

BAB II

LANDASAN TEORI

2.1 Penelitian Terdahulu

Pada tabel 2.1 menampilkan ringkasan penelitian sebelumnya yang fokus pada sentimen analisis dengan objek vaksin covid-19 dan proses vaksinasi.

Tabel 2. 1 Penelitian Terdahulu

Pengarang	Judul	Hasil
Cotfas et al., 2021	<i>The Longest Month: Analyzing COVID-19 Vaccination Opinions Dynamics From Tweet in the Month Following the First Vaccine Announcement</i>	Pengumpulan <i>tweet</i> dilakukan pada periode 9 November 2020 – 8 Desember 2020, dengan topik “ <i>covid</i> ” dan “ <i>vaccination</i> ” dengan kata kunci spesifik: “ <i>covid19, covid-19, coronavirus, coronaoutbreak, coronaviruspandemic, wuhanvirus, vaccination, vaccinate, vaccinating, vaccinated</i> ”. Polaritas sentimen dibagi menjadi tiga : “ <i>against, neutral, dan in favor</i> ”. Metode klasifikasi yang digunakan adalah <i>Multinomial Naive Bayes (MNB), Random Forest (RF), Support Vector Machine (SVM), Bidirectional Long Short-Term Memory(Bi-LSTM), dan Convolutional Neural Network(CNN)</i> .
Fitriana et al., 2021	Analisis Sentimen Opini Terhadap Vaksin Covid - 19 pada Media Sosial Twitter Menggunakan <i>Support Vector Machine (SVM)</i> dan <i>Naive Bayes (NB)</i>	<i>Crawling tweet</i> dilakukan dengan menggunakan kata kunci “ vaksin covid-19”. Metode yang digunakan SVM dan NB. SVM memiliki nilai akurasi, presisi, dan <i>recall</i> lebih baik dibandingkan dengan <i>Naive Bayes</i> .
Naraswati et al., 2021	Analisis Sentimen Publik dari Twitter Tentang Kebijakan Penanganan Covid-19 di Indonesia dengan <i>Naive Bayes Classification</i>	<i>Crawling tweet</i> dilakukan pada periode 5-7 Oktober 2020 dengan kata kunci “PSBB, wajib masker, dan jam malam”. Polaritas dibagi dua: positif dan negatif. <i>Naive Bayes Classification</i> memiliki nilai cukup baik dengan nilai akurasi 87,34%, sensitivitas 93,43%, dan spesifisitas 71,76%.

Sattar & Arifuzzaman, 2021	<i>COVID-19 Vaccination Awareness and Aftermath: Public Sentiment Analysis on Twitter Data and Vaccinated Population Prediction in the USA</i>	Analisis sentimen menggunakan metode <i>Unsupervised Lexicon-Based Approaches</i> dan tools yang digunakan <i>TextBlob</i> dan <i>VADER</i> . <i>Crawling tweet</i> menggunakan kata kunci “ <i>a safe, healthy lifestyle after vaccination</i> ” dan 7 jenis vaksin covid-19 yaitu: <i>Pfizer, Moderna, Johnson & Johnson, Oxford-AstraZeneca, Sputnik V, Covavin, Sinovac</i> . <i>Crawling tweet</i> dilakukan pada rentang 10 April 2021 – 17 May 2021.
Hikmawan et al., 2020	Sentimen Analisis Publik Terhadap Joko Widodo terhadap wabah Covid-19 menggunakan Metode <i>Machine Learningpp</i>	<i>Crawling tweet</i> menggunakan kata kunci “Jokowi” dan “covid” dalam bahasa Indonesia. Polaritas dibagi kedalam tiga klasifikasi: positif, negatif, dan netral. Metode klasifikasi yang digunakan <i>Naive Bayes (NB)</i> , <i>Support Vector Machine (SVM)</i> , dan <i>K-Nearest Neighbor (KNN)</i> . SVM menghasilkan nilai akurasi terbaik dari ketiga metode yang dipakai yaitu sebesar 84,58%.
Nomleni, 2015	Sentiment Analysis Menggunakan <i>Support Vector Machine (SVM)</i>	<i>Support Vector Machine</i> menghasilkan nilai akurasi yang sangat baik, dengan nilai diatas 80% dari 7 kali percobaan yang dilakukan.
Kunal et al., 2018	<i>Textual Dissection of Live Twitter Reviews Using Naive Bayes</i>	Algoritma yang digunakan adalah algoritma <i>Naive Bayes</i> dengan dua <i>library python</i> yaitu <i>Tweepy</i> Dan <i>TextBlob</i> . <i>TextBlob</i> untuk menghitung nilai polaritas dan <i>Tweepy</i> untuk akses ke <i>live tweet</i> pada <i>Twitter</i> .
Villavicencio et al., 2021)	<i>Twitter Sentiment Analysis towards COVID-19 Vaccines in the Philippines Using Naive Bayes</i>	<i>Tweet</i> yang dijadikan objek adalah <i>tweet</i> dalam bahasa Inggris dan tagalog. Polaritas diklasifikasi menjadi tiga : positif, negatif, dan netral. Hasil akurasi <i>Naive Bayes Classification</i> cukup baik yaitu 81,77%.

Nurdeni et al., 2021	<i>Sentiment Analysis on Covid19 Vaccines in Indonesia: From The Perspective of Sinovac and Pfizer</i>	<i>Tweet</i> yang digunakan adalah <i>tweet</i> dalam bahasa Indonesia yang <i>crawling</i> pada rentang Oktober – November 2020 dengan kata kunci “Sinovac” dan “Pfizer”. Pelabelan polaritas dilakukan secara manual, yang terdiri dari tiga label yaitu positif, negatif, dan netral. Dengan menggunakan evaluasi kinerja model yang digunakan yaitu <i>Support Vector Machine</i> dengan menggunakan <i>10-fold cross-validation</i> nilai akurasi untuk <i>dataset Sinovac</i> sebesar 85% dan nilai akurasi untuk <i>dataset Pfizer</i> sebesar 78%.
Laurensz & Sedyono, 2021, p. 19	Analisis Sentimen Masyarakat terhadap Tindakan Vaksinasi dalam Upaya Mengatasi Pandemi Covid-19	<i>Tweet</i> di <i>crawling</i> menggunakan dua kata kunci “vaksin merah putih” dan “vaksin sinovac” yang dilakukan pada periode rentang November 2020 – Februari 2021. Polaritas dibagi menjadi dua kelas: positif dan negatif. <i>Naive Bayes</i> menghasilkan nilai akurasi terbaik sebesar 85,59% sedangkan SVM sebesar 84,41%.
Aljameel et al., 2021	<i>A Sentiment Analysis Approach to Predict an Individual's Awareness of the Precautionary Procedures to</i>	<i>Tweet</i> menggunakan bahasa Arab dengan kata kunci #covid #corona #stay_home, #wearing_mask #sterilization #distance_learning #washing_hands #social_distance. <i>Dataset tweet</i> di download dari <i>Twitter</i> menggunakan <i>TWINT</i> pada periode 23 Maret 2020 - 21 Juni 2020 di Arab Saudi. Dalam tulisan dan bahasa Arab. Polaritas diklasifikasi

Prevent COVID-19 Outbreaks in Saudi Arabia ke dalam tiga kelas: positif, negatif, dan netral. Metode *Machine Learning* yang digunakan *Naïve Bayes*, *K-Nearest Neighbor*, dan *Support Vector Machine*. Menggunakan *N-Gram model* (*unigram*, *unigram TF-IDF*, *bigram*, *bigram TF-IDF*) untuk komparasi model *Natural Language Processing (NLP)*. Dari ketiga metode yang digunakan, SVM, KNN, dan NB, SVM mendapatkan nilai akurasi terbaik yaitu 85%.

2.2 Machine Learning

Machine learning merupakan sub-bagian dari *Artificial Intelligence (AI)*. *Machine learning* merupakan sekumpulan metode yang secara otomatis dapat mendeteksi pola di dalam data, dan kemudian menggunakan pola yang telah ditemukan untuk memprediksi data di masa mendatang, atau untuk melakukan pengambilan keputusan dalam ketidakpastian [13]. *Machine learning* merupakan penggabungan dari sekumpulan metode pembelajaran yang telah dilatih kedalam sebuah mesin [14]. *Machine learning* dapat dikategorikan menjadi dua yaitu *supervised learning* and *unsupervised learning*.

a. Supervised Learning

Supervised learning merupakan salah satu metode yang ada di dalam *machine learning* [15]. Di dalam *supervised learning*, algoritma dibangun berdasarkan *dataset* yang sudah diberi label [16]. Tujuan dari *supervised learning* adalah memprediksi nilai label (t) untuk input (x) yang tidak ada di dalam set pelatihan. Dengan kata lain supervised learning bertujuan untuk menggeneralisasi pengamatan dalam kumpulan *data* (D) ke input baru [17] seperti pada persamaan 2.2.

$$D = \{x_n, t_n\} \quad n = 1 \quad (2.2)$$

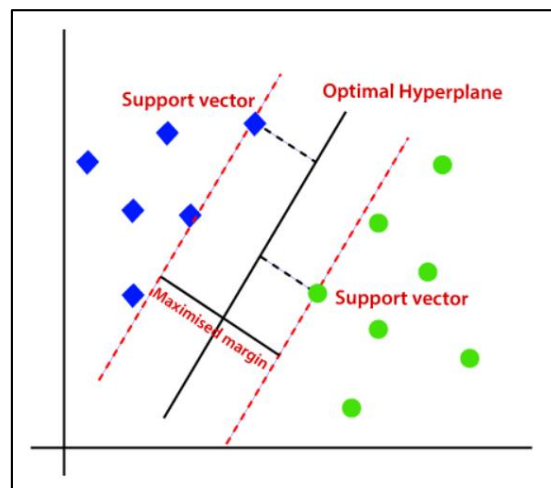
Algoritma *supervised learning* memanfaatkan *dataset* uji dari *data input* dan *output*. Algoritma ini mempelajari hubungan antara *data input* dan *data output* dari *dataset* uji dan kemudian menggunakan hubungan ini untuk memprediksi *output* pada data baru. Salah satu tujuan *supervised learning* yang paling umum adalah klasifikasi. Salah satu tujuan klasifikasi adalah untuk menggunakan informasi yang telah dipelajari untuk memprediksi keanggotaan kelas tertentu [18].

Algoritma yang masuk kedalam kelompok algoritma *supervised learning* adalah *Decision Tree*, *K-Nearest Neighbor*, *Naive Bayes Classifier*, *Artificial Neural Network*, *Support Vector Machine*, dan *Fuzzy K-Nearest Neighbor*.

Logistic Regression, Support Vector Machine, Naive Bayes, dan Neural Network masuk ke dalam grup klasifikasi (*classification*) algoritma *supervised learning* [19].

1. Support Vector Machine

Support Vector Machine (SVM) pertama kali dikenalkan oleh Cortes and Vapnik [17]. SVM adalah sebuah algoritma *supervised learning* yang menganalisis data, mengenali pola, yang digunakan didalam *classification* dan *regression analysis*. SVM memiliki reputasi yang baik dalam hal klasifikasi [20] dan secara efisien melakukan klasifikasi *non-linier* menggunakan apa yang disebut trik kernel, secara implisit memetakan inputnya ke dalam ruang fitur berdimensi tinggi. *Support Vector Machine* akan memaksimalkan margin diantara dua kelas [14] seperti pada gambar 2.1.



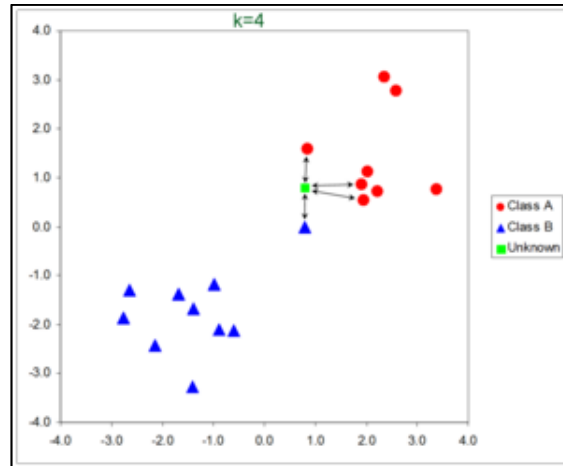
Gambar 2. 1 Support Vector Machine

2. K-Nearest Neighbor

K-Nearest Neighbor (KNN) merupakan model pengenalan pola yang dapat digunakan untuk klasifikasi (*classification*) maupun regresi (*regression*).

K-NN didasarkan pada gagasan bahwa pola terdekat dengan pola target x_1 , yang dicari labelnya, akan memberikan informasi berguna dalam pemberian

label yang diilustrasikan pada gambar 2.2. KNN memberikan label kelas dari sebagian besar informasi label yang berguna [21].



Gambar 2. 2 K –Nearest Neighbor

3. Naive Bayes Classification

Naive Bayes (NB) merupakan metode pengklasifikasian dengan metode probabilistik dan statistik yang pertama kali dikemukakan oleh ilmuwan inggris Thomas Bayes. NB masuk kedalam kelompok *supervised learning* dengan berdasarkan teori Bayes [22] seperti pada persamaan 2.3.

$$P(c|d) = p(c) \prod_{1 \leq k \leq n, d \leq n} p(t_k | c) \quad (2.3)$$

Naive Bayes Classification (NBC) biasanya digunakan dalam klasifikasi teks dan dokumen [23]. Dua varian algoritma *Naive Bayes* yang digunakan dalam klasifikasi sentimen adalah *Bernoulli* dan *MultinomialNB* [24].

b. Text Mining

Text mining atau *Text Analytics* merupakan sebuah penggabungan disiplin ilmu *language science* dan *computer science* dengan menggunakan ilmu statistik dan

teknik mesin *learning*. *Text mining* digunakan untuk analisis teks dan mengubahnya ke dalam bentuk yang terstruktur [16].

Identifikasi topik, pengelompokan teks, *spam filtering*, dan analisa sentimen merupakan bagian dari *text mining*. Konsep penting di dalam *text mining* adalah *Bag of Words*, yang merupakan untuk mempresentasikan dokumen dengan menyatakannya dalam bentuk vektor dari kata-kata yang terkandung di dalamnya. *Bag of words* merupakan representasi dokumen yang selanjutnya dapat diolah untuk berbagai tujuan seperti klasifikasi dokumen dan lain-lainnya.

c. Sentiment Analysis

Sentimen Analisis adalah studi komputasi mengenai sikap, emosi, pendapat, penilaian, pandangan dari sekumpulan teks yang difokuskan adalah mengekstraksi, mengidentifikasi atau menemukan karakteristik sentimen dalam unit *text* menggunakan metode *Natural Language Processing*, statistik, atau *machine learning* [25]. Sentimen analisis banyak digunakan di dalam dunia bisnis, kesehatan, dan politik.

Tools yang banyak menyediakan dukungan untuk berbagai tugas pada *Natural Language Processing (NLP)* termasuk analisa sentimen adalah *TextBlob*. *Textblob* menggunakan bahasa *Python*, yaitu bahasa pemrograman tingkat tinggi yang berorientasi objek dengan semantik yang dinamis.

Analisa sentimen oleh *Textblob* berdasarkan nilai *polarity* dan *subjectivity*. Nilai *polarity* terletak pada rentang $[-1, 1]$, -1 menunjukkan sentimen negatif, 1 menunjukkan sentimen positif, dan 0 menunjukkan sentimen netral. Nilai *Subjectivity* terletak pada rentang $[0, 1]$. *Subjectivity* pada umumnya mengacu pada pendapat, emosi, atau penilaian pribadi, sedangkan penilaian objektif mengacu pada informasi fakta. Sentimen yang lebih objektif daripada subjektif menerima skor yang lebih rendah, yang menunjukkan kemungkinan fakta yang lebih akurat.

d. TF/IDF

Term Frequency – Inverse Document Frequency (TF-IDF) merupakan teknik dalam memberikan bobot pada kata terhadap dokumen seperti pada persamaan 2.3 berikut [16].

$$TF = f_{t,d} \quad (2.3)$$

TF adalah frekuensi (f) dari (t) di dalam dokumen (d).

Rumus dari IDF dengan penskalaan logarithmic yang paling umum dengan persamaan 2.4 dibawah ini.

$$IDF = \log(N/|\{d \in D: t \in d\}|) \quad (2.4)$$

Dengan N jumlah total dokumen di dalam *Corpus* dan $(|\{d \in D: t \in d\}|)$ menjadi jumlah dokumen (d) dimana (t) muncul.

$$\frac{TF}{IDF} = f_{t,d}/\log(N/|\{d \in D: t \in d\}|) \quad (2.5)$$

$TF-IDF$ pada persaman 2.5 memperhitungkan semua dokumen. IDF memberikan gambaran tentang seberapa umum kata tersebut di seluruh *Corpus*: semakin tinggi frekuensi dokumen, semakin umum, dan semakin banyak kata yang kurang informatif. Beberapa langkah dapat dilakukan untuk memanipulasi data sebelum sampai ke tahapan *Bag of words* yaitu proses *tokenization*.

e. API Twitter

API Twitter merupakan fitur yang disediakan oleh *Twitter* sehingga pihak pengembang aplikasi dapat mengakses informasi *web Twitter* seperti *tweet*, *users*, *direct messages*, *lists*, *trends*, *media*, dan *places* [26].

f. WordCloud

Word Cloud merupakan teknik untuk memvisualisasikan kata yang sering muncul dalam teks di mana ukuran kata mewakili frekuensinya.