

## Prediksi Klasifikasi jumlah Pembaca Sebuah Artikel pada Jurnal Biram Samtani Dengan Metode Bayesian Classification

Richasanty Septima S<sup>1\*</sup>, Ira Zulfa<sup>2</sup>, Husna Gemasih<sup>3</sup>, Mahmuda Saputra<sup>4</sup>, Al Israk<sup>5</sup>

<sup>1,2,3,4,5</sup>Universitas Gajah Putih, Indonesia

E-mail: richaseptima@gmail.com

### Abstract

One of the open access Open Journal System (OJS)-based reading media is the Biramsamtani Journal, which is the official publication media of scientific articles published by the journal manager, namely the Institute for Research and Community Service of Gajah Putih University, Central Aceh, Indonesia. Until now, the number of scientific articles that have been published is 76 articles with a total of 26,938 viewers read in 10 editions, 7 volumes, and for 5 years of publication. In the Biram Samtani Journal, an uneven and erratic number of readers was found. Where the causative factors that cannot be determined are due to too many abstracts, the position of the article in the published edition or the category of the article that needs to be adjusted to the month of publication. This problem will be studied using data mining by analyzing data using Bayesian Classification. From the test results obtained to predict the number of readers, the best obtained is model 2 with the composition of training data and 80:20 test data with an accuracy of 79.92%, accuracy of 80.15% and recall of 72.36% while for the classification of the number of readers in Biramsamtani journal articles based on the probability value of 47.7% (Low), 30.76% (High), and 21.53% (Medium).

**Keywords:** Journal, Number of Readers, Data Mining, Bayesian Classification

### Abstrak

Salah satu media bacaan berbasis Open Journal System (OJS) yang akses terbuka adalah Jurnal Biramsamtani, yang merupakan media publikasi resmi artikel ilmiah yang diterbitkan oleh pengelola jurnal, yaitu Lembaga Penelitian dan Pengabdian Masyarakat Universitas Gajah Putih, Aceh Tengah, Indonesia. Hingga saat ini, jumlah artikel ilmiah yang telah diterbitkan adalah 76 artikel dengan total 26.938 pembaca dalam 10 edisi, 7 volume, dan selama 5 tahun publikasi. Di Jurnal Biram Samtani, ditemukan jumlah pembaca yang tidak merata dan tidak teratur. Faktor penyebab yang tidak dapat ditentukan meliputi terlalu banyaknya abstrak, posisi artikel dalam edisi yang diterbitkan, atau kategori artikel yang perlu disesuaikan dengan bulan publikasi. Masalah ini akan diteliti menggunakan data mining dengan menganalisis data menggunakan Klasifikasi Bayesian. Dari hasil pengujian diperoleh prediksi jumlah pembaca, model terbaik yang diperoleh adalah model 2 dengan komposisi data pelatihan dan data uji 80:20, dengan akurasi 79,92%, presisi 80,15%, dan recall 72,36%. Sementara untuk klasifikasi jumlah pembaca pada artikel jurnal Biramsamtani berdasarkan nilai probabilitas adalah 47,7% (Rendah), 30,76% (Tinggi), dan 21,53% (Sedang).

**Kata kunci:** Jurnal, Jumlah Pembaca, Data Mining, Klasifikasi Bayesian

## 1. Pendahuluan

Teknologi informasi dan komunikasi pada saat ini sangat cepat berkembang dalam kehidupan sehari-hari manusia. Dimana hampir semua aspek kebutuhan manusia pada saat sangat tergantung dengan teknologi informasi, terutama pada bidang media baca dimana sebelumnya media baca terbatas pada media cetak saja, namun saat ini media tersebut sudah memanfaatkan TIK berbasis online, sehingga dapat menjangkau pembaca

yang tak terbatas oleh ruang dan waktu. Kelebihan dan manfaat yang didapat dari media baca *online* adalah pengelola media dapat melihat seberapa banyak pembaca pada setiap artikel yang telah diterbitkan, sehingga dapat menjadi acuan langkah dan strategi pembuatan artikel yang lebih menarik bagi pembaca. Apalagi saat ini pada umumnya media jurnal di Indonesia kebanyakan menggunakan media yang bersifat *open access* yang bebas diakses dan berbasis *Open Journal System* (OJS) dan juga terindeks pada *Data of Open Acces Journal* (DOAJ).

Berdasarkan pada DOAJ atau direktori daring yang mengindeksikan dan menyediakan akses ke jurnal yang berkualitas, *open acces*, dan *peer-reviewed*, bahwa jurnal Indonesia menduduki posisi pertama dalam jumlah jurnal yang terindeks DOAJ dengan jumlah 1.867 pada tahun 2021. Disusul dengan Brazil yang saat itu memiliki jumlah jurnal 1.634 jurnal, kemudian Turki 412 jurnal, dan China 177 Jurnal (penerbitdeepublish.com). Salah satu media baca berbasis *Open Journal System* (OJS) yang bersifat *open access* adalah Jurnal Biramsamtani, yang merupakan media publikasi resmi artikel ilmiah yang diterbitkan oleh pengelola jurnal yaitu Lembaga Penelitian dan Pengabdian kepada Masyarakat Universitas Gajah Putih, Aceh Tengah, Indonesia. Sampai saat ini jumlah artikel ilmiah yang telah diterbitkan berjumlah 76 artikel dengan total jumlah *viewer* 26.938 dibaca dalam 10 edisi, 7 volume, dan selama 5 tahun terbit.

Pada edisi ke-3 tahun 2019 terdapat artikel dengan jumlah 1.126 *view*, sedangkan ada juga artikel pada edisi yang sama hanya memperoleh 98 *view*. Kemudian pada edisi Juli tahun 2020 terdapat artikel dengan jumlah 9.678 *view*, namun pada artikel lain hanya memperoleh antara 63 – 700 saja. Fenomena ini menunjukkan jumlah pembaca yang tidak merata dan tidak menentu. Dimana faktor penyebab yang belum dapat ditentukan apakah karena abstrak yang terlalu banyak, posisi artikel pada edisi terbitan atau kategori artikel yang perlu disesuaikan dengan bulan terbit.

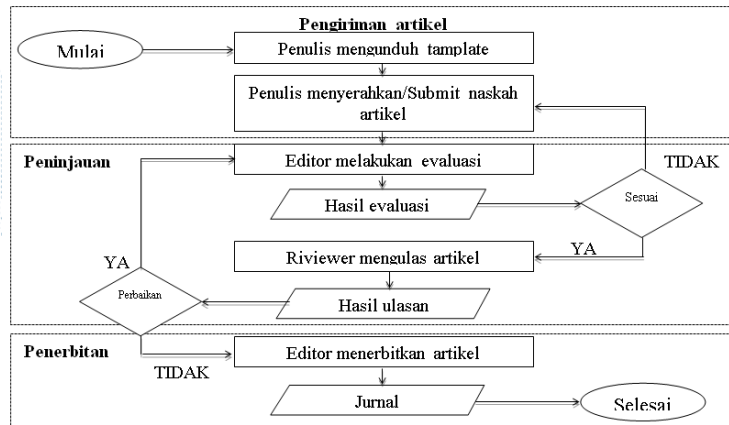
Berdasarkan pada fakta tersebut dianggap perlu adanya analisis mendalam terhadap artikel-artikel yang telah terbit pada jurnal Biramsamtani, baik dari segi bulan terbit, kategori, posisi artikel, dan jumlah kata pada abstrak dan juga klasifikasi jumlah pembaca pada setiap artikel dengan klasifikasi sedikit, sedang, dan banyak, menggunakan metode dan teknik *data mining*. Sehingga penerbitan setiap artikel atau jurnal dapat memperoleh pembaca yang optimal dan merata.

Teknik *data mining* merupakan salah satu cara dalam menganalisis data yang besar yang sering digunakan oleh pada data analisis dalam menggali informasi atau pola yang terdapat dalam data yang besar. Tahapan dalam teknik *data mining* melalui beberapa tahapan dari mulai pembersihan dan integrasi data, seleksi dan transformasi data, *data mining*, dan evaluasi.

Salah satu metode yang populer dan mudah digunakan dalam teknik *data mining* adalah *Bayesian Classification*. Penulis menggunakan metode algoritma *Bayesian Classification* untuk klasifikasi, karena algoritma ini sudah terbukti akurat, seperti pada penelitian (Sutoyo dan Almaarif, 2020) pada prediksi kelulusan mahasiswa dengan akurasi 73.72%. Kemudian pada penelitian (Liliana dkk, 2021) pada prediksi status pasien Covid-19 dengan akurasi mencapai 96.67%. pada penelitian terbaru bahwa penggunaan metode *Bayesian Classification* atau juga disebut *Naïve Bayes Classifier* (NBC) masih sangat populer digunakan oleh para peneliti, seperti pada penelitian (Macfud, dkk., 2023) yang memprediksi klasifikasi tingkat minat barang menunjukkan tingkat akurasi yang tinggi yaitu sebesar 82,76%. Begitu juga pada penelitian (Nurhidayati, dkk., 2023) dalam prediksi klasifikasi penerima beasiswa menghasilkan akurasi yang sangat tinggi yaitu 91.43% dengan tingkat diagnosa mencapai *excellent classification*.

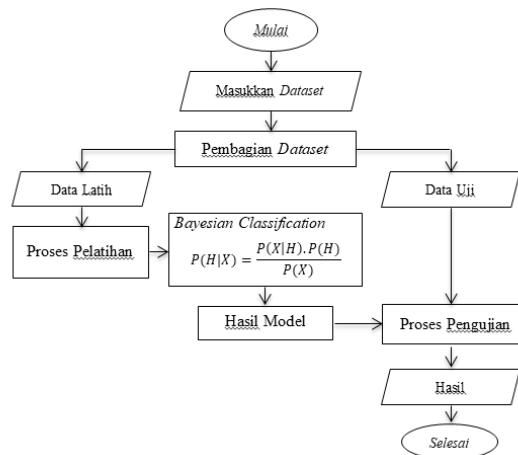
## 2. Metodologi Penelitian

Berikut adalah alur penerbitan artikel di jurnal Biram Samtani:



**Gambar 1.** Alur Penerbitan Artikel

Selanjutnya adalah proses pemodelan dengan metode Bayesian Classification dengan tahapan seperti Gambar berikut:



**Gambar 2.** Tahapan flowchart metode Bayesian Classification

Adapun objek penelitian berdasarkan data historis jumlah *viewers* diambil dari web jurnal Biram Samtani (<https://jurnal.ugp.ac.id/index.php/jbss>). Berdasarkan objek penelitian terdapat 76 jurnal yang sudah terbit sejak tahun 2017 sampai tahun 2023 dengan jumlah *viewer* yang berbeda-beda. Data dikumpulkan berdasarkan tahun, bulan, posisi jurnal pada setiap terbitan, kategpri, jumlah kata pada abstrak, dan jumlah *viewer* yang nantinya akan diolah sesuai kebutuhan metode *Bayesian Classification*.

### 3. Hasil dan Pembahasan

#### 3.1. Hasil

Adapun skenario yang dilakukan untuk mendapatkan hasil klasifikasi dengan metode Bayesian Classification menggunakan tools Rstudio sebagai berikut:

##### a) Pemanggilan *dataset*

pemanggilan *dataset* dengan format (\*.xlsx) menggunakan *library (readxl)*, kemudian menyimpan datanya ke variabel data (`data=DATASET_BIRAMSAMTANI`), dan menampilkan ke bentuk Tabel (`head(data)`).



|   | BULAN    | POSISI_ARTIKEL | KATEGORI   | ABSTRAK | KLASS_VIEWER |
|---|----------|----------------|------------|---------|--------------|
| 1 | NOPEMBER | Awal           | PERTANIAN  | Sedikit | Tinggi       |
| 2 | NOPEMBER | Awal           | EKONOMI    | Sedikit | Rendah       |
| 3 | NOPEMBER | Awal           | AKUNTANSI  | Sedang  | Tinggi       |
| 4 | NOPEMBER | Tengah         | KOMUNIKASI | Sedikit | Menengah     |
| 5 | NOPEMBER | Tengah         | MANAJEMEN  | Sedikit | Rendah       |

**Gambar 3.** Hasil Pemanggilan *dataset*

#### b) Pembagian *dataset*

Dataset Biramsamtani dibagi dalam tiga model yaitu model 1 dengan komposisi data latih dan data uji yaitu 90:10, model 2 yaitu 80:20, dan model 3 yaitu 70:30. Dengan menggunakan kode berikut:

```
#model 1
set.seed(1234)
ind <- sample(2, nrow(data), replace = T, prob = c(0.9, 0.1))
train <- data[ind == 1,]
test <- data[ind == 2,]

#model 2
set.seed(1234)
ind <- sample(2, nrow(data), replace = T, prob = c(0.8, 0.2))
train <- data[ind == 1,]
test <- data[ind == 2,]

#model 3
set.seed(1234)
ind <- sample(2, nrow(data), replace = T, prob = c(0.7, 0.3))
train <- data[ind == 1,]
test <- data[ind == 2,]
```

**Gambar 4.** Pembagian *dataset*

Dari proses pembagian dataset diatas maka dihasilkan data pada model 1 yaitu 73 data latih dan 3 data uji, model 2 yaitu 65 data latih dan 11 data uji, dan pada model 3 yaitu 59 data latih dan 17 data uji.

#### c) Pengklasifikasian dengan *Bayesian classification*

Setelah melakukan pembagian dataset ke dalam tiga model, maka akan dilakukan pengklasifikasian dengan kode berikut:

```
model <- naive_bayes(KLASS_VIEWER ~ ., data = train, usekernel =
T)model
```

**Gambar 5.** Klasifikasi model

Maka probabilitas untuk model 1 diperoleh:

```
-----
A priori probabilities:
Menengah    Rendah    Tinggi
0.1917808 0.5068493 0.3013699
-----
```

**Gambar 6.** Probabilitas model 1

Berdasarkan pada Gambar 6 dapat dijelaskan bahwa probabilitas tertinggi untuk klass view adalah view rendah dengan nilai 0,506 atau 50,6%, view menengah dengan nilai 0,191 atau 19,1%, dan view tinggi dengan nilai 0,301 atau 30,1%. Adapun untuk probabilitas untuk model 2 diperoleh:

A priori probabilities:

| Menengah  | Rendah    | Tinggi    |
|-----------|-----------|-----------|
| 0.2153846 | 0.4769231 | 0.3076923 |

**Gambar 7.** Probabilitas model 2

Berdasarkan pada Gambar 7 dapat dijelaskan bahwa probabilitas tertinggi untuk kelas view adalah view rendah dengan nilai 0,476 atau 47,6%, view menengah dengan nilai 0,215 atau 21,5%, dan view tinggi dengan nilai 0,307 atau 30,7%. Adapun untuk probabilitas untuk model 3 diperoleh:

A priori probabilities:

| Menengah  | Rendah    | Tinggi    |
|-----------|-----------|-----------|
| 0.2372881 | 0.4745763 | 0.2881356 |

**Gambar 8.** Probabilitas model 3

Berdasarkan pada Gambar 8 dapat dijelaskan bahwa probabilitas tertinggi untuk kelas view adalah view rendah dengan nilai 0,474 atau 47,4%, view menengah dengan nilai 0,237 atau 23,7%, dan view tinggi dengan nilai 0,288 atau 28,8%.

#### d) Perhitungan kinerja dan akurasi model

Setelah menghitung probabilitas dari level klas view pada setiap model, maka akan diuji tingkat akurasinya terhadap data latih dengan kode sebagai berikut:

```
#PREDIKSI
p <- predict(model, train, type = 'prob')
head(cbind(p, train))

#CONFUSION MATRIX DATA LATIH
p1 <- predict(model, train)
(tab1 <- table(p1, train$KLASS_VIEWER))
p1
1 - sum(diag(tab1)) / sum(tab1)
```

**Gambar 9.** Akurasi model

Maka akurasi yang didapat untuk model 1 adalah:

```
p1      Menengah Rendah Tinggi
Menengah      7      4      0
Rendah       6     30      4
Tinggi       1      3     18
> p1
[1] Tinggi Menengah Tinggi Rendah Rendah Menengah Rendah Rendah Rendah Rendah Rendah
[12] Tinggi Rendah Rendah Menengah Rendah Rendah Rendah Rendah Rendah Rendah Rendah
[23] Rendah Rendah Rendah Rendah Rendah Rendah Rendah Tinggi Rendah Rendah Rendah
[34] Menengah Rendah Rendah Tinggi Rendah Menengah Rendah Menengah Menengah Rendah Menengah
[45] Tinggi Rendah Rendah Menengah Rendah Rendah Rendah Rendah Rendah Rendah Rendah
[56] Tinggi Tinggi Tinggi Tinggi Tinggi Rendah Rendah Tinggi Rendah Tinggi Tinggi
[67] Tinggi Tinggi Rendah Tinggi Tinggi Rendah Tinggi
Levels: Menengah Rendah Tinggi
> 1 - sum(diag(tab1)) / sum(tab1)
[1] 0.2465753
> |
```

**Gambar 10.** Akurasi model 1

Berdasarkan pada Gambar 10 dapat dijelaskan bahwa prediksi level klas view menengah benar yaitu 7, salah 4. Untuk prediksi level klas view rendah benar yaitu 30,

salah 10, Dan prediksi level klas view tinggi benar yaitu 18, salah 4. Sehingga tingkat kesalahannya adalah 0,246 atau 24,6%.

```
p1
Menengah Rendah Tinggi
Menengah 7 1 0
Rendah 6 27 4
Tinggi 1 3 16
> p1
[1] Tinggi Menengah Tinggi Rendah Menengah Rendah Rendah Tinggi Menengah Menengah Tinggi
[12] Rendah Rendah Rendah Rendah Rendah Rendah Rendah Rendah Rendah Rendah Rendah
[23] Rendah Rendah Tinggi Rendah Rendah Rendah Rendah Rendah Rendah Rendah Rendah
[34] Rendah Rendah Menengah Menengah Rendah Menengah Rendah Rendah Rendah Menengah Rendah
[45] Rendah Rendah Rendah Rendah Rendah Tinggi Tinggi Tinggi Tinggi Tinggi Tinggi
[56] Tinggi Rendah Tinggi Tinggi Tinggi Tinggi Tinggi Tinggi Tinggi Tinggi
Levels: Menengah Rendah Tinggi
> 1 - sum(diag(tab1)) / sum(tab1)
[1] 0.2307692
```

**Gambar 11.** Akurasi model 2

Berdasarkan pada Gambar 11 dapat dijelaskan bahwa prediksi level klas view menengah benar yaitu 7, salah 1. Untuk prediksi level klas view rendah benar yaitu 27, salah 10, Dan prediksi level klas view tinggi benar yaitu 16, salah 4. Sehingga tingkat kesalahannya adalah 0,230 atau 23%.

```
p1
Menengah Rendah Tinggi
Menengah 9 2 2
Rendah 4 23 2
Tinggi 1 3 13
> p1
[1] Tinggi Menengah Tinggi Rendah Menengah Rendah Rendah Tinggi Menengah Menengah Tinggi
[12] Rendah Rendah Menengah Menengah Rendah Rendah Rendah Rendah Rendah Rendah Rendah
[23] Rendah Rendah Tinggi Rendah Rendah Rendah Rendah Rendah Rendah Rendah Rendah
[34] Rendah Menengah Menengah Rendah Menengah Tinggi Menengah Menengah Menengah Rendah Rendah
[45] Rendah Rendah Tinggi Tinggi Tinggi Tinggi Menengah Rendah Tinggi Tinggi Tinggi
[56] Tinggi Tinggi Tinggi Tinggi
Levels: Menengah Rendah Tinggi
> 1 - sum(diag(tab1)) / sum(tab1)
[1] 0.2372881
```

**Gambar 12.** Akurasi model 3

Berdasarkan pada Gambar 12 dapat dijelaskan bahwa prediksi level klas view menengah benar yaitu 9, salah 4. Untuk prediksi level klas view rendah benar yaitu 23, salah 6, Dan prediksi level klas view tinggi benar yaitu 13, salah 4. Sehingga tingkat kesalahannya adalah 0,237 atau 23,7%.

### 3.2. Pembahasan

Dari hasil diatas dapat kita bahas perbandingan hasil akurasi model terhadap data latih dan data uji menggunakan *confusion matrix* untuk model 1, model 2, dan model 3 sebagai berikut:

**Tabel 1.** *Confusion matrix* data latih model 1

| Model 1      |          | Nilai Prediksi |        |        |
|--------------|----------|----------------|--------|--------|
| Nilai Aktual | Menengah | Menengah       | Rendah | Tinggi |
|              | Rendah   | 6              | 30     | 4      |
|              | Tinggi   | 1              | 3      | 18     |

Berdasarkan perhitungan, maka dapat disimpulkan bahwa akurasi untuk model 1 adalah 73,34 %, presisi 73,48 %, dan recall 70,96 %.

**Tabel 2.** *Confusion matrix* data latih model 2

| Model 2      |          | Nilai Prediksi |        |        |
|--------------|----------|----------------|--------|--------|
| Nilai Aktual | Menengah | Menengah       | Rendah | Tinggi |
|              | Rendah   | 6              | 27     | 4      |
|              | Tinggi   | 1              | 3      | 16     |

Berdasarkan perhitungan diatas maka dapat disimpulkan untuk bahwa akurasi untuk model 2 adalah 76,92%, presisi 80,15%, dan recall 72,36%.



**Tabel 3.** *Confusion matrix* data latih model 3

| Model 3      |          | Nilai Prediksi |        |        |
|--------------|----------|----------------|--------|--------|
| Nilai Aktual | Menengah | Menengah       | Rendah | Tinggi |
|              | Rendah   | 9              | 2      | 2      |
|              | Tinggi   | 4              | 23     | 2      |
|              |          | 1              | 3      | 13     |

Berdasarkan perhitungan diatas maka dapat disimpulkan untuk bahwa akurasi untuk model 3 adalah 76,27%, presisi 75%, dan recall 74,29%.

**Tabel 4.** *Confusion matrix* data latih model 4

| No. | Model   | Akurasi       | Presisi       | Recall        |
|-----|---------|---------------|---------------|---------------|
| 1.  | Model 1 | 73,34%        | 73,48%        | 70,96%        |
| 2.  | Model 2 | <b>76,92%</b> | <b>80,15%</b> | 72,36%        |
| 3.  | Model 3 | 76,27%        | 75%           | <b>74,29%</b> |

Berdasarkan pada Tabel 4. maka dapat disimpulkan bahwa akurasi dan presisi tertinggi diperoleh oleh model 2 dengan nilai 76,92% dan 80,15%, sedangkan untuk recall tertinggi diperoleh oleh model 3 dengan nilai 74,29%. Adapun data uji dalam penelitian ini adalah 20% dari dataset berdasarkan model terpilih sebagai berikut:

**Tabel 5.** Data uji

| No. | Bulan     | Posisi Artikel | Kategori   | Abstrak | Klas Viewer |
|-----|-----------|----------------|------------|---------|-------------|
| 1   | Nopember  | Tengah         | Manajemen  | Banyak  | Rendah      |
| 2   | Juli      | Akhir          | Pertanian  | Sedang  | Tinggi      |
| 3   | Januari   | Awal           | Politik    | Sedang  | Rendah      |
| 4   | September | Awal           | Teknologi  | Sedang  | Rendah      |
| 5   | September | Awal           | Manajemen  | Sedang  | Rendah      |
| 6   | September | Tengah         | Manajemen  | Sedang  | Rendah      |
| 7   | Juli      | Tengah         | Teknologi  | Sedikit | Rendah      |
| 8   | Juli      | Tengah         | Pendidikan | Sedikit | Tinggi      |
| 9   | Februari  | Awal           | Manajemen  | Sedang  | Tinggi      |
| 10  | Februari  | Tengah         | Pendidikan | Sedang  | Rendah      |
| 11  | Februari  | Awal           | Ekonomi    | Sedang  | Rendah      |

Perhitungan dilakukan sampai data uji ke-11, dengan mengulangi langkah-langkah perhitungan diatas, sehingga hasil perhitungan dapat dilihat pada Tabel berikut:

**Tabel 6.** Hasil pengujian data uji

| Data ke- | Probabilitas |             |             | Klas          |                 |
|----------|--------------|-------------|-------------|---------------|-----------------|
|          | Rendah       | Menengah    | Tinggi      | Aktual        | Prediksi        |
| 1.       | 0.000594914  | 0           | 0           | Rendah        | Rendah          |
| 2.       | 0.001735165  | 0.000840996 | 0.001225    | <b>Tinggi</b> | <b>Rendah</b>   |
| 3.       | 0            | 0.000784929 | 0           | <b>Rendah</b> | <b>Menengah</b> |
| 4.       | 0.003966091  | 0           | 0.00105     | Rendah        | Rendah          |
| 5.       | 0.007932183  | 0.007849293 | 0.0021      | Rendah        | Rendah          |
| 6.       | 0.004759311  | 0.002242654 | 0.00245     | Rendah        | Rendah          |
| 7.       | 0.000607308  | 0           | 9.42308E-05 | Rendah        | Rendah          |
| 8.       | 0.000151827  | 0           | 0.000282692 | <b>Tinggi</b> | <b>Tinggi</b>   |
| 9.       | 0.00132203   | 0           | 0.0021      | <b>Tinggi</b> | <b>Tinggi</b>   |
| 10.      | 9.91523E-05  | 0           | 0.003675    | <b>Rendah</b> | <b>Tinggi</b>   |
| 11.      | 0.00132203   | 0           | 0.0021      | <b>Rendah</b> | <b>Tinggi</b>   |

Dari Tabel 6 dapat dijelaskan bawa dari 11 data yang diklasifikasi benar sebanyak 7 data, dan yang salah klasifikasi sebanyak 4, sehingga akurasi untuk data uji adalah 63,63%. Berdasarkan pada klasifikasi data latih dan data uji dan juga data actual dapat diperoleh

susunan yang baik dalam menerbitkan artikel pada jurnal Biramsamtani sehingga tujuan jumlah pembaca yang tinggi adalah:

- a) Menghindari jumlah kata pada abstrak yang banyak;
- b) Untuk abstrak yang sedang hanya dapat diterapkan pada kategori Akuntansi, Manajemen, Pendidikan, Pertanian, Peternakan dan Teknologi;
- c) Untuk abstrak yang sedikit berlaku pada kategori Pertanian, Agama, Peternakan, Ekonomi, dan Pendidikan.

#### 4. Kesimpulan

Kesimpulan yang dapat diambil dari analisis ini menunjukkan bahwa model klasifikasi Bayesian Classification untuk prediksi jumlah pembaca yang paling efektif adalah model 2, yang menggunakan komposisi data latih dan data uji 80:20, dengan tingkat akurasi mencapai 79,92%, presisi 80,15%, dan recall 72,36%. Selanjutnya, klasifikasi jumlah pembaca pada artikel jurnal Biramsamtani berdasarkan nilai probabilitasnya menunjukkan bahwa 47,7% memiliki potensi pembaca rendah, 30,76% tinggi, dan 21,53% menengah. Rekomendasi untuk menyusun variabel artikel agar dapat meningkatkan jumlah pembaca adalah dengan menghindari penggunaan jumlah kata yang berlebihan pada abstrak. Untuk abstrak dengan jumlah kata sedang, saran ini berlaku untuk kategori Akuntansi, Manajemen, Pendidikan, Pertanian, Peternakan, dan Teknologi, sementara untuk abstrak dengan jumlah kata sedikit, hal ini dapat diterapkan pada kategori Pertanian, Agama, Peternakan, Ekonomi, dan Pendidikan.

#### Daftar Pustaka

- [1] Faisal, M. R., & Nugrahadi, D. T. (2019). *Belajar DATA SCIENCE Klasifikasi dengan Bahasa Pemrograman R*. Banjarbaru: Scripta Cendekia.
- [2] Febriyanti dan Februariyanti (2023) "Analisis Sentimen Terhadap Program Kampus Merdeka Menggunakan Algoritma *Naive Bayes Classifier* Di Twitter" Jurnal TEKNO KOMPAK, Vol. 17, No. 1, Hal: 25-38
- [3] Firmansyah dan Yulianto (2023). "Prediksi Hasil Belajar Menggunakan *Naive Bayes Classifier* pada Tingkat Sekolah Dasar". Remik: Riset dan E- urnal Manajemen Informatika Komputer Volume 7, Nomor 2, April 2023, Hal: 1174-1182
- [4] Liliana dkk, 2021. "Data Mining untuk Prediksi Status Pasien Covid-19 dengan Pengklasifikasi *Naive Bayes*", Jurnal Multinetics Vol. 7 No. 1 Mei 2021
- [5] Marisa, Fitri., dkk (2021). *Data Mining Konsep dan Penerapannya* Deepublish, Yogyakarta.
- [6] Macfud, dkk (2023) "Analisis Algoritma *Naive Bayes Classifier* (NBC) Pada Klasifikasi Tingkat Minat Barang Di Toko Violet Cell". Jurnal Mahasiswa Teknik Informatika Vol. 7 No. 1. Hal: 87-94
- [7] Nurhidayati, dkk (2023) "Implementasi Algoritma *Naive Bayes* Untuk Klasifikasi Penerima Beasiswa (Studi Kasus Universitas Hamzanwadi)". Jurnal Informatika dan Teknologi Vol. 6 No. 1, Hal: 177-188
- [8] Pebdika, dkk (2023) "Klasifikasi Menggunakan Metode *Naive Bayes* Untuk Menentukan Calon Penerima PIP". Jurnal Mahasiswa Teknik Informatika) Vol. 7 No. 1, Hal:452-258
- [9] Raza, Rohit., dkk, (2022). *Data mining and machine learning applications*. Wiley & Scrivener Publishing, New Jersey U.S
- [10] Roza, dkk., (2020). *Tutorial Sistem Informasi Prediksi Jumlah Pelanggan Menggunakan Metode Regresi Linier Berganda Berbasis Web Menggunakan Framework Codeigniter*. Kreatif, Indonesia
- [11] Sutoyo, E., & Almaarif, A. (2020). Educational Data Mining for Predicting Student Graduation Using the *Naive Bayes Classifier* Algorithm". *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 4(1), 95 - 101



- [12] Wardhono, Adhitya dkk. *Analisis Data Time Series Dalam Model Makroekonomi*. Jawa Timur: CV. Pustaka Abadi. 2019.
- [13] Admin. Daftar Jurnal Indonesia Terindeks DOAJ. [online] (penerbitdeepublish.com). diakses 6 Mei 2023.