

KOMPARASI METODE KLASIFIKASI PADA ANALISIS SENTIMEN USAHA WARALABA BERDASARKAN DATA TWITTER

Tati Mardiana¹; Hafiz Syahreva²; Tuslaela³

¹Sistem Informasi
Universitas Bina Sarana Informatika
www.bsi.ac.id
tati.ttm@bsi.ac.id

^{2,3}Sistem Informasi
STMIK Nusa Mandiri
www.nusamandiri.ac.id
hafizsyahreva22@gmail.com; tuslaela.tll@nusamandiri.ac.id

Abstract—At present, the franchise business in Indonesia has a relatively high attractiveness. However, many business actors also failed. For someone who wants to start a business needs to consider the public sentiment towards the franchise business. Although, it is not easy to do with sentiment analysis because of the large number of conversations on Twitter about franchising and unstructured data. The purpose of this study is to compare the accuracy of Neural Network, K-Nearest Neighbor, Naïve Bayes, Support Vector Machine, and Decision Tree methods in extracting attributes in documents or text containing comments to find out the expressions there and to classify them into positive and negative comments. This research uses real-time data from tweets on Twitter. Next process the data by first cleaning it of noise using Python. The test results obtained by the confusion matrix obtained Neural Network accuracy value of 83%, K-Nearest Neighbor by 52%, Support Vector Machine by 83%, and Decision Tree by 81%. This study shows that the support vector machine and Neural Network methods have the best accuracy for classifying positive and negative comments related to franchising.

Keywords : Franchise, Sentiment, Python, Twitter, Comparison.

Intisari— Saat ini usaha waralaba di Indonesia memiliki daya tarik yang relatif tinggi. Namun, para pelaku usaha banyak juga yang mengalami kegagalan. Bagi seseorang yang ingin memulai usaha perlu mempertimbangkan sentimen masyarakat terhadap usaha waralaba. Meskipun demikian, tidak mudah untuk melakukan analisis sentimen karena banyaknya jumlah percakapan di Twitter terkait usaha waralaba dan tidak terstruktur. Tujuan penelitian ini adalah melakukan komparasi akurasi metode *Neural Network*, *K-Nearest Neighbor*, *Naïve Bayes*, *Support Vector Machine*, dan *Decision Tree* dalam

mengeksktraksi atribut pada dokumen atau teks yang berisi komentar untuk mengetahui ekspresi didalamnya dan mengklasifikasikan menjadi komentar positif dan negatif. Penelitian ini menggunakan data *realtime* dari *tweets* pada Twitter. Selanjutnya mengolah data tersebut dengan terlebih dulu membersihkannya dari *noise* dengan menggunakan *Phyton*. Hasil pengujian dengan *confusion matrix* diperoleh nilai akurasi *Neural Network* sebesar 83%, *K-Nearest Neighbor* sebesar 52%, *Support Vector Machine* sebesar 83%, dan *Decision Tree* sebesar 81%. Penelitian ini menunjukkan metode *Support Vector Machine* dan *Neural Network* paling baik untuk mengklasifikasikan komentar positif dan negatif terkait usaha waralaba.

Kata Kunci: Waralaba, Sentimen, *Phyton*, Twitter, Komparasi.

PENDAHULUAN

Saat ini usaha waralaba di Indonesia memiliki daya tarik yang relatif tinggi. Berdasarkan data dari Kementerian Perdagangan Republik Indonesia, tercatat ada 698 waralaba pada tahun 2016. Dengan jumlah gerai sebanyak 24.400 yang terdiri dari 37% waralaba mancanegara dan 63% waralaba bisnis online lokal (BO), Omzet yang dicapai hingga 172 triliun. Dari data yang dipaparkan, industri waralaba di Indonesia mengalami pertumbuhan sebesar 37% dan menjadi salah satu pasar yang paling prospektif untuk bisnis ini (Fathurahman, Windarti, & Purwanto, 2018). Namun demikian, para pelaku usaha banyak juga yang mengalami kegagalan. Persentase kegagalannya berkisar pada rentang 50-60% per tahun untuk kategori bisnis lokal dan kategori bisnis asing menyentuh angka 12-13% per tahun (Imanuwelita, Putri, & Amalia, 2018). Bagi seseorang yang ingin memulai usaha perlu

mempertimbangkan popularitas usaha waralaba berdasarkan opini masyarakat.

Pesatnya perkembangan internet telah mempengaruhi komunikasi masyarakat modern. Masyarakat modern cenderung memiliki kebebasan dalam memberikan opini melalui jejaring sosial. Saat ini jejaring sosial Twitter cukup populer di kalangan pengguna internet. Pengguna Twitter pada kuartal I di tahun 2017 mencapai sekitar 328 juta bertambah 6% atau 9 juta pengguna aktif (Attabi, Muflikhah, & Fauzi, 2018). Setiap pengguna Twitter bebas memberikan opini melalui komentar atau disebut dengan *tweets* dengan batasan 140 karakter. Opini melalui tweet inilah yang dapat dimanfaatkan untuk melihat bagaimana sentimen masyarakat terhadap usaha waralaba.

Seorang analis dapat menentukan sentimen masyarakat terhadap usaha waralaba secara manual. Akan tetapi, seiring bertambahnya opini pada twitter dan datanya bersifat tidak terstruktur menjadi semakin banyak waktu dan usaha yang dibutuhkan untuk mengklasifikasikan polaritas opini tersebut (Nurhuda, Sihwi, & Doewes, 2013). Oleh karena itu, dibutuhkan metode untuk mengekstraksi atribut pada dokumen atau teks yang berisi komentar untuk mengetahui ekspresi didalamnya dan mengklasifikasikan menjadi polaritas positif dan negatif.

Analisis sentimen merupakan bidang interdisipliner yang terdiri dari pemrosesan bahasa alami, analisis teks, dan linguistik komputasi untuk mengidentifikasi sentimen teks (Vinodhini & Chandrasekaran, 2016) berdasarkan opini-opini, sentimen, serta emosi yang diekspresikan dalam teks (Ling, Kencana, & Oka, 2014). Salah satu teknik yang dapat melakukan klasifikasi secara cepat adalah dengan text mining. Text mining merupakan bagian dari bidang data mining yang bertujuan untuk menggali dan menemukan informasi-informasi, pola, dan tren yang tersembunyi dari jumlah data yang besar (Widaningsih & Suheri, 2018).

Ada beberapa penelitian sebelumnya tentang analisis sentimen menggunakan berbagai metode dalam mesin pembelajaran untuk melakukan klasifikasi. (Ling et al., 2014) mengusulkan penyeleksian fitur menggunakan *chi square* dalam proses memilih subset dari fitur-fitur yang relevan untuk digunakan dalam konstruksi model probabilistik NBC. Hasil penelitian menunjukkan bahwa frekuensi kemunculan fitur yang diharapkan dalam kategori benar dan dalam kategori salah memiliki peran penting dalam pemilihan fitur *chi-square*. Kemudian klasifikasi oleh *Naïve Bayes classifier* memperoleh akurasi 83% dan rata-rata harmonis 90,713%.

(Nurhuda et al., 2013) mengusulkan analisis sentimen masyarakat terhadap calon presiden 2014 berdasarkan opini dari Twitter menggunakan metode *Naïve Bayes*. Hasil dari penelitian menunjukkan pasangan capres dan cawapres Prabowo Subianto dan Hatta Rajasa mendapatkan jumlah percakapan sebesar 53% dengan 47,7% untuk sentimen positif, 26,4% sentimen negatif dan 25,9% sentimen netral. Sedangkan pasangan Joko Widodo – Jusuf Kalla mendapatkan jumlah percakapan sebesar 47% dengan 37,6% sentimen positif, 34,4% sentimen negatif, dan 27,9 sentimen netral.

(Muthia, 2017) mengusulkan analisis sentimen pada *review* restoran dengan mengintegrasikan metode *Naïve Bayes* dengan Genetic Algorithm. Hasil pengujian kedua metode ini, menunjukkan peningkatan akurasi metode *Naïve Bayes* dengan pemilihan fitur *Generic Information* dari 86.50% menjadi 90.50%. Model yang terbentuk menghasilkan bentuk *review* positif dan negatif yang membantu seorang untuk menghemat waktu saat mencari *review* suatu restoran.

Penelitian berikutnya yang dilakukan oleh (Attabi et al., 2018) mengusulkan analisis sentimen untuk penilaian produk kecantikan dan perawatan kulit Mustika Ratu menggunakan *Naïve Bayes Classifier*. Metode *Naïve Bayes Classifier* memiliki kemudahan dalam implementasinya, dan memiliki performa cepat dalam proses pembelajaran. Namun, dalam proses klasifikasi teks terdapat permasalahan dimensi tinggi dari ruang fitur. Oleh karena itu, penambahan *Information Gain* diperlukan untuk proses seleksi fitur dengan mengurangi keberadaan kata yang tidak relevan pada data yang digunakan. Hasil penelitian menunjukkan peningkatan akurasi *Naïve Bayes* dengan metode *Information Gain* dari 70% menjadi 74%.

Penelitian lainnya, (Romadloni, Santoso, & Budilaksono, 2019) yang mengusulkan analisis sentimen transportasi umum KRL *Commuter Line* Jabodetabek menggunakan metode *Naïve Bayes Classifier*, KNN dan *Decision Tree*. Hasil dari pengujian yang dihasilkan menunjukkan bahwa nilai akurasi terhadap masing-masing metode adalah metode *Naïve Bayes Classifier* dan KNN sebesar 80%, sedangkan nilai akurasi pada metode *Decision Tree* sebesar 100%.

Berdasarkan penelitian sebelumnya maka untuk mengetahui sentimen masyarakat terhadap usaha waralaba dapat menggunakan klasifikasi teks. Oleh karena itu, penelitian ini akan melakukan komparasi akurasi metode *Neural Network*, *K-Nearest Neighbors*, *Naïve Bayes*, *Support Vector Machine*, dan *Decision Tree* dalam mengekstraksi atribut pada dokumen atau teks

yang berisi komentar untuk mengetahui ekspresi didalamnya dan mengklasifikasikan menjadi komentar positif dan negatif.

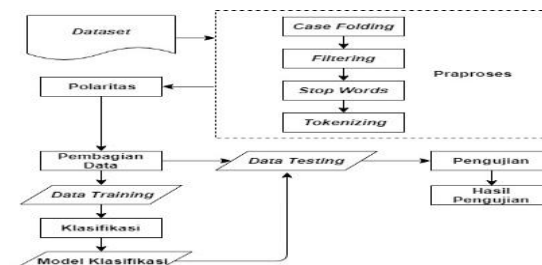
BAHAN DAN METODE

Objek penelitian ini adalah usaha waralaba yang saat ini banyak diminati masyarakat yang akan memulai usaha tetapi tidak memiliki pengalaman. Penelitian ini menggunakan teknik analisis kualitatif deskriptif, yaitu penelitian yang jenis datanya kualitatif berupa kumpulan fakta atau data pada suatu latar alamiah. Latar alamiah yang dimaksud adalah tuturan kalimat yang mengandung kata *franchise* atau waralaba yang dijadikan sebagai sumber data langsung. Mengingat data yang diperoleh merupakan opini yang terdapat pada Twitter berupa komentar positif dan negatif. Maka dari itu, unsur sentimen inilah yang dijadikan sebagai pengolahan data yang hasil akhirnya nanti dapat digunakan sebagai acuan dasar untuk membangun usaha. Rancangan pada penelitian ini dapat dilihat pada gambar 1.

Pada penelitian ini data yang diperoleh merupakan data yang mengenai opini pengguna media sosial Twitter terhadap jenis usaha waralaba atau *franchise*. Ada beberapa cara dalam pengambilan data yang banyak sekaligus dari Twitter. Yang pertama dengan menghubungkan Twitter API. Twitter API tersebut bisa digunakan pada *Software Aplikasi* yang menyediakan pengolahan *Text Mining* pada pemrosesannya. Dan cara yang kedua dengan langsung menggunakan modul pada bahasa pemrograman Python. Data yang diperoleh merupakan data *realtime* yang diambil berdasarkan komentar di Twitter menggunakan modul "twitterscraper" pada Bahasa pemrograman *Python*. Modul pada Bahasa pemrograman *Python* yang tersedia, adalah :

1. Pandas, modul ini digunakan untuk membaca *file* ke dalam *Data Frame* dan diolah dengan menggunakan Python.
2. Numpy, banyak yang memanfaatkan modul ini untuk menghitung dan membuat suatu program namun pada penelitian modul numpy berguna untuk perhitungan metode.
3. Re dan String, modul ini digunakan untuk mengolah kalimat yang nantinya akan dihitung pada tahap selanjutnya.
4. NLTK (*Natural Language Tool Kit*), salah satu modul yang harus disediakan adalah modul ini, modul ini digunakan untuk memproses dan membersihkan data dari *noise* yang mengganggu inti dari kalimat seperti *punctuation* dan *stopwords*,
5. Matplotlib dan Sklearn, modul ini digunakan untuk pengklasifikasian data dan memvisualisasikan data yang telah diolah.

6. Textblob, modul ini digunakan untuk mengklasifikasikan kalimat pada *Data Frame*, pada penelitian ini, modul tersebut digunakan untuk mencari polaritas yang terdapat pada kalimat tersebut.



Sumber : (Mardiana, Syahreva, & Tuslaela, 2019)

Gambar 1. Rancangan Penelitian

Tahapan yang dilakukan pada penelitian ini meliputi beberapa proses, yaitu:

1. Pemahaman Data

Pada tahap ini ada beberapa hal yang dilakukan untuk memahami data yang akan diolah. Data yang diambil merupakan data yang secara langsung diambil dari Twitter yang mengandung kata waralaba atau *franchise*. Setelah data terkumpul, data yang dipilih hanya data berupa teks tanpa membawa nama *user*, tanggal *posting* dsb.

2. Persiapan Data

Ada beberapa proses yang dilakukan pada tahap persiapan data antara lain, mengubah jenis huruf menjadi huruf kecil (*case folding*), *filtering*, menghilangkan *stopwords*, *Tokenizing* dan mengidentifikasi polaritas sentimen dari setiap kalimat *tweets* dengan memilah polaritasnya, (-1) untuk kalimat negatif, (0) untuk kalimat netral dan (1) untuk kalimat positif. Data yang diperoleh dari hasil ini sebanyak 2615 data karena telah melalui tahap *filtering* yang menghilangkan kata yang bersifat duplikasi. Setelah data ditentukan polaritasnya kalimat yang berpolaritas netral akan dihilangkan, sedangkan kalimat negatif akan diganti dengan (0). Proses penukaran polaritas ini dilakukan dengan Ms. Excel atau *software* pengolahan kata lainnya, proses ini bertujuan untuk memfokuskan data menjadi dua polaritas saja pada saat pengklasifikasian pemodelan.

3. Pemodelan

Pada tahap ini, dilakukan pembagian data terlebih dahulu sebelum diproses ditahap selanjutnya. Data dibagi menjadi dua yaitu data *training* dan data *testing*. Data tersebut dibagi sebesar 20% untuk *data testing* dan sisanya dijadikan data *training*. Metode klasifikasi yang digunakan adalah *Multinomial Naïve Bayes* untuk memprediksi peluang mengenai topik yang dibicarakan mengenai jenis usaha waralaba atau *franchise* yang didapatkan dari media sosial

Twitter. Selain itu, percobaan menggunakan metode lainnya juga dilakukan untuk melihat keakuratan dari masing-masing metode tersebut sebagai perbandingan. Metode yang digunakan antara lain, *Support Vector Machine* (SVM), *neural network*, *decision tree* dan *K-nearest neighbor* (kNN).

4. Hasil

Setelah diketahui hasil dari prediksi pada proses pengklasifikasian, tahap selanjutnya dilakukan pengujian. Data yang dihasilkan diuji keakuratannya dengan membuat *confusion matrix* dan dihitung tingkat akurasi. Setelah keakuratan telah diketahui, data yang dihasilkan dapat diketahui nilai presisi, *recall* dan *f1-score*nya dan divisualisasikan dengan *Area Under Curve* (AUC) dan *ROC Curve*. Selain itu, tingkat keakuratan terhadap masing-masing kurva dapat diuji kembali untuk melihat tingkat presentasinya.

HASIL DAN PEMBAHASAN

A. Hasil

Salah satu proses yang paling menentukan dalam pengolahan analisa sentimen dalam bentuk teks adalah proses pembentukan data awal. Data yang akan diolah terlebih dulu dibersihkan dari *noise* yang ada agar pengolahan dapat berjalan maksimal.

a. Case Folding

Pertama *file* harus dimasukan ke dalam *data frame*, selanjutnya data pada kolom *text* diubah dan karakter *enter* dihilangkan

```
In [2]: #-----Preprocessing-----#
In [3]: #import data hasil crawl csv
#menghilangkan semua enter
#membuat huruf kecil
data = pd.read_csv('data-twitter-olah.csv')
data = data.replace('\n', '', regex=True)
data.loc[:, 'text'] = data.text.str.lower()
data.head() #----- data mentah
```

	text
0	dunia yg kita cintai akan pergi, sementara kub...
1	sesuaikan anggaran dg lokasi yg dimiliki, perh...
2	jangan pernah biarkan orang lain mengatakan ka...
3	pria kalo sudah dewasa bukan mliki siapapun k...
4	usaha waralaba murah jakarta, usaha waralaba m...

Sumber : (Mardiana et al., 2019)

Gambar 2. Case Folding

b. Filtering

Setelah keseluruhan data dikonversi ke huruf kecil, data diolah pada proses *filtering* ini bertujuan untuk membersihkan data dari karakter atau elemen yang tidak dibutuhkan seperti URL, *stopwords* (ke, di, dan, dari, dsb), *punctuation* dan angka. Berikut kode yang dapat di masukan pada proses ini.

```
In [4]: #menghilangkan punctuation
data.loc[:, 'text'] = data.text.apply(lambda x: ' '.join(re.sub("[@|
```

Sumber : (Mardiana et al., 2019)

Gambar 3. Filtering Punctuation

Setelah *punctuation* dihilangkan, langkah selanjutnya yaitu menghilangkan *stopwords* seperti pada gambar 4.

```
In [5]: #menghilangkan stopwords
stop_words = get_stop_words('id') #stopwords indonesia
def remove_stopwords(s):
    '''For removing stop words'''
    s = ' '.join(word for word in s.split() if word not in stop_words)
    return s
data.loc[:, 'text'] = data.text.apply(lambda x: remove_stopwords(x))
from stop_words import get_stop_words
stop_words = get_stop_words('en') #stopwords english
def remove_stopwords(s):
    '''For removing stop words'''
    s = ' '.join(word for word in s.split() if word not in stop_words)
    return s
data.loc[:, 'text'] = data.text.apply(lambda x: remove_stopwords(x))
```

Sumber : (Mardiana et al., 2019)

Gambar 4. Filtering Stopwords

Dikarenakan banyaknya data yang sama, pada proses ini data tersebut disisakan hanya satu data.

```
In [6]: #menghapus data yang sama
data = data.drop_duplicates(keep=False)
#menampilkan
data.head() #----- melewati praproses
#menyimpan
data.to_csv('sentimen.csv')
```

	text
43	menginjak lantai waralaba dompet perkasa untuk...
56	waralaba kuliner terlaris harga terjangkau mem...
63	download one piece movie neimaki jima dabouken...
82	pendapatan pasif didapat investasi konglomeras...
83	waralaba melanggar didenda dicabut izinnnya

Sumber : (Mardiana et al., 2019)

Gambar 5. Filtering Duplicates

c. Tokenizing

Toeknisasi merupakan proses memisahkan setiap kata yang menyusun suatu kalimat sehingga memudahkan dalam mengidentifikasi polaritas sentimen dari setiap kalimat.

```
In [14]: #tokenisasi dan perubahan data bernomor-label
re_tok = df['text']
def tokenize(s): return re_tok.sub(r'\s+', ' ').split()
vect = CountVecorizer(tokenize)
tf_train = vect.fit_transform(X_train)
tf_test = vect.transform(X_test)
#menampilkan
In [15]: tf_train
Out[15]: <1413x6593 sparse matrix of type '<class 'numpy.int64'>'
with 22414 stored elements in Compressed Sparse Row format>
In [16]: tf_test
Out[16]: <354x6593 sparse matrix of type '<class 'numpy.int64'>'
with 4721 stored elements in Compressed Sparse Row format>
In [17]: #-----Persiapan Data Selesai-----#
```

Sumber : (Mardiana et al., 2019)

Gambar 6. Tokenisasi

d. Polarity

Proses selanjutnya adalah mengidentifikasi opini yang diberikan bersifat positif atau negatif berdasarkan kata-kata yang telah dipisahkan seperti pada gambar 7.

```

In [8]: #-----Polaritas (positif, negatif, netral)-----
In [9]: #mencari polaritas dari masing-masing postingan
sentimen = data[['text']]

def detect_polarity(text):
    analysis = TextBlob(text)
    if analysis.sentiment.polarity > 0:
        return 1
    elif analysis.sentiment.polarity == 0:
        return 0
    else:
        return -1

sentimen['polarity'] = sentimen.text.apply(detect_polarity)
#sentimen.to_csv('sentimen-data-twitter.csv', encoding='utf8')
sentimen.head()

Out[9]:
      text  polarity
43  menginjak lantai waralaba dompet perkasa untuk...      0
56  waralaba kuliner terlaris harga terjangkau mem...      0
63  download one piece movie neimaki jma dabouken...      0
82  pendapatan pasif didapat investasi konglomeras...      0
83  waralaba melanggar didenda dicabut izinnya      0

In [10]: #-----Polaritas (positif, negatif, netral) selesai-----

```

Sumber : (Mardiana et al., 2019)

Gambar 7. Polarity Process

Hasil pengolahan data diperoleh sebanyak 1767 opini yang terdiri dari 1265 opini bersifat positif dan 502 opini bersifat negatif. Data *training* yang digunakan dalam pengklasifikasian sebanyak 1414 opini dan 353 opini digunakan sebagai data *testing*. Hasil pengujian dari ke lima model klasifikasi pada analisis sentimen usaha waralaba sebagai berikut :

1. Metode Naïve Bayes

Hasil pengujian terhadap opini usaha waralaba dengan metode *Naïve Bayes* disajikan pada gambar 8. Nilai akurasi diperoleh sebesar 0.799 atau 80%. Ketepatan antara informasi yang diminta dengan prediksi yang diberikan (*class precision*) opini bersifat positif sebesar 0.80 dan opini bersifat negatif sebesar 0.78. Tingkat keberhasilan sistem dalam menemukan kembali sebuah informasi (*class recall*) menggunakan metode *Naïve Bayes* menghasilkan opini bersifat positif sebesar 0.97 dan opini bersifat negatif sebesar 0.28. Visualisasi Kurva ROC pada gambar 9 nilai AUC model klasifikasi dengan metode *Naïve Bayes* sebesar 0.836 dengan diagnosa klasifikasi baik.

```

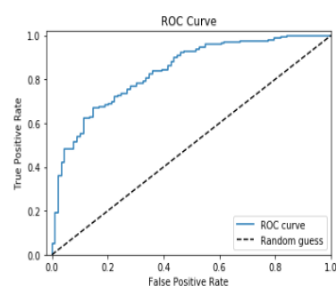
In [24]: #klasifikasi precision, recall, score MNB
print(classification_report(y_test, classifierMNB_pred))

```

	precision	recall	f1-score	support
0	0.78	0.28	0.41	89
1	0.80	0.97	0.88	265
accuracy			0.80	354
macro avg	0.79	0.63	0.65	354
weighted avg	0.80	0.80	0.76	354

Sumber : (Mardiana et al., 2019)

Gambar 8. Pengujian Model Klasifikasi Dengan Metode Naïve Bayes



```

In [26]: #ROC AUC Score MNB
print(roc_auc_score(y_test, MNB_pred_prob))

0.836124655501378

```

Sumber : (Mardiana et al., 2019)

Gambar 9. Kurva ROC Dengan Metode Naïve Bayes

2. Metode Neural Network

Hasil pengujian terhadap opini usaha waralaba dengan metode *Neural Network* disajikan pada gambar 10. Nilai akurasi diperoleh sebesar 0.827 atau 83%. Ketepatan antara informasi yang diminta dengan prediksi yang diberikan (*class precision*) opini bersifat positif sebesar 0.84 dan opini bersifat negatif sebesar 0.74. Tingkat keberhasilan sistem dalam menemukan kembali sebuah informasi (*class recall*) menggunakan metode *Neural Network* menghasilkan opini bersifat positif sebesar 0.94 dan opini bersifat negatif sebesar 0.48. Visualisasi Kurva ROC pada gambar 11 nilai AUC model klasifikasi dengan metode *Neural Network* sebesar 0.833 dengan diagnosa klasifikasi baik.

```

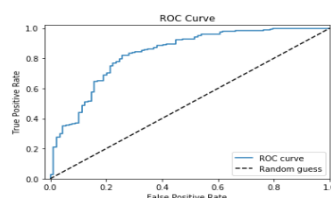
In [27]: #klasifikasi precision, recall, score NN
print(classification_report(y_test, classifierNN_pred))

```

	precision	recall	f1-score	support
0	0.74	0.48	0.59	89
1	0.84	0.94	0.89	265
accuracy			0.83	354
macro avg	0.79	0.71	0.74	354
weighted avg	0.82	0.83	0.81	354

Sumber : (Mardiana et al., 2019)

Gambar 10. Pengujian Model Klasifikasi Dengan Metode Neural Network



```

In [29]: #ROC AUC Score NN
print(roc_auc_score(y_test, NN_pred_prob))

0.8327750688997244

```

Sumber : (Mardiana et al., 2019)

Gambar 11. Kurva ROC Dengan Metode Neural Network

3. Metode *K-Nearest Neighbors*

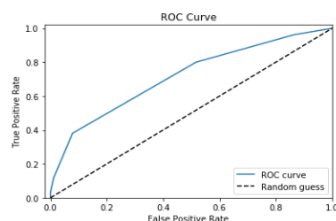
Hasil pengujian terhadap opini usaha waralaba dengan metode *K-Nearest Neighbor* disajikan pada gambar 12. Nilai akurasi diperoleh sebesar 0.516 atau 52%. Ketepatan antara informasi yang diminta dengan prediksi yang diberikan (*class precision*) opini bersifat positif sebesar 0.94 dan opini bersifat negatif sebesar 0.33. Tingkat keberhasilan sistem dalam menemukan kembali sebuah informasi (*class recall*) menggunakan metode *K-Nearest Neighbor* menghasilkan opini bersifat positif sebesar 0.38 dan opini bersifat negatif sebesar 0.92. Visualisasi Kurva ROC pada gambar 13 nilai AUC model klasifikasi dengan metode *K-Nearest Neighbor* sebesar 0.716 dengan diagnosa klasifikasi cukup.

```
In [34]: #klasifikasi precision, recall, score KNN
print(classification_report(y_test, classifierKNN_pred))
```

	precision	recall	f1-score	support
0	0.33	0.92	0.49	89
1	0.94	0.38	0.54	265
accuracy			0.52	354
macro avg	0.63	0.65	0.52	354
weighted avg	0.78	0.52	0.53	354

Sumber : (Mardiana et al., 2019)

Gambar 12. Pengujian Model Klasifikasi Dengan Metode *K-Nearest Neighbors*



```
In [36]: #ROC AUC Score KNN
print(roc_auc_score(y_test, KNN_pred_prob))
```

0.71562433750265

Sumber : (Mardiana et al., 2019)

Gambar 13. Kurva ROC Dengan Metode *K-Nearest Neighbors*

4. Metode *Support Vector Machine*

Hasil pengujian terhadap opini usaha waralaba dengan metode *Support Vector Machine* disajikan pada gambar 14. Nilai akurasi diperoleh sebesar 0.827 atau 83%. Ketepatan antara informasi yang diminta dengan prediksi yang diberikan (*class precision*) opini bersifat positif sebesar 0.88 dan opini bersifat negatif sebesar 0.66. Tingkat keberhasilan sistem dalam menemukan kembali sebuah informasi (*class recall*) menggunakan metode *Support Vector Machine* menghasilkan opini bersifat positif sebesar 0.89 dan opini bersifat negatif sebesar 0.64. Visualisasi Kurva ROC pada gambar 15 nilai AUC model klasifikasi dengan *Support Vector*

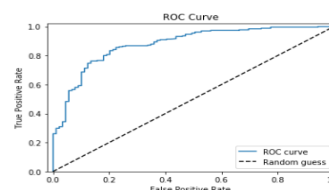
Machine sebesar 0.879 dengan diagnosa klasifikasi baik.

```
In [41]: #klasifikasi precision, recall, score SVM
print(classification_report(y_test, classifierSVM_pred))
```

	precision	recall	f1-score	support
0	0.66	0.64	0.65	89
1	0.88	0.89	0.89	265
accuracy			0.83	354
macro avg	0.77	0.77	0.77	354
weighted avg	0.83	0.83	0.83	354

Sumber : (Mardiana et al., 2019)

Gambar 14. Pengujian Model Klasifikasi Dengan Metode *Support Vector Machine*



```
In [43]: #ROC AUC Score SVM
print(roc_auc_score(y_test, SVM_pred_prob))
```

0.879499682001272

Sumber : (Mardiana et al., 2019)

Gambar 15. Kurva ROC Dengan Metode *Support Vector Machine*

5. Metode *Decision Tree*

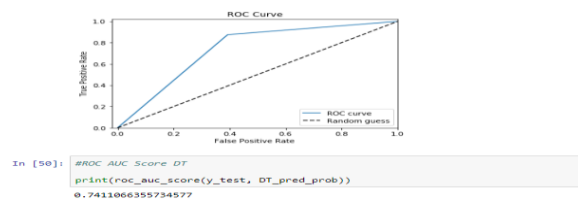
Hasil pengujian terhadap opini usaha waralaba dengan metode *Decision Tree* disajikan pada gambar 16. Nilai akurasi diperoleh sebesar 0.807 atau 81%. Ketepatan antara informasi yang diminta dengan prediksi yang diberikan (*class precision*) opini bersifat positif sebesar 0.87 dan opini bersifat negatif sebesar 0.62. Tingkat keberhasilan sistem dalam menemukan kembali sebuah informasi (*class recall*) menggunakan metode *Decision Tree* menghasilkan opini bersifat positif sebesar 0.88 dan opini bersifat negatif sebesar 0.61. Visualisasi Kurva ROC pada gambar 17 nilai AUC model klasifikasi dengan *Decision Tree* sebesar 0.741 dengan diagnosa klasifikasi cukup.

```
In [48]: #klasifikasi precision, recall, score DT
print(classification_report(y_test, classifierDT_pred))
```

	precision	recall	f1-score	support
0	0.62	0.61	0.61	89
1	0.87	0.88	0.87	265
accuracy			0.81	354
macro avg	0.74	0.74	0.74	354
weighted avg	0.81	0.81	0.81	354

Sumber : (Mardiana et al., 2019)

Gambar 16. Pengujian Model Klasifikasi Dengan Metode *Decision Tree*



Sumber : (Mardiana et al., 2019)

Gambar 17. Kurva ROC Dengan Metode *Decision Tree*

B. Pembahasan

Setiap model klasifikasi memiliki akurasi yang berbeda antara metode satu dengan yang lainnya. Hasil pengujian akurasi pada metode *Naïve Bayes* sebesar 80%, metode *Neural Network* sebesar 83%, *K-Nearest Neighbor* sebesar 52%, *Support Vector Machine* 83%, dan *Decision Tree* mendapatkan hasil sebesar 81%.

Tabel 3. Nilai Akurasi Model Klasifikasi

Model	Akurasi	Klasifikasi
<i>Naïve Bayes</i>	80%	Good
<i>Neural Network</i>	83%	Good
<i>K-Nearest Neighbors</i>	52%	Fail
<i>Support Vector Machine</i>	83%	Good
<i>Decision Tree</i>	81%	Good

Sumber : (Mardiana et al., 2019)

Selain itu, evaluasi pengukuran dengan menggunakan *precision*, *recall* dan *F1-score* dihasilkan untuk memastikan model yang telah dibuat dapat berkinerja dengan baik. Nilai rata-rata dari tabel evaluasi pengukuran untuk klasifikasi positif mendapat nilai sebesar 83% dari semua data evaluasi pengukuran dari masing-masing metode klasifikasi.

Tabel 4. Evaluasi Pengukuran

Model Klasifikasi		Negatif	Positif
<i>Naïve Bayes</i>	<i>Precision</i>	0.78	0.80
	<i>Recall</i>	0.28	0.97
	<i>F1-Score</i>	0.41	0.88
<i>Neural Network</i>	<i>Precision</i>	0.74	0.84
	<i>Recall</i>	0.48	0.94
	<i>F1-Score</i>	0.59	0.89
<i>K-Nearest Neighbors</i>	<i>Precision</i>	0.33	0.94
	<i>Recall</i>	0.92	0.38
	<i>F1-Score</i>	0.49	0.54
<i>Support Vector Machine</i>	<i>Precision</i>	0.66	0.88
	<i>Recall</i>	0.64	0.89
	<i>F1-Score</i>	0.65	0.89
<i>Decision Tree</i>	<i>Precision</i>	0.62	0.87
	<i>Recall</i>	0.61	0.88
	<i>F1-Score</i>	0.61	0.87
Rata-rata		0.59	0.83

Sumber : (Mardiana et al., 2019)

Klasifikasi yang dihasilkan dari kurva ROC juga menunjukkan persentase yang beragam. Pada

klasifikasi ini *score* tertinggi diperoleh dengan perhitungan menggunakan metode *Support Vector Machine*. Ini menunjukkan bahwa, pada penelitian ini penggunaan metode *Support Vector Machine* lebih baik dibandingkan dengan empat metode lainnya. Hasil selengkapnya dapat di lihat pada Tabel 5.

Tabel 5. Perbandingan Kurva ROC

Model	AUC	Diagnosa ROC
<i>Naïve Bayes</i>	84%	Good
<i>Neural Network</i>	83%	Good
<i>K-Nearest Neighbors</i>	72%	Fair
<i>Support Vector Machine</i>	88%	Good
<i>Decision Tree</i>	74%	Fair

Sumber : (Mardiana et al., 2019)

Pada Tabel 6 merupakan hasil dari perhitungan untuk mencari nilai rata-rata *precision score* dari masing-masing metode klasifikasi yang ada pada Tabel 5.

Tabel 6. Average Class Precision

Model	Class Precision
<i>Naïve Bayes</i>	93%
<i>Neural Network</i>	92%
<i>K-Nearest Neighbors</i>	86%
<i>Support Vector Machine</i>	95%
<i>Decision Tree</i>	85%

Sumber : (Mardiana et al., 2019)

Berdasarkan hasil penelitian yang telah dilakukan, nilai akurasi yang terdapat pada Tabel 3 menunjukkan bahwa, penggunaan metode *Support Vector Machine* lebih baik dibandingkan empat metode lainnya, khususnya untuk pengolahan data berbasis *text*. Metode ini mendapatkan nilai akurasi sebesar 83% dan akurasi dari kurva ROC sebesar 88%. Data ini menunjukkan, pengklasifikasian dengan metode tersebut dinilai cukup baik. Selain itu *precision score* yang dihasilkan juga menunjukkan klasifikasi yang baik dengan nilai sebesar 95%. Akan tetapi, penggunaan metode *Naïve Bayes* juga menghasilkan nilai sebesar 80% untuk akurasinya dan 84% untuk akurasi dari kurva ROC. Ini membuktikan bahwa penggunaan metode ini pun juga disarankan terhadap penelitian analisa teks. Selain mendapatkan nilai yang baik atau *good classification* pada kelasnya, metode ini juga terbilang simpel dengan performa yang relatif kuat untuk metode yang mengandalkan asumsi secara keseluruhan.

KESIMPULAN

Dalam penelitian ini, kami telah membandingkan lima metode klasifikasi dalam menentukan sentimen usaha waralaba berdasarkan opini dari Twitter. Kami menerapkan pra-pemrosesan yang sama pada kelima metode

klasifikasi. Berdasarkan hasil pra-pemrosesan, kami memproses data ke dalam metode klasifikasi untuk membandingkan lima metode dan memperoleh metode yang memiliki akurasi tertinggi. Pada pembahasan menunjukkan bahwa *Naive Bayes*, *Neural Network*, *K-Nearest Neighbor*, *Support Vector Machines*, dan *Decision Tree* memperoleh hasil yang berbeda dalam akurasi serta dalam recall dan skor f1 terutama dalam sentimen positif yang memiliki skor kecil. Namun, di antara kelima metode klasifikasi, dapat dilihat bahwa *Neural Network* dan *Support Vector Machine* menghasilkan akurasi tertinggi sebesar 83%. Sedangkan *Decision Tree* memperoleh 81%, *Naive Bayes* sebesar 80%, dan *K-Nearest Neighbor* sebesar 52%. Berdasarkan akurasi yang dihasilkan dari masing-masing metode perhitungan, *Neural Network*, *Support Vector Machine* dan *Naive Bayes* dapat dikategorikan sebagai *Good Clasification*, sedangkan *Decision Tree* dan *K-Nearest Neighbor* dikategorikan sebagai *Fair Clasification*. Oleh karena itu, algoritma *Neural Network*, *Support Vector Machine*, dan *Naive Bayes* cocok untuk analisis sentimen dalam penelitian ini. Dataset juga menunjukkan bahwa usaha waralaba memiliki banyak sentimen positif. Ini berarti usaha waralaba masih menjadi pilihan bagi masyarakat ingin memulai usaha. Penelitian lanjutan dalam analisis sentimen usaha waralaba dapat menggunakan pendekatan *hybrid* untuk mengoptimasi metode *Neural Network*, *Support Vector* dan *Naive Bayes* sehingga memiliki akurasi yang lebih baik.

REFERENSI

- Attabi, A. W., Muflikhah, L., & Fauzi, M. A. (2018). Penerapan Analisis Sentimen untuk Menilai Suatu Produk pada Twitter Berbahasa Indonesia dengan Metode Naive Bayes Classifier dan Information Gain. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer (J-PTIIK) Universitas Brawijaya*, 2(11), 4548–4554.
- Fathurahman, M. F., Windarti, A., & Purwanto, I. (2018). Pengaruh Value dan Physical Benefit Produk Waralaba Terhadap Kepuasan Konsumen. *Journal of Applied Business and Economics*, 4(4), 305–319.
- Imanuwelita, V., Putri, R. R. M., & Amalia, F. (2018). Penentuan Kelayakan Lokasi Usaha Franchise Menggunakan Metode AHP dan VIKOR. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 2(1), 122–132.
- Ling, J., Kencana, I. P. E., & Oka, T. B. (2014). Analisis Sentimen Menggunakan Metode Naive Bayes Classifier Dengan Seleksi Fitur Chi Square. *E-Jurnal Matematika*, 3(3), 92. <https://doi.org/10.24843/mtk.2014.v03.i03.p070>
- Mardiana, T., Syahreva, H., & Tuslaela, T. (2019). *Laporan Akhir Penelitian: Komparasi Metode Klasifikasi Pada Analisis Sentimen Usaha Waralaba Berdasarkan Data Twitter*. Jakarta.
- Muthia, D. A. (2017). Analisis Sentimen Pada Review Restoran Dengan Teks Bahasa Indonesia Menggunakan Algoritma Naive Bayes. *Jurnal Ilmu Pengetahuan Dan Teknologi Komputer*, 2(2), 39–45. <https://doi.org/10.1515/HUMOR.2006.009>
- Nurhuda, F., Sihwi, S. W., & Doewes, A. (2013). Analisis Sentimen Masyarakat terhadap Calon Presiden Indonesia 2014 berdasarkan Opini dari Twitter Menggunakan Metode Naive Bayes Classifier. *Jurnal Teknologi & Informasi ITSmart*, 2(2), 35–42. <https://doi.org/10.20961/its.v2i2.630>
- Romadloni, N. T., Santoso, I., & Budilaksono, S. (2019). Perbandingan Metode Naive Bayes , Knn Dan Decision Tree Terhadap Analisis Sentimen Transportasi Krl. *Jurnal IKRA-ITH Informatika*, 3(2), 1–9.
- Vinodhini, G., & Chandrasekaran, R. M. (2016). A comparative performance evaluation of neural network based approach for sentiment classification of online reviews. *Journal of King Saud University - Computer and Information Sciences*, 28(1), 2–12. <https://doi.org/10.1016/j.jksuci.2014.03.024>
- Widaningsih, S., & Suheri, A. (2018). Klasifikasi Jurnal Ilmu Komputer Berdasarkan Pembagian Web of Science Dengan Menggunakan Text Mining. *Seminar Nasional Teknologi Informasi Dan Komunikasi 2018 (SENTIKA 2018)*, 2018(Sentika), 23–24.