

## BAB II

### TINJAUAN PUSTAKA DAN LANDASAN TEORI

#### 2.1 Tinjauan Pustaka

Penelitian ini berfokus pada analisis sentimen dan topik di media sosial menggunakan algoritma BERT (*Bidirectional Encoder Representations from Transformers*), *Latent Dirichlet Allocation* (LDA), dan *Geo-Location* dengan data dari X untuk memahami persepsi publik terhadap kebijakan daerah di Jawa Timur. Beberapa penelitian sebelumnya telah menggunakan metode serupa, seperti Naïve Bayes dan LDA dalam studi oleh Bilal dkk. (2016), Akbar dkk. (2015), dan Kusumaningrum dkk. (2014) yang juga menggunakan model berbasis BoW (*Bag of Words*). Meskipun pendekatan tersebut sama, penelitian ini menambahkan fitur *POS Tagging* dan *Geo-Location* untuk memperkaya analisis, suatu hal yang belum banyak diterapkan dalam studi terdahulu. Sumber data yang digunakan juga menjadi pembeda, di mana sebagian besar studi sebelumnya, seperti penelitian Tripathy dkk. (2016) dan Lin dan He (2009), menggunakan data IMDb yang lebih spesifik pada sentimen film, bukan opini publik terkait kebijakan. Selain itu, meskipun Alsaqer dkk. (2023) juga memanfaatkan *Geo-Location* pada data X, pendekatan mereka lebih umum tanpa pemisahan analisis topik dan sentimen mendetail. Dalam penelitian ini, penggunaan *Geo-Location* memungkinkan pemetaan sentimen berdasarkan wilayah tertentu, menawarkan wawasan yang lebih kontekstual mengenai bagaimana kebijakan dipersepsikan oleh masyarakat di berbagai lokasi. Penelitian ini menghadirkan inovasi dengan memadukan LDA untuk mengidentifikasi topik yang muncul di media sosial terkait kebijakan dengan sentimen yang diklasifikasikan berdasarkan lokasi geografis. Dengan mengkombinasikan BERT untuk klasifikasi sentimen, LDA untuk topik, dan analisis geospasial, pendekatan ini menawarkan cara pandang baru yang lebih mendalam tentang persepsi publik di wilayah tertentu. Hasilnya diharapkan dapat memberikan dasar rekomendasi kebijakan publik yang lebih efektif dan sesuai dengan karakteristik masyarakat setempat, menjadikan penelitian ini lebih

komprehensif daripada studi sebelumnya yang terbatas pada analisis sentimen atau topik tanpa dimensi spasial yang terintegrasi.

## 2.2 Landasan Teori

### 2.2.1 *Machine Learning*

*Machine learning* atau dikenal dengan pembelajaran mesin adalah ilmu komputer yang bisa bekerja tanpa diprogram secara eksplisit (Samuel, 1959). Dalam arti lain, *machine learning* adalah kemampuan komputer untuk belajar dari data tanpa harus diberi tahu secara eksplisit apa yang harus dilakukan. Komputer dapat belajar dari data untuk membuat prediksi, klasifikasi, atau keputusan (Deepublish, 2020). Program tersebut memanfaatkan data untuk membangun model dan mengambil keputusan berdasarkan model yang telah dibangun. Menurut (Samuel, 1959) ML berisi sebuah algoritma yang bersifat general (umum) dimana algoritma tersebut dapat menghasilkan sesuatu yang menarik atau bermanfaat dari sejumlah data tanpa harus menulis kode yang spesifik. Pada intinya, algoritma yang general tersebut ketika diberikan sejumlah data maka akan dapat membangun sebuah aturan atau model atau inferensi dari data tersebut (Id, 2021).

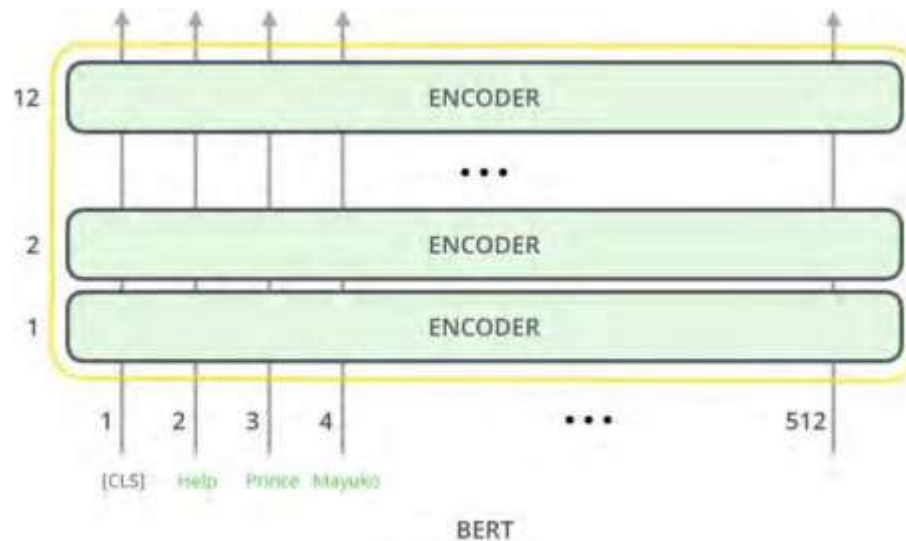
### 2.2.2 Algoritma BERT

Algoritma BERT merupakan singkatan dari *Bidirectional Encoder Representations from Transformers* sebuah arsitektur model Bahasa yang dikembangkan oleh Google pada tahun 2018. BERT menggunakan pendekatan transformasi yang berbasis pada transformer, yaitu arsitektur neural network yang populer dalam NLP. Arsitektur transformer memungkinkan BERT untuk mengatasi kendala dalam pemrosesan Bahasa seperti pemahaman konteks, hubungan antar kata, dan ambiguitas dalam kalimat (Kurniawan, 2022). Adapun keunggulan dari BERT adalah kemampuannya dalam memahami konteks kata yang lebih baik. BERT melakukan pelatihan pada korpus dengan menggunakan text besar dengan metode dengan metode *self-supervised learning*, di mana model belajar untuk memprediksi kata-kata yang hilang dalam kalimat (Fitria, 2022). Hal ini membantu

model BERT untuk mengembangkan representasi kata yang lebih kaya dan pemahaman yang lebih baik tentang hubungan antara kata-kata dalam kalimat. Selain itu, BERT juga mengimplementasikan pendekatan "*bidirectional*" yang memungkinkannya untuk melihat konteks kata dari kedua arah, baik sebelum dan sesudah kata tersebut. Pendekatan ini memungkinkan BERT untuk menangkap konteks yang lebih luas dan kompleks dalam kalimat. BERT telah berhasil dalam berbagai tugas pemrosesan bahasa alami, termasuk analisis sentimen, pemahaman pertanyaan, penandaan entitas, dan pemodelan bahasa.

Model BERT yang telah dilatih dapat digunakan dalam tugas-tugas ini dengan penyesuaian pada lapisan output sesuai dengan kebutuhan spesifik dari tugas tersebut. Dengan kemampuannya yang luar biasa dalam pemahaman bahasa, BERT telah menjadi landasan penting dalam pengembangan berbagai aplikasi NLP dan telah mencapai hasil yang sangat baik dalam meningkatkan kinerja dalam berbagai tugas pemrosesan bahasa alami (Elhan dkk., 2022) .

BERT bekerja dengan memanfaatkan arsitektur *transformer* untuk mempelajari hubungan kontekstual antara kata-kata atau sub kata-kata dalam sebuah teks. Arsitektur *transformer* terdiri dari dua bagian utama, yaitu *encoder* dan *decoder*. *Encoder* BERT bertanggung jawab untuk membaca input teks dan membangun representasi kontekstual dari kata-kata di dalamnya, sedangkan *decoder* digunakan untuk menghasilkan prediksi atau output untuk tugas tertentu. Arsitektur BERT seperti ditunjukkan pada Gambar 2.1.



Gambar 2.1 Arsitektur BERT

Cara kerja BERT dapat dijelaskan sebagai berikut:

1. *Tokenisasi*: Teks input dibagi menjadi token-token, yang bisa berupa kata-kata atau subkata-kata. Setiap token diberi token *embedding* numerik yang merepresentasikan kata tersebut.
2. *Self-Attention*: BERT menggunakan mekanisme *self-attention* dalam lapisan-lapisan transformernya. *Self-attention* memungkinkan model untuk memperoleh pemahaman konteks yang lebih baik dengan memperhatikan hubungan antara token-token dalam teks.
3. *Encoding Kontekstual*: Setiap token dienkripsi menjadi vektor representasi yang memuat informasi tentang konteks dan makna token tersebut. Vektor representasi ini diperbarui berdasarkan perhitungan *self-attention* di lapisan-lapisan *transformer*.
4. *Fine-tuning*: Setelah pelatihan awal, model BERT dapat disesuaikan (*fine-tuned*) untuk tugas-tugas khusus seperti klasifikasi sentimen atau pemodelan bahasa.
5. Proses *Fine Tuning* melibatkan pelatihan tambahan menggunakan dataset yang telah dilabeli dengan benar untuk tugas tersebut.

Mekanisme representasi Input BERT seperti ditunjukkan pada Gambar 2.2



Gambar 2.2 Representasi Input BERT

Dengan memanfaatkan arsitektur *transformer* dan mekanisme *attention*, BERT mampu memahami konteks kata-kata dalam teks secara bidirectional, yaitu melihat konteks sebelum dan sesudah kata tersebut. Hal ini membantu model BERT dalam mengatasi ambiguitas dan menghasilkan representasi kata dengan memperhatikan hubungan kontekstual yang kompleks.



### 2.2.3 Analisis Sentimen

Analisis sentimen dikenal pula dengan *opinion mining*. Pendapat, ekspresi dan perasaan para pelanggan terhadap sebuah produk dianalisis melalui pendapat-pendapat yang mereka tulis untuk mendapatkan kesimpulan tentang opini yang berkembang. Media sosial seperti Blog, Facebook dan X menjadi sesuatu yang memiliki peran penting, karena didalamnya terkandung informasi tentang berbagai ulasan dan komentar yang beragam dari banyak orang dan berbagai macam produk. Pendapat menjadi penting karena seringkali seseorang ataupun suatu organisasi perlu mendengar pendapat orang lain terlebih dahulu sebelum mengambil keputusan (Pang dan Lee, 2008).

Pada analisis sentimen, bentuk analisis data yang digunakan adalah klasifikasi. Klasifikasi merupakan proses untuk mengelompokkan data kedalam suatu kelas yang telah didefinisikan sebelumnya. Tujuan utama analisis sentimen adalah menemukan polaritas dalam sebuah dokumen teks, sehingga proses analisis data dengan bentuk klasifikasi digunakan untuk memprediksi suatu data masuk kedalam kelas positif atau negatif.

Terdapat tiga level dalam analisis sentimen yang dapat dijadikan acuan dalam sebuah penelitian, yaitu level dokumen, level kalimat dan level aspek (Liu, 2012).

#### 1. Level Dokumen.

Pada level ini tujuannya adalah menentukan apakah dokumen secara keseluruhan (misalnya artikel berita, review produk, bernada sentimen positif, negatif atau netral). Analisis sentimen dari ulasan suatu produk, maka hasilnya adalah opini positif atau negatif dari produk tersebut dengan didasarkan pada keseluruhan ulasan. Contoh yang lain adalah artikel berita tentang program Jatim Cerdas yang memuji kualitas Pendidikan secara menyeluruh, berarti dokumen keseluruhan bersentimen positif.

## 2. Level Kalimat

Pada level ini fokus setiap kalimat dalam teks. Tujuannya adalah menilai apakah satu kalimat mengandung opini positif, negative, atau netral. Contoh Pernyataan warga tentang program subsidi transportasi. : “ Bus gratis untuk pelajar sangat membantu. Tetapi jadwalnya sering tidak tepat/terlambat”. kalimat tersebut selanjutnya dianalisis untuk menentukan apakah kalimat tersebut berisi sentimen positif ataukah negatif. Kalimat 1 bersentimen positif, kalimat 2 bersentimen Negatif.

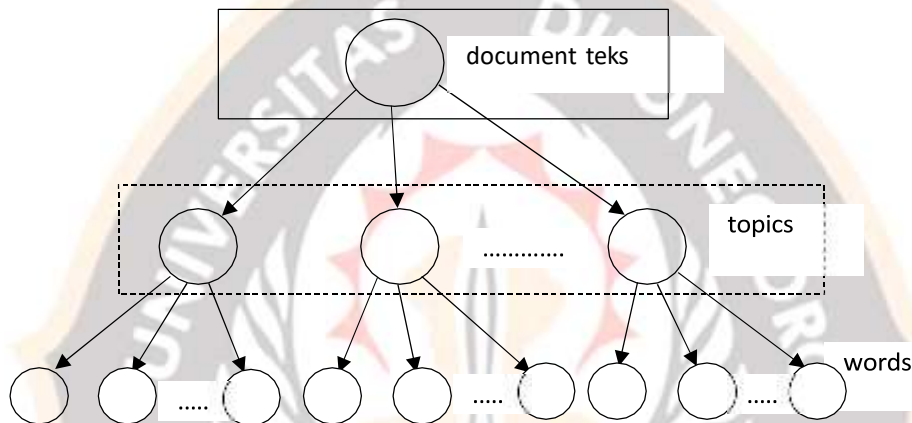
## 3. Level Aspek

Pada level ini, proses analisis dilakukan lebih mendalam terhadap sebuah dokumen. Pada level dokumen dan kalimat, terbatas hanya untuk mengetahui polaritas dari dokumen atau kalimat tersebut, kedua model tersebut belum mampu menggambarkan target opini yang ada pada keseluruhan dokumen. Misalkan Komentar Publik terkait kebijakan “Jatim Sehat”. “Pelayanan di puskesmas sudah lebih baik, tetapi obat sering kosong”. kalimat ini menunjukkan terdapat dua entitas dari Kebijakan Publik Jatim Sehat yaitu terkait aspek pelayanan dan aspek ketersediaan obat.

### 2.2.4 *Latent Dirichlet Allocation (LDA)*

LDA adalah model umum probabilitas dari sebuah korpus. Tugas utama dari LDA adalah menemukan variabel yang tersembunyi (*latent*) yang ada di dalam sebuah dokumen teks yaitu topik. Topik dapat digambarkan dari distribusi dari kata-kata yang ada dalam korpus. Kata dengan probabilitas yang tinggi, sangat besar kemungkinannya untuk menjadi sebuah topik. (Blei dkk, 2003). Pengelompokkan data dalam jumlah besar ke dalam topik-topik tertentu akan mempermudah proses penggalian informasi dari data tersebut. Misalkan untuk mendapatkan topik utama dari sebuah dokumen atau menggali makna dari dokumen tersebut (Blei, 2012).

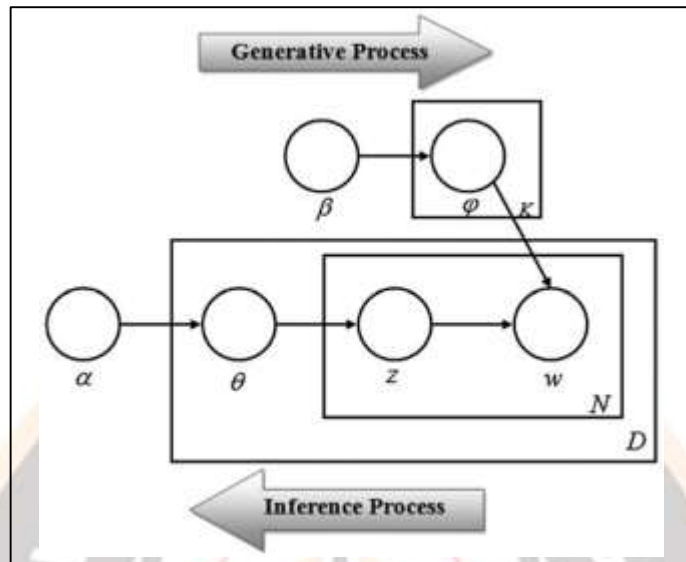
LDA merupakan salah satu bentuk probabilitas pemodelan topik (*probabilistic topic model*). Pada probabilitas pemodelan topik, suatu dokumen dapat menghasilkan beberapa topik. Topik tersebut terbentuk berdasarkan keterkaitan dan probabilitas antar kata-kata yang ada dalam sebuah dokumen tersebut (Liu, 2013). Secara umum topologi jaringan LDA seperti ditunjukkan pada Gambar 2.3.



Gambar 2.3 Topologi Jaringan LDA (Liu, 2013)

LDA memiliki dua bentuk, yaitu LDA dengan proses generatif (*generative process*) dan LDA dengan proses inferensi (*inference process*). Pada LDA dengan model generatif, tujuannya adalah membentuk suatu dokumen berdasarkan nilai distribusi probabilitas yang telah diketahui sebelumnya. Sedangkan LDA dengan proses inferensi merupakan kebalikan dari proses generatif. LDA dengan model inferensi bertujuan untuk mengidentifikasi variabel laten dengan mendapatkan nilai distribusi probabilitas kata  $w$  untuk masing-masing topik  $k$  ( $\phi_k$ ) dan nilai proporsi topik  $k$  untuk setiap dokumen  $d$  ( $\theta_d$ ) dengan menambahkan hiperparameter  $\alpha$  dan  $\beta$ . Proses umum pada LDA inferensi seperti ditunjukkan pada Gambar 2.4.





Gambar 2.4 *Latent Dirichlet Allocation* (Kusumaningrum dkk, 2014)

### 2.2.5 *Colapsed Gibbs Sampling* (CGS)

Pada penelitian ini, model LDA yang digunakan adalah LDA dengan proses inferensi. LDA dengan proses inferensi memiliki dua proses utama, yaitu pelatihan dan pengujian. Proses pertama adalah pelatihan, proses ini bertujuan mendapatkan nilai  $\phi_k$  atau dapat disebut PWZ dan nilai  $\theta_d$  atau dapat disebut PZD dengan memberikan masukan berupa hiperparameter  $\alpha$ ,  $\beta$  dan topik. Untuk memperoleh nilai PWZ dan PZD, semua kata pada korpus terlebih dahulu harus didistribusikan ke dalam topik-topik berdasarkan keterkaitan dan probabilitas antar kata-kata yang ada dalam korpus (Kusumaningrum dkk, 2014).

Untuk memperoleh distribusi probabilitas antar kata tersebut digunakan metode LDA dengan model *Colapsed Gibbs Sampling* (CGS) (Griffiths dan Steyvers, 2004). Metode ini dipilih karena mudah dalam penggunaan dan implementasinya. Konsep utama pada model CGS adalah pembangkitkan sampel distribusi bersyarat dari masing-masing variabel yang nilainya telah diketahui. Setiap variabel selalu diperbaharui nilainya sesuai dengan nilai probabilitas variabel

lainnya. Nilai probabilitas di hitung menggunakan Persamaan (2.1) (Heinrich, 2009).

$$p(z_i = k | z^{(-i)}, w) = \frac{n_{k,-i}^{(w)} + \beta}{\sum_{w=1}^V (n_{k,-i}^{(w)} + \beta)} * (n_{d,-i}^{(k)} + \alpha) \quad (2.1)$$

Keterangan :

- W : Jumlah kata dalam dokumen korpus  
 K : Jumlah topik  
 k : Topik  
 V : Jumlah kata dalam *vocabulary*  
 $\alpha$  : *Hyperparameter* untuk proporsi topik dari topik k  
 $\beta$  : *Hyperparameter* untuk probabilitas kata-topik dari kata w di dalam *vocabulary*  
 i : Token (d,n), yaitu token untuk kata ke-n pada dokumen ke-d  
 $n_{k,-i}^{(w)}$  : Banyaknya kata w yang masuk dalam topik ke-k, kecuali token ke-i  
 $n_{k,-i}^{(j)}$  : Banyaknya kata dalam topik ke-k, kecuali token ke-i  
 $n_{d,-i}^{(k)}$  : Banyaknya kata pada dokumen ke-d yang ditetapkan sebagai topik k, kecuali token ke-i

Hasil distribusi akhir kata  $w$  untuk tiap-tiap topik  $k$  selanjutnya dapat digunakan untuk membentuk model klasifikasi. Model klasifikasi terdiri dari empat macam model yaitu PZD, PWZ, nilai rata-rata probabilitas topik terhadap dokumen untuk setiap kelas yang telah ditentukan (PZC), dan nilai probabilitas topik (PZ). Nilai PZD ( $\theta_{d,k}$ ) dan PWZ ( $\varphi_{w,k}$ ) diperoleh dengan Persamaan (2.2) dan (2.3) (Heinrich, 2009; Kusumaningrum dkk, 2014).

$$\theta_{d,k} = \frac{n_{d,k}^{(d)} + \alpha}{\sum_{k=1}^K n_{d,k}^{(d)} + \alpha} \quad (2.2)$$

SEKOLAH PASCASARJANA

$$\varphi_{w,k} = \frac{n_k^{(w)} + \beta}{\sum_{w=1}^7 n_k^{(w)} + \beta} \quad (2.3)$$

Model yang ketiga adalah PZC. Nilai PZC dihitung menggunakan rata-rata harmonik (*harmonic mean*). Untuk menghitung rata-rata harmonik digunakan Persamaan (2.4).

$$PZC_{c,k} = \frac{n}{\frac{1}{x_1} + \frac{1}{x_2} + \frac{1}{x_3} + \dots + \frac{1}{x_n}} \quad (2.4)$$

Selanjutnya model yang terakhir adalah PZ. Nilai PZ dihitung menggunakan Persamaan (2.5).

$$PZ_k = \frac{n_k^{(*)}}{\sum_{d=1}^n \frac{n_k^{(*)}}{n_*^{(*)}}} \quad (2.5)$$

Proses kedua pada LDA *inference* adalah pengujian. Proses ini bertujuan melakukan klasifikasi terhadap dokumen uji untuk menemukan opini dari dokumen tersebut. Proses klasifikasi dilakukan dengan mengukur kemiripan pola distribusi distribusi antara Nilai PZC yang diperoleh pada proses pelatihan dengan nilai PZC dari dokumen uji. Nilai PZC dari dokumen uji diperoleh dari Persamaan (2.6).

$$\theta_{d,k} = PZD_{d,k} = p(z = k) * \sum_{i=1}^n \varphi_{w,k}^i \quad (2.6)$$

$\theta_{d,k}$  adalah nilai PZD dokumen latih.  $p(z = k)$  adalah nilai PZ yang diambil dari dokumen latih.  $\sum_{i=1}^n \varphi_{k,w}^i$  adalah Nilai PWZ dokumen uji. Kemiripan pola distribusi diukur menggunakan *Kullback Leibler Diverence* (KLD). Nilai KLD dihitung menggunakan persamaan di bawah ini. Nilai KLD di hitung menggunakan Persamaan (2.7), (2.8) dan (2.9).

$$DPQ = \sum_i p_i \log \left( \frac{p_i}{q_i} \right) \quad (2.7)$$

$$DQP = \sum_i q_i \log \left( \frac{q_i}{p_i} \right) \quad (2.8)$$

$$KLD = \frac{DPQ + DQP}{2} \quad (2.9)$$

Nilai  $p_i$  merupakan nilai PZD dokumen uji dan  $q_i$  merupakan nilai PZC dokumen latih.

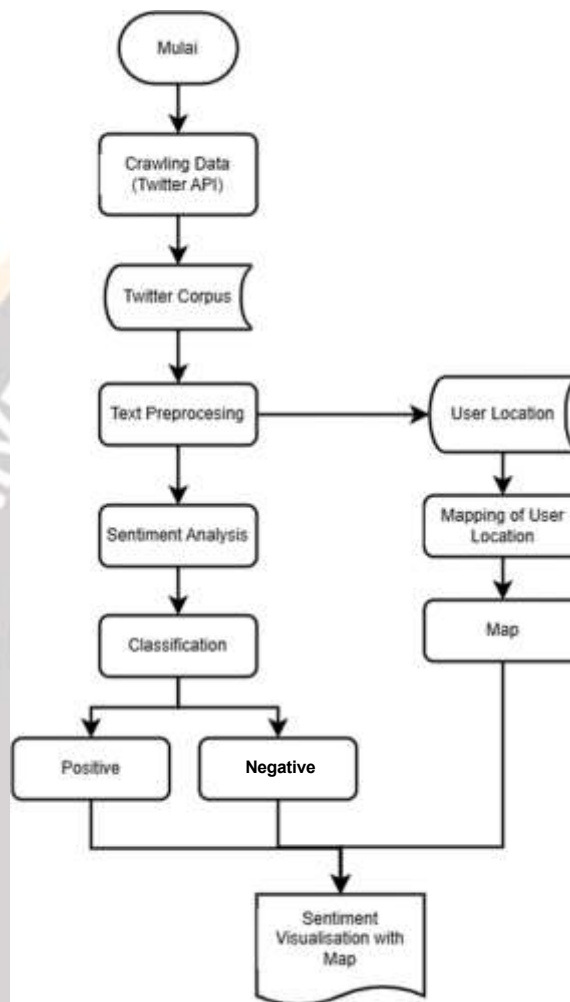
### 2.2.6 Geographic Information System (GIS)

*Geographic Information System* (GIS) atau nama lain dari Sistem Informasi Geografis (SIG) adalah struktur untuk mengumpulkan, mengelola, dan mengevaluasi data, menggunakan ilmu dasar geografi, yang mengintegrasikan banyak jenis teknologi bersama dengan data (Shekhar, S., & Xiong, H., 2007). Selanjutnya informasi koordinat lokasi spasial divisualisasikan menggunakan peta dan pemandangan 3D (Ingle, V., 2020).

GIS mengungkapkan persepsi yang lebih dalam ke dalam data bentuk pola, hubungan, dan situasi yang membantu pengguna untuk membuat keputusan. Peta dapat berguna dalam memahami berbagai struktur sosial. Mereka dapat mengungkap struktur kerumunan dan menyoroti lokasi atau peran strategis dalam jaringan koneksi ini. Dengan memetakan ruang pengguna media sosial, peneliti dapat mempelajari tentang penggunaan paling umum dan terbaik untuk layanan komunikasi ini. Informasi lokasi pengguna X ditemukan menggunakan nama pengguna X. Koordinat diperoleh dengan mencocokkan kota dan negara dengan koordinat yang tepat melalui Google Map API.

Teknologi *Geotagging* digunakan dalam menentukan lokasi pengguna akun sosial media. *Geotagging* merupakan proses akumulasi metadata yang berisi data geografis tentang sesuatu posisi ke dalam peta digital. Data umumnya terdiri dari koordinat lintang serta bujur, namun barangkali juga menyertakan stempel waktu, dan tautan ke data tambahan. Metadata *Geotagging* bisa ditambahkan secara

manual ataupun terprogram. Secara umum alur kerja dari metode penentuan lokasi pengguna akun sosial media X terlihat pada Gambar 2.5.



Gambar 2.5 Alur Kerja Penentuan Lokasi Pengguna X

### 2.2.7 *Term Frequency-Inverse Document Frequency (TF-IDF)*

TF-IDF adalah metode pembobotan sebuah kata/*term* yang memberikan bobot yang berbeda pada setiap *term* dalam dokumen berdasarkan frekuensi *term* per dokumen dan frekuensi *term* pada seluruh dokumen (Cahyani & Patasik, 2021).

Pada proses TF-IDF ini dilakukan pembobotan pada kata atau pengubahan data teks menjadi representasi numerik.



*Term Frequency* (TF) adalah mengukur seberapa sering suatu kata muncul dalam sebuah dokumen. Pendekatan umum untuk menghitung TF adalah dengan menghitung jumlah kemunculan kata tersebut dibagi dengan jumlah total kata dalam dokumen. Dalam beberapa kasus, TF dapat diubah dengan menerapkan skema penimbangan yang lebih kompleks. Sedangkan *Inverse Document Frequency* (IDF) adalah mengukur seberapa penting suatu kata dalam konteks koleksi dokumen yang lebih besar. Kata-kata yang muncul lebih jarang di seluruh koleksi dokumen cenderung memiliki IDF yang lebih tinggi. IDF dihitung dengan membagi jumlah total dokumen dalam koleksi dengan jumlah dokumen yang mengandung kata tersebut. Hasilnya kemudian diambil logaritma untuk memperhalus skala.

Dalam metode TF-IDF, nilai TF dan IDF dikalikan bersama-sama untuk menghasilkan bobot kata (*term weight*) untuk setiap kata dalam sebuah dokumen. Bobot ini mencerminkan tingkat pentingnya kata dalam dokumen tersebut dibandingkan dengan koleksi dokumen yang lebih besar. Adanya proses TF-IDF ini karena algoritma klasifikasi hanya dapat memproses data numerik. Persamaan 2.10 merupakan formula TF-IDF dengan penyesuaian logaritmik untuk menghindari nilai negatif atau nol saat  $df_t = 1$ . Nilai  $tdf_t$  menunjukkan frekuensi kemunculan *term t* pada dokumen  $d$ , sementara  $\log(N/df_t) + 1$  merupakan bagian IDF yang memberikan bobot lebih tinggi pada kata-kata yang jarang muncul di seluruh koleksi dokumen. Penambahan konstanta 1 pada komponen IDF bertujuan untuk *smoothing*, agar kata yang muncul di seluruh dokumen tidak menghasilkan bobot nol, dan tetap dapat diproses dalam model klasifikasi yang hanya menerima data numerik.

$$wtd = tftd \times (\log(N/df_t) + 1) \quad (2.10)$$

Keterangan :

$d$  : dokumen ke- $d$

$t$  : kata ke- $t$  dari kata kunci

$w$  : bobot  $d$  sampai  $d$  terhadap kata  $t$

$tf$  : jumlah kata yang dicari pada sebuah dokumen

IDF : *Inverse Document Frequency*

$N$  : jumlah total dokumen

$df$  : jumlah dokumen yang berisi token

### 2.2.8 Confusion Matrix

*Confusion matrix* adalah alat penting dalam evaluasi model klasifikasi dalam pembelajaran mesin. Tabel *confusion matrix* tersebut menggambarkan data yang benar (*correct*)/ data yang salah (*incorrect*) dari hasil klasifikasi. Elemen diagonal matriks mewakili prediksi yang benar, sementara elemen di luar diagonal menunjukkan kesalahan klasifikasi. Baris dan kolom dari *confusion matrix* berisi data terkait dengan informasi kelas yang sebenarnya (*actual classes*) dan kelas yang diprediksi (*predicted classes*). Dengan menganalisis nilai-nilai seperti *true positive* (TP), *false positive* (FP), *true negative* (TN), dan *false negative* (FN) dalam matriks, praktisi dapat memperoleh wawasan mendalam tentang performa model, termasuk kekuatan dan kelemahannya (Nurdin dkk, 2024). Tabel *confusion matrix* seperti ditunjukkan pada Table 2.1. (Manning dkk, 2008).

Tabel 2.1 *Confusion matrix* dengan dua buah kelas

|                                             |         | Kelas Prediksi ( <i>Predicted Class</i> ) |         |
|---------------------------------------------|---------|-------------------------------------------|---------|
|                                             |         | Negatif                                   | Positif |
| Kelas Sebenarnya<br>( <i>Actual Class</i> ) | Negatif | TP                                        | FN      |
|                                             | Positif | FP                                        | TN      |

*True Positive* (TP) dan *True Negatif* (TN) adalah jumlah kelas data yang benar hasil klasifikasinya. Nilai *False Positive* (FP) adalah nilai dimana hasilnya diprediksi sebagai kelas positif namun sebenarnya merupakan kelas negatif

sedangkan *False Negatif* (FN) adalah nilai dimana nilai dimana hasilnya diprediksi sebagai kelas negatif namun faktanya termasuk dalam kelas positif.

Kinerja model klasifikasi dengan pendekatan *confusion matrix* diukur dari nilai akurasi (*accuracy*). Akurasi digunakan untuk mengetahui proporsi data yang terklasifikasi dengan benar. Nilai akurasi diperoleh dengan menggunakan Persamaan (2.10). Selain dengan pendekatan *confusion matrix*, Kinerja model dapat pula dinilai dengan pendekatan *Receiver Operating Characteristic* (ROC). Ukuran yang digunakan dalam ROC adalah nilai *True Positive Rate* (TPR) atau dapat disebut *sensitivity*, *False Positive Rate* (FPR), dan *Area Under the ROC Curve* (AUC). Nilai TPR, FPR dan AUC diperoleh dengan menggunakan Persamaan (2.11), (2.12), dan (2.13)

- a) *Accuracy*, yaitu nilai yang menunjukkan tingkat akurasi system dalam mengklasifikasikan secara benar (Sholawati dkk., 2022). Berikut rumus perhitungan *Accuracy*:

$$\text{Akurasi} = \left( \frac{TP+T}{TP+TN+FP+F} \right) \times 100\% \quad (2.11)$$

- b) *Recall (Sensitivity)*, yaitu perbandingan jumlah data yang mungkin dikenali dengan jumlah seluruh data yang dikenali. (Sholawati dkk., 2022). Berikut rumus perhitungan *recall*:

***True Positive Rate (TPR)***

$$\text{TPR} = \left( \frac{TP}{TP+F} \right) \times 100\% \quad (2.12)$$

***False Positive Rate (FPR)***

$$\text{FPR} = \left( \frac{FP}{FP+TN} \right) \times 100\% \quad (2.13)$$

***True Negatif Rate (TNR)***

$$\text{TNR} = \left( \frac{TN}{TN+FP} \right) \times 100\% \quad (2.14)$$

- c) *Precision*, yaitu perbandingan jumlah data yang mungkin dikenali dengan jumlah data yang dikenali (Sholawati dkk., 2022). Berikut rumus perhitungan

*Precision*:

$$= \left( \frac{TP}{TP+F} \right) \times 100\% \quad (2.15)$$

- d) *F1-Score*, yaitu rata-rata *harmonic* dari nilai *Precision* dan *Recall* (Fadiyah Basar dkk., 2022). Berikut rumus *F1-score*:

$$= 2 \times \left( \frac{Precision \times Recall}{Precision + Recall} \right) \times 100\% \quad (2.16)$$

### 2.2.9 Kebijakan Pemerintah Jawa Timur

Peningkatan kesejahteraan masyarakat masih menjadi fokus utama dalam Pembangunan Provinsi Jatim di 2023. Khususnya terkait penanganan kemiskinan ekstrem, Pendidikan dan pengembangan SDM serta penurunan angka stunting. Jawa Timur memiliki program unggulan yang dinamakan dengan Nawa Bhakti Satya. Nawa Bhakti Satya memiliki 9 program unggulan yaitu Jatim Sejahtera, Jatim Kerja, Jatim Cerdas dan Sehat, Jatim Akses, Jatim Berkah, Jatim Agro, Jatim Bedaya, Jatim Amanah dan Jatim Harmoni.

#### 1. Jatim Sejahtera:

Mengentaskan Kemiskinan Menuju Keadilan dan Kesejahteraan Sosial.

- a. PKH Plus untuk penduduk miskin 38 kab/kota, 664 kec, 5674 desa dan 2827 kelurahan. Disabilitas, lansia terlantar, perempuan kepala keluarga rentan.
- b. Mengurangi beban 26 Penyandang Masalah Kesejahteraan Sosial (PMKS) dengan subsidi provinsi anggaran ini akan meningkat mengikuti peningkatan pendapatan APBD provinsi.

#### 2. Jatim Kerja :

- a. *Millineal Job Center* dengan cara memberikan *job training*, Pendidikan vokasi, membantu *starting up* usaha, membantu promosi bagi usahawan muda, dan membantu pembiayaan usaha pada tahap awal usaha.

- b. *Dream Team Science Techno Park STP* (5-10 anak SMK dan 2-4 anak (D3/S1), membentuk STP bagi kelompok rintisan usaha di berbagai daerah.

**3. Jatim Cerdas dan Sehat :**

- a. Tis-Tas (gratis dan berkualitas) dengan memperluas cakupan bantuan siswa miskin, bantuan biaya sekolah, dana insentif operasional akreditasi, tunjangan kinerja bagi guru tidak tetap.
- b. *Program Desa Sehat* untuk memperkuat layanan kesehatan pedesaan
- c. Memberikan akses pendidikan berbasis Pesantren bagi anak petani, anak nelayan, anak buruh, anak yatim dan anak yatim piatu yang kurang mampu.

**4. Jatim Akses :**

Membangun infrastruktur dalam Kerangka Pengembangan Wilayah Terpadu, dan Keadilan akses Bagi Masyarakat Pesisir dan Desa Terluar

- a. Kawasan Lingkar Wilis, Lingkar Bromo, Lingkar Ijen, Gerbang Kertasusila, Koridor Maritim Pantura Jawa-Madura, Koridor Maritim Selatan Jawa
- b. Dermaga perintis pulau-pulau Sumenep, penguatan layanan transportasi laut pulau Bawean, pengembangan pesisir selatan.

**5. Jatim Berkah :**

Membangun Karakter Masyarakat yang Berbasis Nilai-Nilai Kesalehan Sosial, Budi pekerti Luhur.

- a. Memberi Tunjangan kehormatan Imam Masjid di Kampung, Pesisir dan Pulau Terluar.
- b. Perluasan tunjangan hafidz dan hafidzah
- c. Penguatan peran Pondok Pesantren dalam mendorong partisipasi sekolah dan beasiswa guru diniyah S2

**6. Jatim Agro :**

Memajukan sektor pertanian, peternakan, perikanan darat dan laut, kehutanan, perkebunan untuk mewujudkan kesejahteraan petani dan nelayan



- a. Memperkuat program petik, olah, kemas, jual, mengembangkan Agropolitan, stabilisasi dan tabungan pangan, asuransi petani, restrukturisasi produk pertanian.
- b. Menjadikan sungai dan hutan sebagai sumber kehidupan dan penguatan SDM pertanian dan GAPOKTAN

**7. Jatim Berdaya :**

Memperkuat ekonomi kerakyatan dengan berbasis UMKM, koperasi dan mendorong pemberdayaan pemerintahan desa

- a. *One Village One Product Corporate & Agropolitan*
- b. *Communal Branding* untuk UMKM
- c. *Supply and Demand Channel*, penataan pasar tradisional, inklusi UMKM retail modern

**8. Jatim Amanah:**

Menyelenggarakan pemerintahan yang bersih, efektif dan anti korupsi.

- a. Membudayakan meritokrasi, menyelenggarakan complain handling system, budaya birokrasi yang melayani dan efektif, menjaga *clean government*, *sound government*, perluasan dan pelayanan berbasis IT.

**9. Jatim Harmoni :**

Menjaga Harmoni sosial dan alam dengan melestarikan Kebudayaan Lingkungan Hidup

- a. Pariwisata partisipatoris, integrasi museum perpustakaan dan galeri seni, ruang kebinekaan, seni tradisional, *clean industries*, *green city*, *halal tourism*.
- b. Dialog antar budaya (seni, seniman dan budayawan)
- c. Dialog Intern & antar umat beragama.

Prestasi paling nyata adalah di pembangunan ekonomi. Pasca Covid 19 kinerja ekonomi Jawa Timur tumbuh 3,57% pada tahun 2021, meningkat dari 2020 minus 2,33 % sedangkan pertumbuhan 2022 naik sebesar 5,34 %.

Angka Tingkat Pengangguran Terbuka (TPT) juga lulusan SMK juga mengalami penurunan. TPT SMK di Jawa Timur pada Agustus 2020 sebesar

11,89 persen turun menjadi 9,54 % pada Agustus 2021. Pada Agustus 2022 turun menjadi 6,70 %



**SEKOLAH PASCASARJANA**