprediksi label <i>class</i> berdasarkan <i>instance</i> paling dekat.	• Menentukan Parameter <i>k</i> (Jumlah Tetangga Terdekat) • Tidak bisa menangani <i>Missing Value</i> • <i>Sensitif</i> terhadap data <i>Outlier</i> • Rentan terhadap atribut yang <i>non-Informatif</i> • Rentan terhadap jumlah data yang besar • Waktu komputasi yang lama

Sumber: Mutrofin, 2010

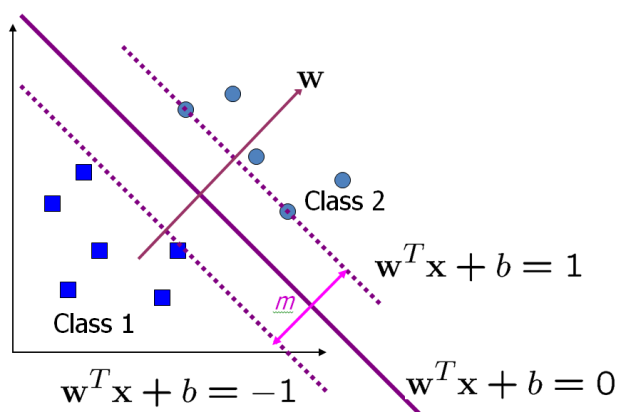
2.4.2 Support Vector Machine

Support Vector Machine (SVM) adalah suatu metode prediksi yang menggunakan teknik klasifikasi maupun *regresi* (Santosa, 2007). SVM memiliki prinsip dasar *linier classifier* yaitu mengklasifikasikan data dengan cara memisahkan secara *linier*, namun SVM telah dikembangkan agar dapat bekerja pada *problem non-linier* dengan memasukkan konsep *kemel* pada ruang kerja berdimensi tinggi.

SVM memisahkan data dengan cara membangun sebuah bidang pemisah (*hyperplane*). *Hyperplane* yang baik tidak hanya mampu digunakan untuk memisahkan data, akan tetapi juga memiliki batasan (*margin*) yang besar. Pencarian *hyperplane* inilah yang menjadi inti dari *support vector machine*. Namun, dalam proses pencarian *hyperplane* tersebut, akan muncul permasalahan baru yaitu sebuah formula yang sangat sulit untuk dipecahkan disebut dengan permasalahan *Quadratic Programming*.

Konsep SVM sebagai solusi untuk mencari *hyperplane* terbaik yang berfungsi sebagai pemisah dua buah *class* yaitu + dan - pada input *space* (*discrimination boundaries*). *Margin* adalah jarak antara *hyperplane* dengan pola terdekat dari masing-masing *class*. Pola yang paling dekat ini disebut sebagai *support vector*. Mencari posisi *hyperplane* ini merupakan inti dari proses pembelajaran pada SVM. Pada ruang berdimensi tinggi, akan dicari *hyperplane* yang dapat memaksimalkan jarak (*margin*) antara *class* data. *Hyperplane* klasifikasi *linier* SVM dinotasikan sebagai berikut (Santosa, 2007).

$$f(x) = w^T x + b \dots\dots\dots (2.6)$$



Gambar 2.2 Model klasifikasi SVM

Tabel 2.11 menunjukkan mengenai kelebihan dan kekurangan metode SVM pada saat proses klasifikasi.

Tabel 2.11 Kelebihan dan kekurangan metode SVM

Kelebihan	Kekurangan
• Mampu mengklasifikasikan data yang tidak masuk dalam <i>class</i> atau kategori • Mampu menghadapi masalah pada suatu pola (<i>Curse of dimensionality</i>). • SVM dapat diimplementasikan dengan mudah.	• SVM sangat sulit jika menggunakan data dengan jumlah yang besar • SVM perlu ditambahkan teknik pengolahan lanjut supaya mampu menyelesaikan masalah pada data <i>multiclass</i>.

Sumber: Nugroho, 2003

2.4.3 Naive Bayes

Naive Bayes merupakan sebuah metode klasifikasi berdasarkan perhitungan dari sekumpulan probabilitas dengan menjumlahkan frekuensi dan kombinasi nilai dari *dataset* yang diberikan. Algoritma menggunakan teorema *Bayes* dan mengasumsikan bahwa semua *atribut* tidak saling ketergantungan (*independen*). Karena diasumsikan sebagai variable *independent*, maka hanya varians dari suatu variabel dalam sebuah *class* yang dibutuhkan untuk menentukan klasifikasi, bukan keseluruhan dari *matriks kovarians*.

Kelebihan menggunakan *Naive Bayes* adalah bahwa metode ini hanya membutuhkan jumlah data pelatihan (*training data*) yang sedikit untuk menentukan *estimasi* parameter yang diperlukan dalam proses pengklasifikasian. Berikut ini persamaan umum dari *Naive Bayes* (Indranandita, 2008).

$$P(Y = y | X_1, X_2, \dots, X_p) = \frac{P(y)P(X_1, X_2, \dots, X_p | y)}{P(X_1, X_2, \dots, X_p)} \dots \dots \dots (2.7)$$

Dengan asumsi *independent* $X_1, X_2, \dots, X_p$ didapat
 $P(X_1, X_2, \dots, X_p | y) = P(X_1 | y)P(X_2 | y)P(X_3 | y) \dots P(X_p | y)$ sehingga probabilitas masing-masing *class* dalam klasifikasi *naïve bayes* adalah

$$P(y | X_1, X_2, \dots, X_p) = \frac{P(y)P(X_1 | y)P(X_2 | y) \dots P(X_p | y)}{P(X_1, X_2, \dots, X_p)} = \frac{P(y) \prod_{i=1}^p P(X_i | y)}{P(X_1, X_2, \dots, X_p)} \dots \dots (2.8)$$

Keterangan:

X = Data dengan *class* yang belum diketahui

Y = Hipotesis data merupakan suatu *class* spesifik

$P(Y|X)$ = Probabilitas hipotesis Y berdasar kondisi X (*posteriori probabilitas*)

$P(Y)$ = Probabilitas hipotesis Y (*prior probabilitas*)

$P(X|Y)$ = Probabilitas X berdasarkan kondisi pada hipotesis Y

$P(X)$ = Probabilitas X

Tabel 2.12 menunjukkan mengenai kelebihan dan kekurangan metode *Naïve Bayes* pada saat proses klasifikasi.

Tabel 2.12 Kelebihan dan kekurangan metode *Naïve Bayes*

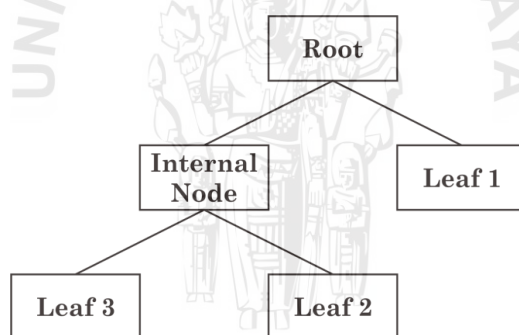
Kelebihan	Kekurangan
• Mudah untuk dipahami • Waktu komputasi Lebih cepat • Mampu menangani data <i>kuantitatif</i> dan data <i>diskrit</i> • Memerlukan jumlah data pelatihan yang sedikit • Mampu menangani <i>missing value</i> • Efisiensi pada ruang memori penyimpanan • Mampu menyelesaikan masalah data <i>imbalanced</i>	• Jumlah data training yang sedikit menyebabkan kurang akurat untuk membuktikan kebenaran prediksi klasifikasi • Tidak berlaku jika probabilitas kondisional nya adalah 0 (nol), apabila nol maka probabilitas prediksi akan bernilai nol juga • Mengasumsikan variabel bebas

Sumber: Manalu, 2017

2.4.4 Decision Tree

Decision tree atau pohon keputusan merupakan salah satu metode klasifikasi yang paling terkenal karena mudah untuk dipahami dan diterapkan oleh manusia dan komputer. Proses dari *decision tree* adalah mengubah bentuk data (*tabel*) menjadi pola model pohon, mengubah pola model pohon menjadi aturan (*rule*), dan menyederhanakan *rule*. Manfaat utama dari penggunaan *decision tree* adalah kemampuannya untuk *membreak down* proses pengambilan keputusan yang kompleks menjadi lebih sederhana sehingga lebih mudah dalam menganalisa pola keputusan. *Decision tree* juga berguna untuk mengobservasi data, menemukan hubungan tersembunyi antara variabel *input* dengan sebuah variabel target.

Pola model *decision tree* dibentuk oleh *node-node* antara lain *root node*, *intemal node* dan *leaf node*. *Root node* merupakan *node* yang berada pada posisi puncak yang menghasilkan cabang kriteria. Cabang kriteria disebut juga dengan *intemal node* yang berasal dari *node* atasnya dan memiliki cabang kebawah. *Leaf node* merupakan target digunakan dalam mengambil keputusan. Konsep dari pohon keputusan ditampilkan pada gambar dibawah ini.



Gambar 2.3. Konsep dasar pohon keputusan

Pada *decision tree* untuk menentukan *node* sebagai pola model pohon keputusan dibutuhkan sebuah algortima. Algoritma yang biasanya digunakan antara lain ID3, C4.5 dan C5.0. Tabel 2.13 merupakan perbedaan antara algoritma ID3, C4.5 dan C5.0 seperti ditunjukkan seperti table di bawah ini.

Tabel 2.13 Perbandingan algoritma ID3, C4.5 dan C5.0

	ID3	C4.5	C5.0
Type of data	Kategori	Kontinue dan kategori	Kontinue dan kategori, <i>Dates, Times, Timestamps</i>
Speed	Lambat	Lebih cepat dari ID3	Paling cepat dari ID3 dan C4.5
Pruning	<i>No</i>	<i>Pre-pruning</i>	<i>Pre-pruning</i>
Boosting	Tidak support	Tidak support	Support
Missing Value	Tidak dapat menyelesaikan	Tidak dapat menyelesaikan	Dapat menyelesaikan
Formula	<i>information entropy</i>	<i>split info and gain ratio</i>	<i>information Gain</i>

Sumber: Kumar, 2015

Tabel 2.14 menunjukkan mengenai kelebihan dan kekurangan metode *Decision Tree* pada saat proses klasifikasi.

Tabel 2.14 Kelebihan dan kekurangan metode *Decision Tree*

Kelebihan	Kekurangan
• Waktu komputasi yang cepat • Pengambilan keputusan lebih simple dan spesifik • Efisien dalam merancang model • Fleksibel sehingga akan meningkatkan kualitas keputusan yang dihasilkan • Dapat mengolah data numerik • Model yang dihasilkan mudah dipahami • Dapat menyelesaikan masalah data <i>imbalanced</i>	• Terjadi <i>overlapping</i> • Pengakumulasian jumlah kesalahan dari setiap tingkat dalam sebuah pohon keputusan yang besar • Kesulitan dalam mendesain pohon keputusan yang optimal • Hasil kualitas keputusan sangat tergantung pada bagaimana pohon tersebut didesain. • Kebutuhan ruang penyimpanan yang besar

Sumber: Astuti, 2019

2.4.4.1 Algoritma ID3

Algoritma Pohon Keputusan ID3 (*Iterative Dichotomiser 3*) merupakan sebuah metode yang digunakan untuk membuat pola model pohon keputusan. Algoritma ini menggunakan konsep dari *entropy informasi* untuk menentukan setiap node. Algoritma ini melakukan pencarian secara menyeluruh pada semua kemungkinan pohon keputusan. Secara ringkas, langkah kerja Algoritma ID3 dapat digambarkan sebagai berikut:

1. Hitung *Entropy* (H) dan *Information gain* (G) dari setiap atribut dengan menggunakan rumus sebagai berikut. (Hssina, 2013)

$$Entropy(S) = - \sum_{i=1}^n p(s_i) \log_2(p(s_i)) \quad \dots\dots\dots (2.9)$$

$$Gain(S, A_i) = Entropy(S) - \sum_{\alpha \in A_i} \frac{|S_{\alpha}|}{|S|} Entropy(S_{\alpha}) \quad \dots\dots\dots (2.10)$$

Keterangan:

S = Himpunan kasus

A = Atribut

n = Jumlah partisi atribut A

$|S_{\alpha}|$ = Jumlah kasus pada partisi ke- i

$|S|$ = Jumlah kasus dalam S

p_i = Proporsi dari S_i terhadap S

2. Bentuk simpul yang berisi atribut tersebut.
3. Proses terus dilakukan perhitungan information gain sampai seluruh data telah masuk dalam kelas yang sama. Atribut yang telah dipilih tidak diikuti lagi dalam perhitungan nilai information gain.

2.4.4.2 Algoritma C4.5

Algoritma pohon keputusan C4.5 adalah pengembangan dari algoritma ID3. Sehingga algoritma C4.5 memiliki prinsip dasar kerja yang sama dengan algoritma ID3. Perbedaan utama C4.5 dari ID3 adalah:

1. C4.5 mampu menangani atribut kontinyu dan diskrit.
2. C4.5 mampu menangani *training* data dengan *missing value*.
3. Hasil pohon keputusan C4.5 akan *pruning* setelah dibentuk.
4. Pemilihan *node* yang dilakukan dengan menggunakan *Gain Ratio*.

Algoritma C4.5 menggunakan *gain ratio* untuk memperbaiki *information gain*, dengan rumus *gain ratio* (Merbah, 2014).

$$GainRatio(S, A) = \frac{Gain(S, A)}{Split\ inf\ o(S, A)} \quad \dots\dots\dots (2.11)$$

Atribut dengan nilai *Gain Ratio* tertinggi dipilih sebagai *node root* untuk simpul. Dengan *gain* adalah *information gain*. Pendekatan ini menerapkan normalisasi pada *information gain* dengan menggunakan apa yang disebut sebagai *split information*. *SplitInfo* menyatakan *entropy* atau informasi potensial dengan rumus:

$$\text{Split info } o(S, A) = - \sum_{i=1}^n \frac{|S_{\alpha}|}{|S|} \log_2 \frac{|S_{\alpha}|}{|S|} \dots\dots\dots (2.12)$$

Keterangan:

S = Himpunan kasus

A = Atribut

n = Jumlah partisi atribut A

$|S_{\alpha}|$ = Jumlah kasus pada partisi ke- i

$|S|$ = Jumlah kasus dalam S

2.4.4.3 Algoritma C5.0

Algoritma C5.0 menjadi penyempura dari algoritma sebelumnya yaitu C4.5 dan ID3. Dibandingkan dengan C4.5, algoritma C5.0 lebih cepat dalam mengambil keputusan dan lebih efektif dalam membangun struktur pohon keputusan. *Information gain* atau biasa disebut *gain info* adalah kriteria pemisahan yang menggunakan pengukuran *entropy*. Untuk memperoleh *information gain* dari suatu atribut dibutuhkan *entropy* keseluruhan *class* atau *Entropy (S)*. *Entropy (S)* adalah banyaknya jumlah *bit* yang dibutuhkan untuk mengekstrak suatu *class* dari sekumpulan data acak pada ruang sample S . Semakin kecil nilai *entropy* maka memiliki hasil yang baik untuk digunakan dalam mengekstraksi suatu *class*. Nilai informasi yang dinyatakan sebagai panjang kode adalah $p \log_2 p$ yang memiliki probabilitas p . *Entropy* merupakan suatu parameter yang digunakan untuk mengukur heterogenitas (keberagaman) dari suatu kumpulan data sampel. Semakin besar nilai *entropy* maka jumlah kriteria pada data sampel semakin heterogen. Secara matematis *entropy* dirumuskan sebagai berikut (Han et al, 2011).

$$H(S) = \sum_{i=1}^n -p(s_i) \log_2(p(s_i)) \dots\dots\dots (2.13)$$

Setelah mendapatkan nilai *entropy* selanjutnya mencari nilai *information gain*. *Information gain* digunakan untuk mengukur efektivitas karakteristik *atribut* dalam mengklasifikasikan *class*. Persamaan 2.14 digunakan untuk menghitung *information gain* sebagai berikut:

$$IG(S, A_i) = H(S) - \sum_{\alpha \in A_i} \frac{|S_\alpha|}{|S|} H(S_\alpha) \dots\dots\dots (2.14)$$

Keterangan:

H = Entropy

S = Himpunan kasus

A = Atribut

n = Jumlah partisi atribut A

$|S_\alpha|$ = Jumlah kasus pada partisi ke- i

$|S|$ = Jumlah kasus dalam S

p_i = Proporsi dari S_i terhadap S

Dengan kata lain, *Gain* (A) adalah jumlah keseluruhan nilai kriteria pada atribut A terhadap jumlah total data sampel. Atribut dengan nilai gain terbesar dipilih sebagai simpul atribut (*node root*). Sedangkan cabang (*node intemal*) dibuat berdasarkan pada kriteria pada node sebelumnya hingga menghasilkan keputusan (*leaf node*).

2.5 Metacost

Cost sensitive leaming adalah sebuah metode pembelajaran yang dapat menyelesaikan permasalahan data *imbalanced*. Konsep metode ini adalah mengurangi *cost* dari *missclassification* sehingga dapat meningkatkan performa pada *classifier*. Metode ini mengansumsikan bahwa minoritas *class* mempunyai *cost* lebih besar dari pada mayoritas *class* sehingga *machine leaming* akan fokus mengkalsifikasikan minoritas *class* dengan akurat. *Output* dari *cost sensitive leaming* berupa nilai *cost function*. *Cost function* merupakan bobot yang diberikan karena kesalahan klasifikasi.

Cost fuction dinotasikan sebagai berikut $C(i, j)$ dimana *missclasification* terjadi ketika aktual *class* i diprediksi menjadi *class* j . *Cost sensitive decision tree* dapat diintegrasikan

dalam empat cara. Pertama, *cost function* digunakan sebagai kriteria untuk pemilihan atribut yang akan displit. Kedua, *cost function* digunakan untuk memilih bagian tree yang akan dipangkas (*pruning*). Ketiga, *cost function* digunakan memanipulasi bobot dari data *training* sehingga menghasilkan tree dengan *cost* minimal. Keempat, *cost function* digunakan sebagai kriteria untuk pelabelan *leaf node*.

Cost sensitive learning dapat dikategorikan menjadi dua kriteria. Kriteria pertama yaitu *direct method*, metode ini antara lain ICET dan *cost sensitive decision tree*. Kriteria kedua yaitu *Meta Learning*, dibagi menjadi dua *thresholding* dan *sampling*. Untuk struktur dari *cost sensitive learning* dapat dilihat pada tabel 2.15.

Tabel 2.15 Struktur Model *Cost Sensitive Decision Tree*

Cost Sensitive Learning	
Direct Method	Meta Learning
1. ICET	Thresholding
2. Cost Sensitive Decision Tree	1. Metacost
	2. Cost Sensitive Classifier
	3. Cost Sensitive Naïve Bayes
	4. Emperical Thresholding
	Sampling
	1. Coasting
	2. Weighting

Sumber: Ling, 2008

Metacost merupakan metode *cost sensitive learning* menggunakan teknik *thresholding meta learning*. Metode ini dapat diimplementasikan pada *classifier decision tree*. Prinsip kerja dari metode ini adalah menghitung *probability* setiap *leaf node*. Ketika nilai *probability* pada *leaf node* > 1 , ada dua cara dalam mengevaluasi model yaitu *pruning* dan *relabeling* untuk mendapatkan model dengan *minimum cost*. *Pruning* merupakan teknik pemangkasan cabang pohon sedangkan teknik *relabeling* merupakan pemberian *label class* baru pada *leaf node* berdasarkan jumlah *class* paling banyak (Domingos, 1999).

Algoritma : Metacost

Output : Decision Tree

Step :

For i = 1 to m

Let S_i be a leaf node of S with n node.

Let M_i = Model produced by applying L to S_i .

For each node x in S

For each class j

$$\text{Let } P(j|x) = \frac{1}{\sum_i 1} \sum_i P(j|x, M_i)$$

if $P(j|x, M_i) = 0$ is produced by M_i

Else $P(j|x, M_i) > 1$

$$\text{Let } x\text{'s class} = \underset{i}{\operatorname{argmin}} \sum_j P(j|x) C(i, j)$$

Let M = Model produced by applying L to S .

Return M

Keterangan:

S = Training set

L = Algoritma classification learning

C = Cost matrix

C = Jumlah leaf node model

m = Jumlah node setiap leaf node

2.6 Validation

Cross validation (CV) adalah sebuah metode analisa yang dapat digunakan untuk mengevaluasi kinerja dari *classifier* dengan cara, *dataset* dibagi menjadi dua *subset* yaitu data pembelajaran dan data pengujian. Selanjutnya pemilihan jenis CV berdasarkan pada besarnya ukuran *dataset*. *K-fold* CV merupakan metode yang baik digunakan karena dapat meminimalkan waktu komputasi dengan tetap mempertimbangkan keakuratan estimasi.

Cara kerja dari *k-fold cross validation* adalah membagi *dataset* menjadi dua kelompok yaitu data *training* dan data *testing*, kemudian dilakukan proses *testing* yang dilakukan pengulangan sebanyak k kali. Hasil *testing* itu kemudian dirata-rata untuk menghasilkan sebuah nilai *accuracy*. Nilai k yang terlalu kecil dapat mempengaruhi nilai *accuracy* menjadi

rendah. Hal tersebut disebabkan karena nilai k yang kecil, mudah terpengaruh oleh *noise*. Sedangkan, nilai k yang terlalu besar menghasilkan proses yang kurang efektif. Hal ini disebabkan nilai k yang besar, membutuhkan waktu yang cukup lama dalam melakukan pengujian. Metode *10 fold cross validation* menjadi metode standar untuk pembelajaran dan pengujian data (Catal & Diri, 2009). Berikut contoh ilustrasi *10-fold cross validation*.

Tabel 2.16 Ilustrasi *10 fold cross validation*

Validasi ke-n	Data Set									
	1	2	3	4	5	6	7	8	9	10
1	test	train	train	train	train	train	train	train	train	train
2	train	test	train	train	train	train	train	train	train	train
3	train	train	test	train	train	train	train	train	train	train
4	train	train	train	test	train	train	train	train	train	train
5	train	train	train	train	test	train	train	train	train	train
6	train	train	train	train	train	test	train	train	train	train
7	train	train	train	train	train	train	test	train	train	train
8	train	train	train	train	train	train	train	test	train	train
9	train	train	train	train	train	train	train	train	test	train
10	train	train	train	train	train	train	train	train	train	test

Sumber: Catal & Diri, 2009

Berdasarkan Tabel 2.16 menjelaskan bahwa nilai *fold* yang digunakan *10-fold cross validation*. Berikut adalah langkah-langkah pengujian data menggunakan *10-fold cross validation*.

- Membagi *dataset* menjadi 10 bagian, yaitu $D_1, D_2, D_3, D_4, \dots, D_{10}$ $t = (1, 2, 3, \dots, 10)$ digunakan sebagai data *training* dan *dataset* lainnya sebagai data *testing*.
- Nilai *accuracy* dihitung pada setiap *iterasi* (*iterasi-1, iterasi-2, \dots, iterasi-10*) kemudian dihitung rata-rata nilai *accuracy* dari seluruh *iterasi* untuk mendapatkan nilai *accuracy* dari data keseluruhan.

2.7 Evaluasi Performansi Model Klasifikasi

Pengukuran terhadap kinerja suatu sistem klasifikasi merupakan hal yang penting. Kinerja sistem klasifikasi menggambarkan seberapa baik sistem dalam mengklasifikasikan

data. Dalam tahap ini ada dua tahap pengukuran performansi pengujian model klasifikasi *cost sensitive decision tree* yaitu dengan menggunakan metode *confusion matrix* dan *cost matrix*.

2.7.1 Confusion Matrix Multi-class

Confusion matrix merupakan salah satu metode yang dapat digunakan untuk mengukur kinerja *classifier* dengan baik. Pada dasarnya *confusion matrix* mengandung informasi dengan membandingkan hasil prediksi yang dilakukan oleh sistem dengan nilai aktual yang sebenarnya. Berikut merupakan gambaran hasil evaluasi klasifikasi *multi-class* seperti ditampilkan pada Tabel 2.17 berikut:

Tabel 2.17 *confusion matrix multi-class*

Class Actual	Class Prediction					Total
	1	2	3	...	K	
1	x_{11}	x_{12}	x_{13}	...	x_{1k}	n_1
2	x_{21}	x_{22}	x_{23}	...	x_{2k}	n_2
3	x_{31}	x_{32}	x_{33}	...	x_{3k}	n_3
...	...	...	...	...	...	...
k	x_{k1}	x_{k2}	x_{k3}	...	x_{kk}	n_k
Total	n_1	n_2	n_3	...	n_k	N_{total}

Sumber : (Akbar, Yudhistira dan Cholissodin, 2014)

$$TP = x_{11} + x_{22} + x_{33} + \dots + x_{kk}$$

$$TN = (x_{11} + x_{22}) + (x_{11} + x_{33}) + (x_{22} + x_{33}) + (x_{11} + x_{kk}) + (x_{22} + x_{kk}) + (x_{33} + x_{kk})$$

$$FP = (x_{21} + x_{31} + \dots + x_{k1}) + (x_{12} + x_{32} + \dots + x_{k2}) + (x_{13} + x_{23} + \dots + x_{k3}) + \dots$$

$$FN = (x_{12} + x_{13} + \dots + x_{1k}) + (x_{21} + x_{23} + \dots + x_{2k}) + (x_{31} + x_{32} + \dots + x_{3k}) + \dots$$

$$n_1 = x_{11} + x_{12} + x_{13} + \dots + x_{1k}$$

$$n_2 = x_{21} + x_{22} + x_{23} + \dots + x_{2k}$$

$$n_3 = x_{31} + x_{32} + x_{33} + \dots + x_{3k}$$

$$n_k = x_{k1} + x_{k2} + x_{k3} + \dots + x_{kk}$$

$$N_{total} = n_1 + n_2 + n_3 + \dots + n_k$$

1. *True Positive* (TP) menunjukkan bahwa *class* yang diprediksi positif sedangkan *class* aktualnya adalah positif.
2. *True Negatif* (TN) menunjukkan bahwa *class* yang diprediksi negatif sedangkan *class* aktualnya adalah negatif.
3. *False Positif* (FP) menunjukkan bahwa *class* yang diprediksi negatif sedangkan *class* aktualnya adalah positif.
4. *False Negatif* (FN) menunjukkan bahwa *class* yang diprediksi positif sedangkan *class* aktualnya adalah negatif.

Pengukuran performa dari data *imbalanced* dapat dilakukan dengan menggunakan perhitungan nilai *recall*, *precision* dan *f-measure*. *Accuracy* klasifikasi menunjukkan performansi model klasifikasi secara keseluruhan, dimana semakin tinggi *accuracy* maka tingkat ketepatan antara nilai prediksi mendekati dengan nilai aktual (Powers, 2011).

$$Accuracy = \frac{TP}{N_{total}} \dots\dots\dots (2.15)$$

Recall digunakan untuk mengukur tingkat positif (*minor*) benar atau mengukur tingkat keberhasilan sistem dalam menemukan lagi informasi pada data acak. *Precision* adalah tingkat negative (*mayor*) benar atau tingkat ketepatan antara informasi yang diminta oleh pengguna dengan jawaban yang diberikan oleh sistem. Persamaa *recall* dan *Precision* adalah sebagai berikut (Powers, 2011).

$$Recall = \frac{TN}{TN + FP} \dots\dots\dots (2.16)$$

$$Precision = \frac{TP}{(TP + FP)} \times 100\% \dots\dots\dots (2.17)$$

F-measure merupakan kombinasi dari nilai *recall* dan *precision* yang digunakan untuk Mengukur performansi prediksi pada *classier* berdasarkan gabungan hasil nilai *precision* dan *recall*. Dengan menggunakan persamaan seperti berikut (Cao dkk dalam Sain, 2013).

$$F - Measure = 2 \times \frac{sensitivity \times specificity}{sensitivity + specificity} \dots\dots\dots (2.18)$$

2.7.2 Cost Matrix Multi-class

Cost matrix adalah representasi dari hasil kesalahan prediksi (*misclassification*) dalam *machine learning*. *Missclassification cost* pada *multiclass* dinotasikan sebagai berikut C_{12} , C_{13} , C_{21} , C_{23} , C_{31} , C_{32} , C_{kk} . Ketika klasifikasi benar dinotasika C_{11} , C_{22} , C_{33} , C_{kk} nilai *cost* adalah 0. *Cost matrix* ditampilkan Tabel 2.18.

Tabel 2.18 *cost matrix multi-class*

Class Actual	Class Prediction					Total
	1	2	3	...	K	
1	C_{11}	C_{12}	C_{13}	...	C_{1k}	n_1
2	C_{21}	C_{22}	C_{23}	...	C_{2k}	n_2
3	C_{31}	C_{32}	C_{33}	...	C_{3k}	n_3
...	...	...	...	...	...	...
k	C_{k1}	C_{k2}	C_{k3}	...	C_{kk}	n_k
Total	n_1	n_2	n_3	...	n_k	N_{total}

Total *cost* menunjukkan performansi model klasifikasi secara keseluruhan berdasarkan hasil *misclassification cost*, dimana semakin kecil nilai *cost* klasifikasi . Hal ini berarti semakin baik performansi model klasifikasi. Persamaan untuk menghitung total *cost* ditunjukan seperti berikut (George, 2017).

$$\text{Cost} = N_{total} [(C_{FN} + FN) * (C_{FP} + FP)] \dots\dots\dots (2.19)$$