

BAB II

KAJIAN LITERATUR

2.1 Tinjauan Pustaka

2.1.1 Program Makan Bergizi Gratis (MBG)

Program Makan Bergizi Gratis (MBG) merupakan program unggulan pemerintahan Prabowo Subianto - Gibran. Salah satu tujuan MBG adalah untuk meningkatkan kesejahteraan masyarakat, khususnya pelajar. Program MBG merupakan program unggulan pemerintah yang membutuhkan anggaran sangat besar yang masih banyak menuai kritik dan pertanyaan [1]. Program MBG merupakan program unggulan pemerintah yang membutuhkan anggaran sangat besar yang masih banyak menuai kritik dan pertanyaan [2].

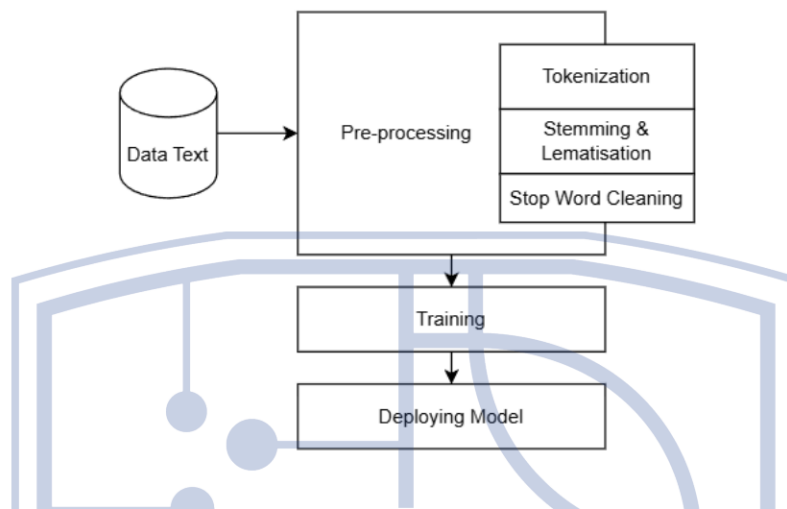
Sampai dengan Oktober 2025, program MBG belum memiliki aturan pelaksana yang menjadi dasar hukum bagi penerapan teknis lapangan. Ombudsman RI menjabarkan terdapat 8 permasalahan utama dalam penyelenggaraan program MBG seperti kesenjangan yang lebar antara target dan realisasi capaian, maraknya kasus keracunan, penetapan mitra yayasan dan SPPG yang rawan konflik kepentingan, keterbatasan dan penataan SDM, mutu bahan baku, standar pengolahan makanan, distribusi makanan yang belum tertib, dan pengawasan yang belum terintegrasi [3]. Pelaksanaan program tersebut memicu berbagai tanggapan masyarakat di media sosial.

Tujuan penelitian ini untuk menghasilkan peta persepsi masyarakat terhadap berbagai aspek dari kebijakan MBG untuk mendukung evaluasi kebijakan publik.

2.1.2 *Natural Language Processing* (NLP)

Natural Language Processing (NLP) atau Pemrosesan Bahasa Alami merupakan cabang dari kecerdasan buatan (*Artificial Intelligence*) yang berfokus pada interaksi antara manusia dan mesin melalui bahasa alami [13]. Perkembangan teknologi dalam beberapa

dekade terakhir telah memungkinkan kemajuan signifikan dalam penelitian dan pengembangan NLP [14].



Gambar 2.1 Cara kerja NLP [27]

Dari gambar tersebut, dapat diketahui cara kerja dari NLP sebagai berikut:

1. *Preprocessing*

Perangkat lunak NLP menggunakan teknik pra-pemrosesan seperti tokenisasi, *stemming*, lematisasi, dan penghapusan kata henti guna menyiapkan data untuk berbagai aplikasi. Berikut penjelasan dari teknik-teknik ini:

- Tokenisasi memecah sebuah kalimat menjadi unit kata atau frasa individual.
- *Stemming* dan lematisasi menyederhanakan kata ke dalam bentuk akarnya. Misalnya, proses ini mengubah "*starting*" menjadi "*start*."
- Dengan penghapusan kata henti, kata yang tidak berkontribusi pada makna penting kalimat, seperti "*for*" dan "*with*" akan dihapus.

2. *Training*

Menggunakan data yang diproses sebelumnya dengan *machine learning* untuk melatih model NLP guna menjalankan aplikasi spesifik berdasarkan informasi tekstual yang

disediakan. Pelatihan algoritma NLP memerlukan pemberian sampel data besar pada perangkat lunak untuk meningkatkan akurasi algoritma.

3. *Deploying*

Model yang ditraining kemudian melakukan *deployment* model atau mengintegrasikan model tersebut ke dalam lingkungan produksi yang sudah ada. Model NLP menerima input dan memprediksi output untuk kasus penggunaan tertentu yang didesain untuk model tersebut.

Secara umum, ruang lingkup NLP dapat dibagi menjadi dua kategori besar, yaitu aspek linguistik (teoretis) dan aspek komputasional (praktis) [15] :

1. Aspek Linguistik (Teoretis)

Mencakup pemahaman terhadap struktur dan makna bahasa, seperti:

- Morfologi: analisis bentuk kata.
- Sintaksis: analisis struktur kalimat.
- Semantik: makna kata dan kalimat.
- Pragmatik: makna yang bergantung pada konteks komunikasi.

2. Aspek Komputasional (Praktis)

Fokus pada penerapan algoritma dan model pembelajaran mesin untuk memproses bahasa, yang meliputi:

- *Tokenization dan Part-of-Speech Tagging*
- *Named Entity Recognition (NER)*
- *Topic Modeling*
- *Sentiment Analysis*
- *Machine Translation*
- *Question Answering dan Chatbot Systems*

Dalam konteks penelitian ini, ruang lingkup NLP difokuskan pada analisis teks komentar publik di media sosial, khususnya melalui:

- a. *Topic Modeling*, untuk menemukan topik atau tema utama dari komentar.
- b. Model IndoBERT, yang digunakan untuk meningkatkan akurasi pemahaman teks berbahasa Indonesia.

Pendekatan tersebut menggambarkan bagaimana NLP tidak hanya berfungsi sebagai alat analisis linguistik, tetapi juga sebagai sarana untuk memahami persepsi publik terhadap kebijakan sosial secara lebih komprehensif.

2.1.3 *Topic Modeling*

Topic modeling merupakan salah satu teknik *unsupervised machine learning* untuk menemukan tema tersembunyi dari kumpulan dokumen teks besar guna mengelompokkan tema-tema tersebut menjadi satu topik [16]. *Topic modeling* memberikan beberapa manfaat, termasuk mengidentifikasi topik-topik penting, otomatisasi pengkodean dan pengukuran fenomena skala besar, serta membantu dalam interpretasi dan validasi hasil [17]. *Topic Modeling* dapat dikatakan efektif dalam membentuk topik-topik yang paling sering dibahas pada media online [18].

2.1.3.1 Konsep Dasar *Topic Modeling*

Topic Modeling bekerja dengan mengidentifikasi pola-pola yang ada dalam dokumen untuk menemukan kumpulan kata-kata yang sering muncul bersama-sama [19]. Dengan menggunakan pemodelan probabilistik untuk mengelompokkan kata-kata pada dokumen ke dalam topik yang berbeda [20].

2.1.3.2 BERTopic

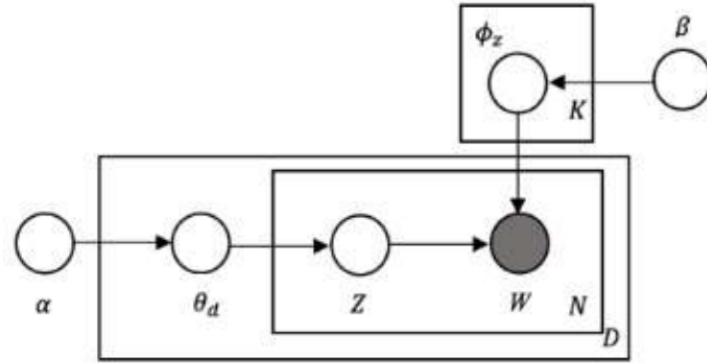
BERTopic diperkenalkan pada tahun 2020 adalah model topik yang menggunakan transformer BERT, teknik *clustering*, dan variasi c-TF-IDF (*class-based TF-IDF*) untuk

menghasilkan representasi topik yang koheren. Model ini dirancang untuk mengatasi keterbatasan model lain yang kurang memperhitungkan hubungan semantik antara kata. BERTopic unggul dalam berbagai benchmark berkat kinerjanya yang dapat diskalakan seiring dengan perkembangan model bahasa baru, fleksibilitasnya dalam memisahkan proses *clustering* dari representasi topik, serta kemampuannya untuk memodelkan evolusi topik melalui distribusi kata [22]. BERTopic menggabungkan embedding BERT, reduksi dimensi UMAP, dan algoritma klasterisasi HDBSCAN untuk mengelompokkan dokumen ke dalam topik yang koheren [31].

BERTopic memiliki 3 langkah utama untuk menghasilkan topik yang telah diolah. Langkah pertama yaitu, *Document Embeddings* yang mana pada tahap ini akan dilakukan konversi document atau biasanya teks menjadi bentuk vektor. Langkah selanjutnya adalah *Document Clustering*, pada tahap ini akan dilakukan *clustering* atau pengelompokkan data berdasarkan jarak terdekat. Langkah terakhir adalah *Topic Representation*, pada tahap ini akan ditentukan sebuah kalimat atau kata sebagai perwakilan kluster yang mana penentuan perwakilan kluster tersebut dilakukan dengan prosedur TF-IDF [32].

2.1.3.3 Latent Dirichlet Allocation (LDA)

Latent Dirichlet Allocation (LDA) merupakan salah satu metode dari pemodelan topik yang banyak digunakan. *Latent* mengacu pada segala sesuatu yang tersembunyi dalam data, *Dirichlet* merujuk pada distribusi probabilitas yang menggambarkan distribusi topik dalam dokumen dan kata-kata dalam topik, dan *Allocation* berarti mengalokasikan topik atau kata-kata menjadi lebih spesifik ke dalam dokumen [19][29]. Model topik LDA menghitung probabilitas distribusi bersyarat (distribusi posterior) dari variabel tersembunyi berdasarkan variabel yang diamati dengan menggunakan distribusi gabungan [22].



Gambar 2.2 Model Representasi Grafis LDA [22]

Model ini bekerja berdasarkan asumsi generatif, di mana dokumen dianggap dihasilkan melalui proses acak berdasarkan dua parameter utama: α (alpha) dan η (eta). Parameter α mengontrol distribusi topik dalam dokumen; nilai α kecil akan menghasilkan dokumen yang cenderung fokus pada sedikit topik. Sementara itu, η mengatur distribusi kata dalam topik; nilai η kecil menghasilkan topik yang lebih spesifik. Nilai-nilai ini dapat diatur secara manual atau ditentukan secara otomatis selama proses pelatihan model [30].

Berdasarkan proses di atas, maka probabilitas dari dataset yang diberikan dapat dirumuskan dengan Persamaan 1.

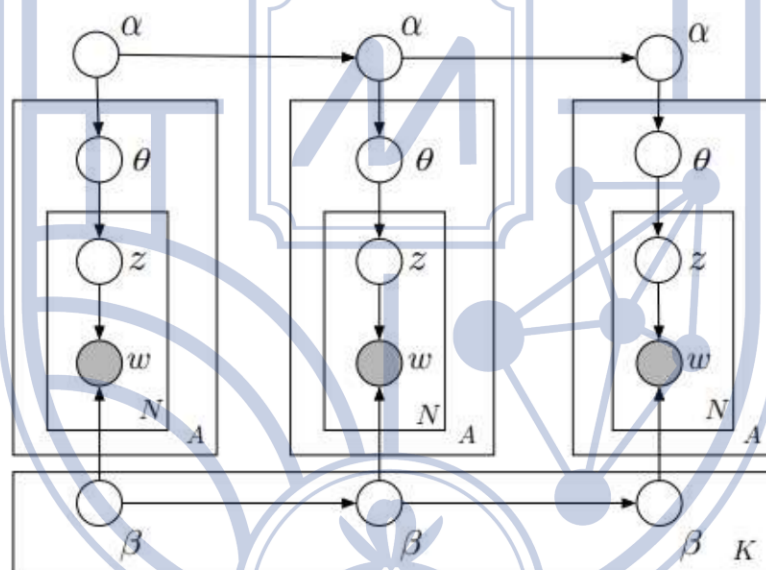
$$p(d, w) = p(d) \sum_z p(z|d) p(w|z) = p(d) \sum_z \theta_d(z) \phi_z(w) \quad (1)$$

Dimana $p(d, w)$ adalah probabilitas LDA dari dataset yang diberikan, yang merupakan gabungan dari probabilitas dokumen d dan kata w dan $p(d)$ probabilitas dari dokumen d . Variabel z merujuk pada indeks topik yang dipilih dari distribusi multinomial untuk dokumen tertentu. $p(z/d)$ adalah probabilitas topik z diberikan dokumen d , yang sama dengan distribusi topik $\theta_d(z)$, sedangkan $p(w/z)$ merupakan probabilitas kata w diberikan topik z , yang sama dengan distribusi kata $\phi_z(w)$ [22].

2.1.3.4 Dynamic Topic Model (DTM)

Dynamic Topic Model (DTM) diperkenalkan oleh Blei dan Lafferty (2006) sebagai pengembangan dari LDA untuk menangkap perubahan topik sepanjang waktu. DTM memandang dokumen sebagai kumpulan yang terbagi ke dalam beberapa *time slices* (misalnya per bulan atau per tahun), di mana setiap time slice memiliki distribusi topik yang saling berhubungan dengan time slice sebelumnya.

Secara formal, DTM memodelkan evolusi distribusi topik menggunakan proses *Gaussian linear dynamical system*, di mana distribusi topik pada waktu t bergantung pada distribusi pada waktu $t-1$. Dengan demikian, topik dapat berubah secara bertahap, mengikuti pola temporal yang diskrit.



Gambar 2.3 Diagram Dynamic Topic Model [11]

Keterangan :

α = Parameter hiper untuk distribusi topik per dokumen

β = Distribusi kata untuk setiap topik (*topic-word distribution*)

θ = Distribusi topik untuk setiap dokumen

z = Variabel laten topik untuk tiap kata

w = Kata yang teramati (*observed word*)

N = Jumlah kata dalam dokumen

A = Jumlah dokumen per waktu

K = Jumlah topik

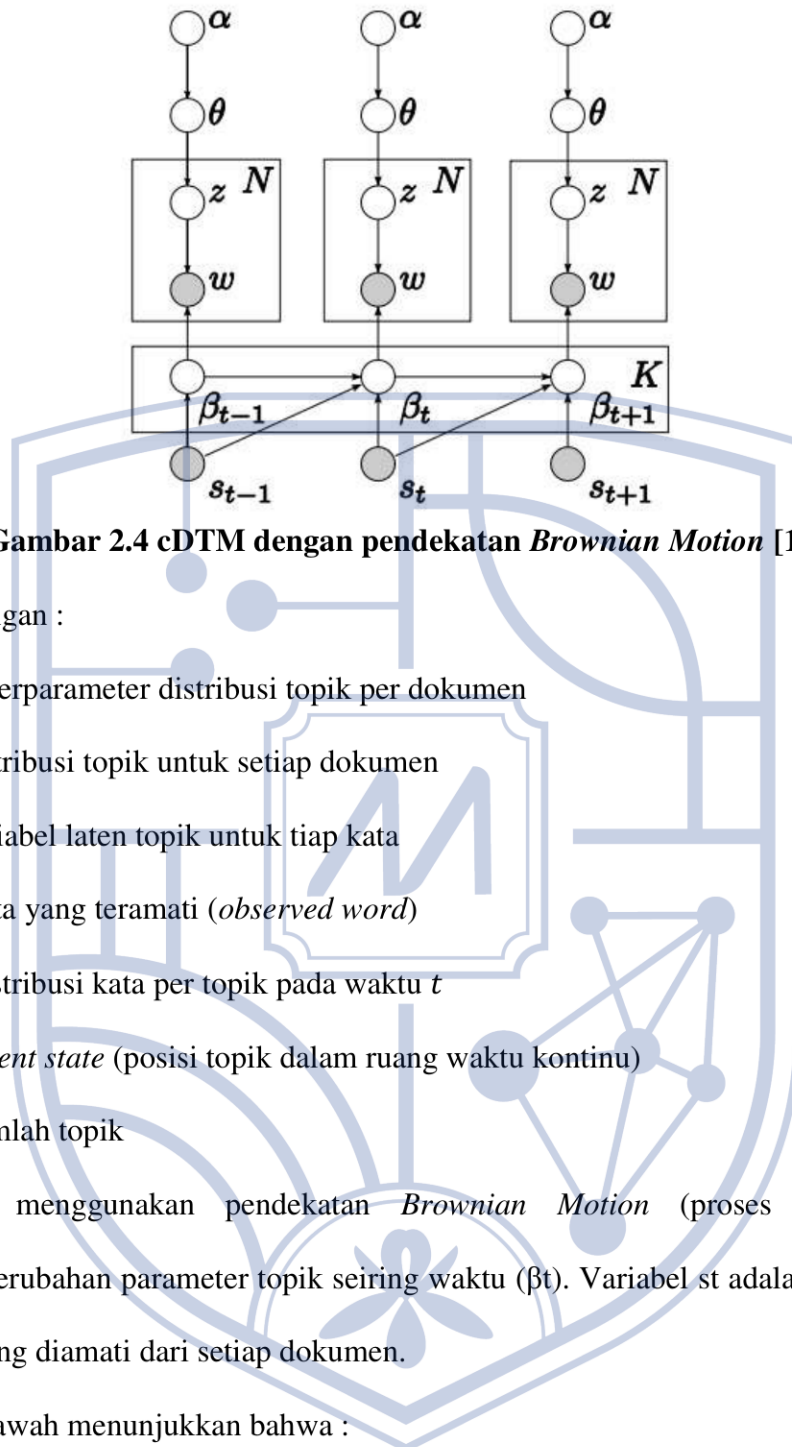
Dari gambar tersebut terdapat arah panah horizontal dari satu kotak ke kotak berikutnya untuk variabel β dan α . Gambar tersebut menunjukkan bahwa :

- β_t bergantung pada β_{t-1} yang artinya distribusi kata untuk topik di waktu t berevolusi dari waktu sebelumnya
- α_t bergantung pada α_{t-1} yang artinya parameter distribusi topik per dokumen juga berubah secara halus dari waktu ke waktu.

Secara intuitif DTM mengasumsikan bahwa topik tidak muncul atau hilang secara tiba-tiba, tapi berubah perlahan mengikuti tren waktu. Namun, karena DTM bekerja berdasarkan pembagian waktu diskrit, model ini memiliki keterbatasan dalam menangani data dengan cap waktu (*timestamp*) yang tidak teratur seperti data media sosial yang memiliki intensitas posting yang fluktuatif setiap hari atau jam.

2.1.3.5 Continuous-Time Dynamic Topic Model (cDTM)

Untuk mengatasi keterbatasan DTM, Wang, Blei, dan Heckerman (2008) memperkenalkan *Continuous-Time Dynamic Topic Model* (cDTM). Model ini mengasumsikan bahwa evolusi topik terjadi secara kontinu terhadap waktu, bukan dalam interval diskrit. Dengan demikian, perubahan topik tidak harus terjadi pada batas waktu tertentu (misalnya per bulan), tetapi dapat diperbarui setiap kali dokumen baru masuk.



Gambar 2.4 cDTM dengan pendekatan *Brownian Motion* [12]

Keterangan :

α = Hiperparameter distribusi topik per dokumen

θ = Distribusi topik untuk setiap dokumen

z = Variabel laten topik untuk tiap kata

w = Kata yang teramati (*observed word*)

β_t = Distribusi kata per topik pada waktu t

st = *Latent state* (posisi topik dalam ruang waktu kontinu)

K = Jumlah topik

cDTM menggunakan pendekatan *Brownian Motion* (proses Wiener) untuk memodelkan perubahan parameter topik seiring waktu (β_t). Variabel st adalah stempel waktu (*timestamp*) yang diamati dari setiap dokumen.

Baris bawah menunjukkan bahwa :

$$\beta_{t-1} \rightarrow \beta_t \rightarrow \beta_{t+1}$$

Gambar 2.5 Hubungan antar waktu cDTM [12]

Setiap β_t bergantung pada keadaan laten st . Variabel st juga mengikuti proses difusi kontinu.

$$st \sim N(st-1, v\Delta_t I)$$

Gambar 2.6 Proses difusi pada variabel st [12]

Di mana:

Δt = jarak waktu antar-dua *event* (misal selisih hari, minggu, dan seterusnya)

v = varians laju evolusi topik

Ini membuat perubahan topik proporsional terhadap jarak waktu nyata, bukan sekedar indeks periode. Setiap topik diwakili oleh distribusi kata yang berubah secara kontinu mengikuti proses stokastik, di mana deviasi kecil dari satu waktu ke waktu berikutnya menggambarkan perubahan alami dalam penggunaan kata.

Model ini lebih realistis untuk data sosial seperti media sosial, berita daring, atau forum publik, di mana topik dapat muncul, berubah, atau menghilang kapan saja tanpa pola waktu tetap. Komponen utama cDTM meliputi:

1. Dokumen dengan cap waktu (*Timestamp*) - Setiap dokumen diberi waktu pembuatan, misalnya tanggal posting di media sosial.
2. Distribusi Topik Dinamis - Distribusi topik tidak tetap, melainkan berevolusi secara kontinu.
3. *Inference* dengan *Variational Kalman Filtering* - cDTM menggunakan metode inferensi *Variational Kalman Filtering* untuk memperkirakan parameter topik yang berubah sepanjang waktu.

Dengan pendekatan ini, model mampu menangkap perubahan gradual maupun cepat dalam tema diskusi, menghubungkan dokumen yang berjauhan secara waktu tetapi masih memiliki kemiripan semantik dan menunjukkan kapan topik baru muncul dan kapan topik lama menurun intensitasnya.

Hasil utama dari penerapan cDTM adalah urutan distribusi kata yang berubah terhadap waktu, memperlihatkan bagaimana topik berevolusi, grafik yang menunjukkan kapan topik tertentu meningkat atau menurun dan korelasi dengan peristiwa nyata dengan mengaitkan pola perubahan topik dengan peristiwa eksternal untuk mengetahui hubungan sebab-akibat atau dampak sosial dari suatu isu.

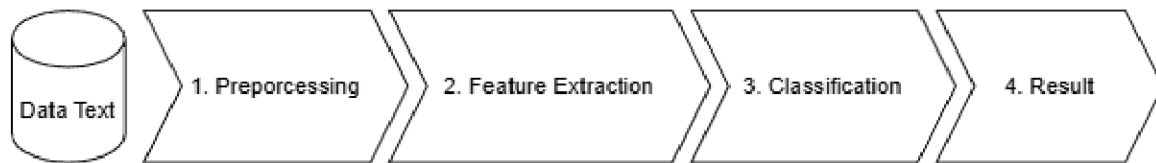
Model topik statis seperti *Latent Dirichlet Allocation* (LDA) ataupun BERTopic mengasumsikan bahwa distribusi topik bersifat tetap sepanjang waktu, sehingga tidak mampu menangkap perubahan fokus pembahasan yang terjadi secara dinamis. *Dynamic Topic Model* (DTM) konvensional memperkenalkan dimensi waktu dengan membagi data ke dalam interval diskrit, seperti bulanan atau tahunan. Namun, pendekatan ini memiliki keterbatasan karena perubahan topik diasumsikan hanya terjadi pada batas interval tersebut.

Continuous-Time Dynamic Topic Model (cDTM) mengatasi keterbatasan tersebut dengan memodelkan evolusi topik secara kontinu berdasarkan cap waktu aktual dokumen. Pendekatan *continuous-time* memungkinkan representasi perubahan topik yang lebih realistis, terutama pada data media sosial yang bersifat *real-time* dan tidak terikat pada interval waktu tetap. Oleh karena itu, cDTM dinilai lebih relevan untuk menganalisis dinamika opini publik terhadap kebijakan publik yang berkembang secara organik di media sosial.

Dalam konteks penelitian ini, cDTM menjadi relevan karena memungkinkan pemodelan perubahan topik pembahasan masyarakat terhadap kebijakan publik secara kontinu, sejalan dengan karakteristik data media sosial yang dinamis dan tidak terikat pada *interval* waktu diskrit.

2.1.4 Analisis Sentimen

Analisis sentimen merupakan sebuah metode pengambilan informasi pada data berupa teks untuk menilai kecenderungan sebuah opini terhadap suatu topik baik itu opini positif ataupun negatif [23].



Gambar 2.7 Cara kerja analisis sentimen [13]

Dari gambar tersebut dapat diketahui proses analisis sentimen meliputi beberapa tahapan:

1. *Preprocessing*, yaitu pembersihan teks dari elemen noninformasi seperti tanda baca, emoji, atau kata tidak penting.
2. *Feature Extraction*, yaitu proses mengubah teks menjadi representasi numerik menggunakan metode seperti TF-IDF, *Word2Vec*, atau *Transformer-based embeddings*.
3. *Sentiment Classification*, yaitu proses klasifikasi menggunakan algoritma pembelajaran mesin seperti *Naïve Bayes*, SVM, LSTM, atau model *Transformer* seperti BERT.

Dengan model seperti IndoBERT, analisis sentimen pada teks berbahasa Indonesia dapat dilakukan dengan akurasi yang lebih tinggi dibandingkan pendekatan konvensional.

2.1.5 Model IndoBERT

Model IndoBERT adalah hasil dari pengembangan model *pre-trained* BERT pada tahun 2020 yang diberi nama IndoBERT [22] yang merupakan model bahasa yang telah dilatih sebelumnya untuk bahasa Indonesia dengan arsitektur BERT [23]. IndoBERT terbukti efektif dalam memahami makna dan konteks dalam teks berbahasa Indonesia dan juga menunjukkan kinerja yang sangat baik dalam tugas klasifikasi sentimen ulasan dengan menunjukkan akurasi sebesar 92,16%, dengan nilai *precision* 93,16%, *recall* 91,09%, dan *F1-score* 92,11% [24].

Perkembangan metode analisis sentimen menunjukkan pergeseran dari pendekatan machine learning tradisional menuju model berbasis transformer. Untuk menunjukkan posisi IndoBERT sebagai bagian dari *State of the Art*, dilakukan perbandingan dengan beberapa

metode yang umum digunakan dalam penelitian analisis sentimen Bahasa Indonesia sebagaimana disajikan pada tabel 2.1 berikut.

Tabel 2.1 Perbandingan Model Analisis Sentimen Pada Bahasa Indonesia

Model	Arsitektur	Kemampuan Kontekstual	Kelebihan	Keterbatasan	Rata-rata F1 Score
<i>Naïve Bayes</i>	<i>Probabilistic</i>	Tidak	Komputasi sangat cepat, efisien untuk dataset kecil.	<i>Bag-of-words</i> ; mengabaikan struktur kalimat dan konteks.	60–75%
SVM	<i>Linear Classifier</i>	Tidak	Efektif pada ruang dimensi tinggi, dan stabil dengan kernel trick.	Sangat bergantung pada <i>feature engineering</i> (TF-IDF/N-gram).	70–82%
LSTM	<i>RNN</i>	Sebagian (<i>Sequential</i>)	Mampu menangkap dependensi jarak jauh (<i>long-term dependencies</i>).	<i>Vanishing gradient</i> ; tidak bisa diparalelisasi (<i>training</i> lambat).	75-85%
BERT	<i>Transformer</i>	Ya (<i>Bidirectional</i>)	Memahami konteks kata dari dua arah secara simultan.	Membutuhkan sumber daya GPU yang besar (<i>vram</i> tinggi).	85-92%
IndoBERT	<i>Transformer (Pre-trained BI)</i>	Ya (<i>Bidirectional</i>)	Optimasi pada kosakata slang/formal Bahasa Indonesia.	Memerlukan <i>fine-tuning</i> spesifik dan data berlabel berkualitas.	90-95%

Sumber: Diolah dari [23] dan [24]

Berdasarkan hasil perbandingan model yang disajikan pada tabel 2.1, terlihat adanya evolusi performa yang signifikan dari model tradisional berbasis probabilitas menuju model berbasis *Transformer*. Model klasik seperti *Naïve Bayes* dan SVM menawarkan keunggulan dalam kecepatan komputasi dan efisiensi pada dataset kecil, namun memiliki keterbatasan mendasar dalam memahami konteks semantik karena sifatnya yang tidak memperhatikan urutan kata (*bag-of-words*). Sebaliknya, penggunaan arsitektur LSTM mulai mampu

menangkap dependensi urutan, meskipun sering terkendala masalah vanishing gradient pada kalimat yang panjang.

Lompatan performa paling drastis ditunjukkan oleh model BERT dan khususnya IndoBERT, yang mencapai estimasi *F1-Score* tertinggi di rentang 90–95%. Hal ini disebabkan oleh kemampuan arsitektur *Transformer* dalam melakukan pemrosesan dua arah (*bidirectional*) secara simultan, sehingga model dapat memahami ambiguitas kata berdasarkan kata-kata di sekitarnya. Penggunaan IndoBERT terbukti paling efektif untuk korpus Bahasa Indonesia karena model ini telah dilatih sebelumnya (*pre-trained*) menggunakan dataset masif bahasa lokal, mencakup variasi bahasa formal hingga *slang*, yang menjadi kelemahan utama pada model mBERT (*multilingual*) maupun model tradisional lainnya [23][24].

2.1.6 *Fine-tuning*

Fine-tuning model yang sudah dilatih sebelumnya terbukti efektif untuk tugas spesifik, terutama dalam NLP [28]. Proses ini menyesuaikan model *pre-trained* agar optimal pada dataset tertentu dengan melatih kembali beberapa lapisan akhir, sementara lapisan awal tetap untuk mempertahankan fitur umum bahasa. Hal ini memungkinkan model mempelajari pola khusus tanpa kehilangan pengetahuan umum selama proses pre-training [28].

2.1.7 *Confusion Matrix*

Confusion matrix adalah metode evaluasi performa model klasifikasi dalam machine learning. Metode ini menyajikan informasi tentang klasifikasi aktual dan prediksi yang dilakukan oleh sistem. Kinerja dievaluasi menggunakan data dalam matriks [28], yang memvisualisasikan perbandingan antara label sebenarnya dan prediksi model, seperti yang ditunjukkan pada Gambar 2.7.

		Actual Values	
		Positive	Negative
Predicted Values	Positive	True Positive (TP)	False Positive (FP)
	Negative	False Negative (FN)	True Negative (TN)

Gambar 2.8 Confusion Matrix [28]

Struktur dari confusion matrix umumnya sebagai berikut:

1. *True Positive* (TP) : Klasifikasi dengan hasil positif pada data asli yang positif.
2. *True Negative* (TN) : Klasifikasi dengan hasil negatif pada data asli yang negatif.
3. *False Positif* (FP) : Klasifikasi dengan hasil positif pada data asli yang negatif.
4. *False Negatif* (FN) : Klasifikasi dengan hasil negatif pada data asli yang positif.

Hasil perbandingan dari kelas prediksi yang dihasilkan oleh algoritma dibandingkan dengan data aktual tersebut kemudian digunakan untuk menghitung nilai metrik evaluasi yaitu *accuracy*, *precision*, *recall*, dan *f1-score* untuk mengukur kinerja model klasifikasi.

Accuracy adalah metrik yang mengukur proporsi dari total jumlah prediksi yang benar terhadap keseluruhan data yang diprediksi.

$$accuracy = \frac{TP+TN}{TP+FP+TN+FN}$$

Gambar 2.9 Accuracy Evaluation Metric [28]

Precision adalah metrik yang menunjukkan proporsi hasil positif yang diprediksi oleh model yang benar-benar positif.

$$precision = \frac{TP}{TP+FP}$$

Gambar 2.10 Precision Evaluation Metric [28]

Recall adalah metrik yang mengukur proporsi dari semua kasus positif aktual yang berhasil diidentifikasi dengan benar oleh model sebagai positif.

$$recall = \frac{TP}{TP+FN}$$

Gambar 2.11 Recall Evaluation Metric [28]

F1-Score adalah *harmonic mean* dari *precision* dan *recall*, yang memberikan keseimbangan antara kedua metrik tersebut. Metrik ini akan memberikan gambaran yang lebih baik ketika kelas tidak seimbang.

$$F1 - score = 2 \times \frac{(precision \times recall)}{(precision + recall)}$$

Gambar 2.12 F1-Score [28]

2.2 Penelitian Terdahulu

Penelitian terkait dengan topik *sentiment analysis* pada Program Makan Bergizi Gratis telah banyak dilakukan dalam waktu 1 tahun terakhir. Bagian ini membahas sejumlah studi relevan yang menjadi dasar teoretis dan metodologis bagi penelitian ini, serta menunjukkan perbedaan (*novelty*) dengan penelitian-penelitian sebelumnya.

Penelitian oleh Zahra Purwanti (2024) Pemodelan *Text Mining* untuk Analisis Sentimen Terhadap Program Makan Siang Gratis di Media Sosial X Menggunakan Algoritma *Support Vector Machine* (SVM). Hasilnya dari total 606 data uji yang digunakan dalam penelitian ini, algoritma SVM mampu memprediksi sebanyak 497 tanggapan yang dikategorikan sebagai sentimen positif, sementara 37 tanggapan dikategorikan sebagai sentimen negatif. Secara keseluruhan, analisis ini memberikan gambaran yang jelas bahwa

program makan siang gratis mendapatkan respons positif dari mayoritas pengguna media sosial.

Penelitian oleh Assaroh Putriyeki (2025) Analisis Multidimensional Sentimen Masyarakat Terhadap Program Makan Bergizi Gratis Pada Media Sosial X. Hasilnya sentimen negatif mendominasi dengan proporsi 54,89%, sentimen netral berada di posisi kedua dengan persentase 27,82%, dan sentimen positif dengan persentase paling sedikit yakni 17,29%. Dominasi sentimen negatif menunjukkan bahwa cuitan yang dipengaruhi aspek sosial berisi ujaran kebencian yang ditujukan kepada pemerintah sebagai bentuk protes dari masyarakat kepada pembuat dan pelaksana kebijakan.

Penelitian oleh Wardianto (2025) Analisis Sentimen Publik Terhadap Program Makan Bergizi Gratis di Instagram Menggunakan Algoritma *Support Vector Machine* (SVM). Hasilnya menunjukkan bahwa algoritma SVM memiliki kinerja terbaik dengan akurasi mencapai 0,90. Hasil analisis menunjukkan bahwa mayoritas respons masyarakat cenderung bersifat netral, di mana hampir setengah dari komentar atau tanggapan yang dianalisis tidak mencerminkan sikap yang jelas, baik positif maupun negatif, terhadap program tersebut.

Penelitian oleh Rizal Abror Munir (2025) Analisis Sentimen Cuitan Di Media Sosial X Tentang Program Makan Bergizi Gratis Dengan Metode NLP. Hasilnya Metode *clustering K-Means* mampu mengelompokkan *tweet* ke dalam tiga kategori sentimen, yaitu negatif, netral, dan positif. Hasil *clustering* menunjukkan bahwa mayoritas tanggapan masyarakat cenderung bernada negatif terhadap program tersebut.

Penelitian oleh Widya Wahyuni (2025) *Implementation of BERTopic for Topic Modeling Analysis of the Free Nutritious Meal Program Based on YouTube Comments*. Menghasilkan sepuluh topik yang mencerminkan beragam tema yang mewakili perbedaan pendapat masyarakat. Metode ini terbukti efektif dalam mengelompokkan percakapan publik

menjadi topik-topik yang bermakna sehingga dapat digunakan sebagai dasar evaluasi program, penyusunan kebijakan, dan strategi komunikasi publik yang lebih tepat sasaran.

Dari berbagai penelitian tersebut, dapat disimpulkan bahwa:

1. Sebagian besar penelitian analisis sentimen terhadap Program Makan Bergizi Gratis (MBG) masih berfokus pada klasifikasi statis menggunakan algoritma seperti SVM atau pendekatan *clustering*.
2. Pendekatan *Topic Modeling* dengan BERTopic mulai diterapkan untuk menggali isu yang dibahas masyarakat, namun belum diintegrasikan dengan dimensi waktu (temporal).
3. Belum ada penelitian yang menggabungkan Analisis Topik, Analisis Sentimen, dan Analisis Temporal secara simultan, padahal dinamika opini publik di media sosial sangat dipengaruhi oleh waktu dan peristiwa sosial yang terjadi.
4. Sebagian penelitian hanya menyoroti proporsi sentimen (positif, negatif, netral), tanpa mengaitkannya dengan konteks topik atau fase waktu kemunculan isu.

Berbeda dengan penelitian-penelitian sebelumnya yang umumnya menggunakan pendekatan statis atau diskrit, penelitian ini mengadopsi pendekatan *continuous-time* melalui cDTM untuk menangkap evolusi topik dan sentimen secara lebih natural sesuai karakteristik data media sosial. Pendekatan *continuous-time* memungkinkan pemodelan evolusi topik dan sentimen yang lebih realistis dibandingkan pendekatan diskrit, terutama pada data media sosial yang bersifat organik dan tidak terikat interval waktu tetap. Dengan demikian, penelitian ini tidak hanya memperluas kajian analisis sentimen statis, tetapi juga menghadirkan perspektif temporal yang relevan untuk pengambilan keputusan kebijakan berbasis data.

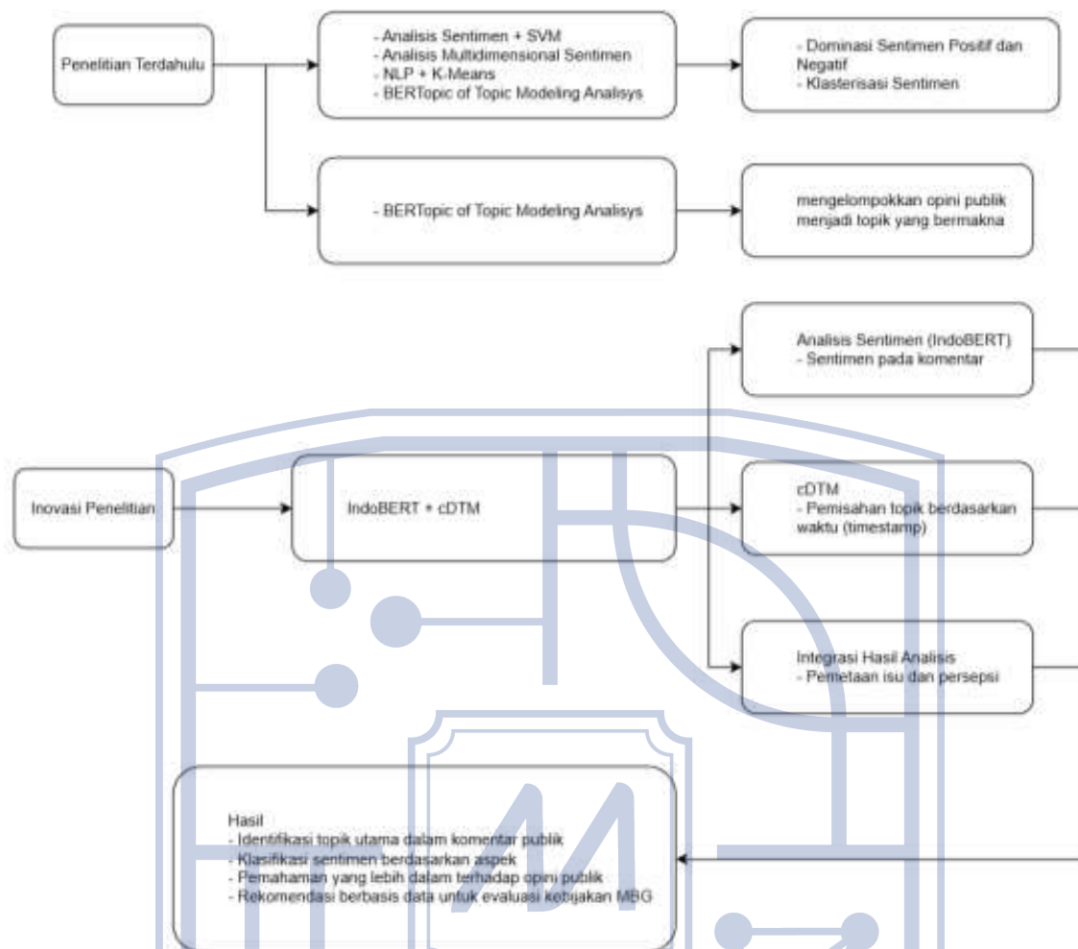
2.3 Kerangka Konseptual

Berdasarkan celah (*research gap*) di atas, penelitian ini mengusulkan pendekatan baru yang menggabungkan kekuatan analisis sentimen berbasis IndoBERT, dan analisis temporal

dengan *Continuous-Time Dynamic Topic Model (cDTM)* dengan kerangka konseptual sebagai berikut:

1. **Pemodelan Topik** : digunakan untuk mengidentifikasi dan mengelompokkan topik utama yang muncul dalam komentar publik mengenai program MBG, menghasilkan representasi tematik yang lebih koheren dan kontekstual.
2. **Analisis Sentimen**: digunakan untuk mengukur opini masyarakat terhadap setiap topik yang ditemukan, dengan klasifikasi sentimen positif, negatif, dan netral.
3. **Integrasi Hasil Analisis**: hasil dari IndoBERT dan cDTM akan digabungkan untuk menghasilkan pemetaan yang menyeluruh tentang isu-isu utama dan persepsi masyarakat terhadap masing-masing aspek pada MBG.
4. **Implikasi Praktis**: hasil penelitian ini dapat menjadi dasar bagi pemerintah dalam mengevaluasi efektivitas komunikasi dan pelaksanaan program MBG berdasarkan data opini publik yang sebenarnya.

Langkah pertama dalam penelitian ini adalah melakukan pengumpulan data komentar publik menggunakan teknik *web scraping*. Data mentah yang diperoleh kemudian diproses melalui tahap pembersihan dan normalisasi teks, yang meliputi *case folding*, *tokenizing*, *stopword removal*, dan *stemming*.



Gambar 2.13 Kerangka Konseptual

Sebagaimana ditunjukkan pada **Gambar 2.12**, kerangka konseptual penelitian ini disusun berdasarkan sintesis teori analisis sentimen berbasis IndoBERT dan pemodelan topik temporal menggunakan cDTM. Kerangka ini menggambarkan alur analisis yang mengintegrasikan hasil klasifikasi sentimen dan evolusi topik untuk memahami dinamika opini publik secara komprehensif. Penelitian terdahulu telah banyak menggunakan berbagai metode seperti analisis sentimen berbasis SVM, analisis multidimensional sentimen, NLP dengan *K-Means*, serta pendekatan *topic modeling* menggunakan BERTopic. Pendekatan-pendekatan tersebut umumnya menghasilkan temuan berupa dominasi sentimen positif dan negatif, klasterisasi sentimen, serta pengelompokan opini publik menjadi topik-topik yang bermakna. Sebagai pengembangan dari penelitian sebelumnya, penelitian ini menghadirkan inovasi melalui kombinasi metode IndoBERT dan *Continuous-Time Dynamic Topic Model* (cDTM).

IndoBERT digunakan untuk melakukan analisis sentimen yang lebih kontekstual dan akurat terhadap komentar publik berbahasa Indonesia, sedangkan cDTM digunakan untuk memisahkan dan melacak evolusi topik berdasarkan dimensi waktu (*timestamp*). Hasil dari kedua analisis ini kemudian diintegrasikan untuk memetakan isu dan persepsi publik secara lebih komprehensif. Melalui pendekatan tersebut, penelitian ini diharapkan dapat menghasilkan identifikasi topik utama dalam komentar publik, klasifikasi sentimen berdasarkan aspek tertentu, pemahaman yang lebih mendalam terhadap opini publik, serta rekomendasi berbasis data yang dapat digunakan untuk evaluasi dan pengambilan keputusan terkait kebijakan Program Makan Bergizi Gratis (MBG).

Dari hasil kombinasi IndoBERT dan cDTM, diperoleh informasi mengenai:

1. Topik utama yang menjadi perhatian publik.
2. Sentimen masyarakat terhadap masing-masing topik.
3. Perubahan distribusi kata untuk tiap topik sepanjang waktu.
4. Distribusi opini publik terhadap kebijakan MBG.

Informasi tersebut kemudian dianalisis untuk menghasilkan interpretasi substantif, yaitu bagaimana persepsi masyarakat terhadap kebijakan MBG dapat menjadi dasar evaluasi bagi pemerintah dalam perbaikan kebijakan di masa mendatang.

Berdasarkan kajian teori dan penelitian terdahulu, penelitian ini menempati posisi yang berbeda dengan mengintegrasikan analisis sentimen berbasis IndoBERT dan pemodelan topik temporal berbasis cDTM dalam konteks kebijakan publik Indonesia, sehingga mampu menangkap dinamika topik dan pergeseran sentimen publik secara simultan dan berkelanjutan.