

Klasifikasi Stance dan Pemodelan Topik Komentar YouTube Terhadap Narasi Risiko Penyakit Menular di Indonesia

Yudha Dwika Sandya*, Indra Budi

Universitas Indonesia, Indonesia

Email: yudha.dwika@ui.ac.id*, indra@cs.ui.ac.id

Kata Kunci	Abstrak
covid; deteksi <i>stance</i> ; hmpv; indobert; pemodelan topik; youtube.	Pasca berakhirnya status darurat kesehatan global COVID-19, komunikasi risiko tetap penting karena potensi kemunculan varian baru dan penyakit menular lain masih terus terjadi. Namun efektivitas komunikasi risiko menghadapi tantangan akibat kelelahan pandemi dan infodemi, sehingga diperlukan pemetaan sikap yang terukur untuk penyesuaian strategi komunikasi di masa pasca pandemi. Penelitian ini memetakan respons publik terhadap narasi risiko penyakit menular COVID-19 dan human metapneumovirus (HMPV) pada komentar YouTube melalui <i>stance detection</i> dan pemodelan topik LDA. Data dikumpulkan dari komentar video bertopik COVID-19/HMPV pada kanal media resmi. Dataset berisi 19.172 komentar dan 9.586 sampel (50%) dianotasi. Penelitian menerapkan klasifikasi dua tahap, yaitu klasifikasi relevansi, diikuti klasifikasi <i>stance</i> pada komentar relevan. Eksperimen membandingkan IndoBERT dan IndoBERTtweet dengan <i>Stratified 5-Fold Cross Validation</i> pada skenario tanpa <i>oversampling</i> dan <i>Random Oversampling</i> (ROS). Hasil menunjukkan IndoBERTtweet memberikan performa terbaik dengan skor <i>F1-macro</i> 0.8541 pada klasifikasi relevansi dan 0.8474 pada klasifikasi <i>stance</i> . Hasil LDA menunjukkan bahwa respons masyarakat pada COVID-19 dan HMPV memperlihatkan pola serupa yang didominasi penolakan. Studi ini menunjukkan analisis <i>stance</i> dan pemodelan topik dengan LDA pada komentar YouTube dapat mendukung formulasi strategi komunikasi risiko, serta mengindikasikan IndoBERTtweet cenderung sesuai dengan karakteristik teks komentar YouTube yang cenderung informal dan bervariasi.
Keywords	Abstract
covid; hmpv; indobert; <i>stance detection</i> ; topic modelling; youtube.	<i>After the end of COVID-19's global health emergency status, risk communication remains essential because new variants and other infectious diseases continue to emerge. However, its effectiveness is challenged by pandemic fatigue and the infodemic, creating a need for measurable mapping of public stances to inform post-pandemic communication strategies. This study maps public responses to infectious-disease risk narratives about COVID-19 and Human Metapneumovirus (HMPV) in YouTube comments using stance detection and LDA topic modeling. Data were collected from comments on COVID-19/HMPV-related videos posted by official media channels. The dataset contains 19,172 comments, with 9,586 samples (50%) manually annotated. We employ a two-stage classification pipeline: relevance classification followed by stance classification on relevant comments. Experiments compare IndoBERT and IndoBERTtweet using Stratified 5-Fold Cross-Validation under two settings: without oversampling and with Random Oversampling (ROS). Results show IndoBERTtweet achieves the best performance, with macro F1 of 0.8541 for relevance and 0.8474 for stance</i>

classification. LDA results indicate similar response patterns for COVID-19 and HMPV, dominated by rejection. These findings suggest that stance detection and LDA topic modelling from YouTube comments can support risk communication strategy formulation, and that IndoBERTweet is more suitable to informal, diverse YouTube text.



PENDAHULUAN

Pengalaman pandemi COVID-19 menunjukkan bahwa penyakit menular tidak hanya menimbulkan krisis kesehatan, tetapi juga berdampak pada stabilitas ekonomi, sosial, dan psikologis masyarakat global. Meskipun status darurat kesehatan global COVID-19 dinyatakan berakhir pada 5 Mei 2023, ancaman penyakit menular tetap perlu diwaspadai karena potensi kemunculan varian COVID-19 serta penyakit menular lainnya yang terus bersirkulasi (Sarker et al., 2023; WHO, 2024). Kewaspadaan relevan dipertahankan mengingat adanya kecenderungan melemahnya persepsi risiko, sebagaimana diindikasikan oleh sejumlah penelitian pada beberapa negara (Delussu et al., 2022; Stechemesser et al., 2023; Tejamaya et al., 2025). Selaras dengan hal tersebut, World Health Organization (WHO) merekomendasikan penguatan komunikasi risiko dan melibatkan komunitas (RCCE) sebagai salah satu pilar utama kesiapsiagaan kesehatan masyarakat, agar masyarakat dapat memahami risiko dan mengadopsi perilaku protektif yang tepat (WHO, 2023).

Mengingat komunikasi risiko berkorelasi positif dengan perilaku protektif (Wang et al., 2020) dan efektivitasnya dipengaruhi oleh pemahaman publik terhadap pesan yang disampaikan (Ong et al., 2023), pemantauan opini dan persepsi masyarakat menjadi krusial sebagai dasar perencanaan dan evaluasi strategi komunikasi risiko. Dalam konteks komunikasi risiko kesehatan, pesan kerap dikemas dalam bentuk narasi yang menggambarkan apa yang terjadi, bagaimana risiko dapat berkembang atau terulang, serta langkah mitigasinya (Berg et al., 2021; Burgess, 2019). Dengan bentuk narasi tersebut, serta selaras dengan rekomendasi WHO yang menekankan pemantauan opini dan persepsi publik (*social listening*) sebagai dasar formulasi strategi komunikasi risiko (WHO, 2025), pemantauan dapat dioperasionalkan melalui analisis respons masyarakat terhadap narasi risiko yang berisi anjuran kewaspadaan terhadap potensi penyakit menular yang beredar di ruang publik digital, termasuk bagaimana pesan kewaspadaan dipahami, diterima, atau diperdebatkan.

Di Indonesia, pengalaman pandemi menunjukkan dampak epidemiologis yang luas dan berkepanjangan, sehingga kebutuhan terhadap peningkatan kewaspadaan dan kesiapsiagaan masyarakat masih diperlukan untuk menghadapi ancaman penyakit menular pasca pandemi (Harapan et al., 2023). Dalam konteks tersebut, komunikasi risiko menjadi komponen kunci untuk menjaga pemahaman risiko dan keberlanjutan perilaku protektif, namun implementasinya menghadapi tantangan berupa kelelahan pandemi (*pandemic fatigue*) yang dapat menurunkan motivasi masyarakat untuk mempertahankan perilaku protektif, serta infodemi (banjir informasi akurat maupun tidak akurat) yang menyulitkan masyarakat membedakan informasi benar dan sumber tepercaya (Kementerian Kesehatan Republik

Indonesia, 2021). Dengan demikian, penguatan komunikasi risiko di Indonesia pasca pandemi perlu mempertimbangkan dinamika respons masyarakat, termasuk bentuk dukungan maupun penolakan terhadap narasi risiko, agar strategi komunikasi yang disusun tetap relevan dan efektif.

Seiring perkembangan teknologi, data dari media sosial memungkinkan pemantauan opini dan persepsi publik secara lebih cepat dibandingkan survei tradisional (Charles-Smith et al., 2015). Sejumlah studi di Indonesia menganalisis respons masyarakat terhadap isu kesehatan berbasis data X (Twitter) melalui pendekatan analisis sentimen, misalnya mengenai COVID-19 (Damanik & Setyohadi, 2021), vaksin COVID-19 (Manuaba, 2023), serta HMPV (Purba & Panjaitan, 2025). Meskipun analisis sentimen berguna untuk memberikan gambaran umum respons masyarakat berupa polaritas emosi (positif dan negatif), komunikasi risiko memerlukan identifikasi sikap masyarakat yang lebih operasional (misalnya dukungan atau penolakan) yang dapat dilakukan dengan pendekatan *stance detection*. Beberapa penelitian terdahulu menunjukkan bahwa *stance detection* dapat digunakan untuk memetakan sikap publik (dukungan, penolakan) terhadap suatu topik dan sering ditekankan sebagai konsep yang berbeda dari analisis sentimen (Burnham, 2023; Glandt et al., 2021).

Berdasarkan penelusuran literatur, studi dengan pendekatan *stance detection* sebelumnya dalam konteks Indonesia maupun di luar Indonesia masih terfokus pada X/Twitter dengan target spesifik misalnya vaksin (Bimantara et al., 2025; Delcea et al., 2022; Purwitasari et al., 2023), kebijakan *lockdown* (Miao et al., 2022), serta penggunaan masker (Cotfas et al., 2021). Sementara kajian yang mengidentifikasi sikap publik (dukungan/penolakan) pada narasi risiko yang berisi anjuran kewaspadaan pada komentar YouTube, khususnya di masa pasca pandemi masih terbatas. Padahal YouTube memiliki basis pengguna yang besar di Indonesia (World Population Review, 2025), serta digunakan untuk komunikasi risiko oleh otoritas kesehatan (Kemenkes, 2021; Shortt et al., 2021). Dengan format konten berbasis video serta fitur interaksi berupa komentar (Gurler & Buyukceran, 2022), YouTube menyediakan data tekstual yang kaya konteks, termasuk respons spontan, diskusi, dan dinamika penerimaan publik. Karena itu, YouTube berpotensi menjadi sumber data penting untuk memahami persepsi publik pada narasi risiko penyakit menular, khususnya pada masa pasca pandemi.

Berdasarkan kesenjangan tersebut, penelitian ini bertujuan menganalisis respons masyarakat terhadap narasi risiko penyakit menular di masa pasca pandemi pada komentar YouTube menggunakan pendekatan *stance detection*. Selain itu, metode *Latent Dirichlet Allocation* (LDA) digunakan untuk mengekstraksi tema percakapan utama pada korpus yang dipilah menurut *stance* dan jenis penyakit, guna memperkaya interpretasi narasi dukungan/penolakan. Pada penelitian ini, korpus komentar difokuskan pada dua isu penyakit menular yang banyak dibahas di YouTube pada tahun 2025 yaitu COVID-19 dan *Human Metapneumovirus* (HMPV). Pemilihan kedua penyakit tersebut dimaksudkan agar memungkinkan analisis perbandingan narasi respons masyarakat dalam konteks penyakit yang familiar dan berkepanjangan (COVID-19) serta penyakit yang relatif belum familiar bagi masyarakat (HMPV).

Penelitian *stance detection* terdahulu dalam konteks isu kesehatan menunjukkan bahwa model *Bidirectional Encoder Representations from Transformers* (BERT) umumnya mengungguli pendekatan pembelajaran mesin klasik seperti *Naive Bayes*, *Support Vector*

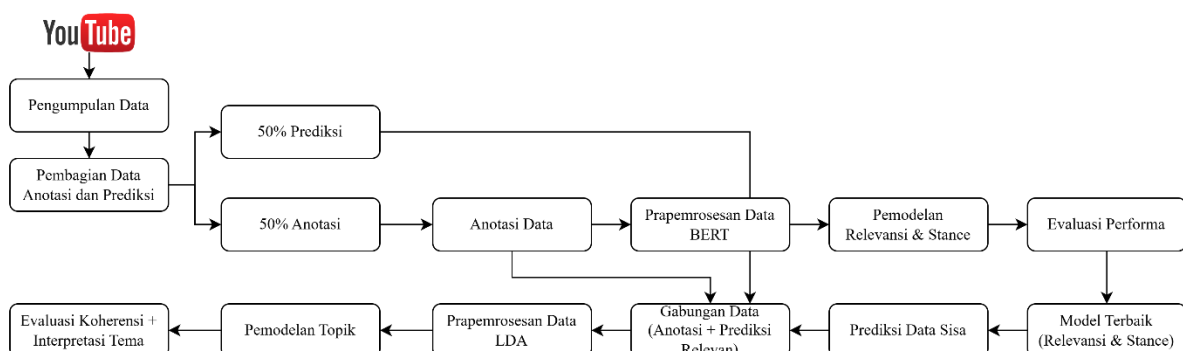
Machine, dan *Random Forest* (Cotfas et al., 2021; Delcea et al., 2022). Penelitian lain juga melaporkan keunggulan model BERT (misalnya DistillBERT) dibanding model pembelajaran mesin klasik seperti *Naive Bayes*, *Logistic Regression*, dan *Support Vector Machine* ketika diterapkan pada tugas *stance detection* di data komentar YouTube (Kang et al., 2024). Selain itu, penelitian (Alrowais, 2024) menekankan pentingnya kesesuaian korpus prapemrosesan model pralatih BERT dengan konteks bahasa penelitian, di mana model pralatih yang dilatih pada bahasa target cenderung memberikan performa lebih baik.

Berangkat dari temuan beberapa penelitian tersebut, penelitian ini akan membandingkan dua model pralatih BERT berbahasa Indonesia, IndoBERT (Wilie et al., 2020) dan IndoBERTweet (Koto et al., 2021) untuk mengklasifikasikan komentar YouTube terkait narasi risiko penyakit menular COVID-19 dan HMPV dengan pendekatan *stance detection*. IndoBERT dipilih karena dilatih pada korpus berbahasa Indonesia yang luas dan mencakup korpus X/Twitter, sedangkan IndoBERTweet dipilih karena prapemrosesannya spesifik menggunakan korpus X/Twitter berbahasa Indonesia yang cenderung informal (misalnya *slang*, singkatan, dan variasi ejaan) dengan topik salah satunya di bidang kesehatan, sehingga relevan untuk diuji pada teks komentar YouTube yang juga memiliki karakteristik yang serupa.

Penelitian ini memberikan kontribusi dalam pemahaman respons publik terhadap narasi risiko penyakit menular melalui *stance detection* dan pemodelan topik tema percakapan pada komentar YouTube terkait COVID-19 dan HMPV yang diharapkan dapat mendukung pemangku kepentingan dalam penguatan *social listening* dan perumusan strategi komunikasi risiko melalui identifikasi pola dukungan, penolakan, dan keraguan publik, termasuk memungkinkan analisis perbedaan pola respons antara konteks penyakit yang lebih familiar (COVID-19) dan yang relatif kurang familiar (HMPV). Dari sisi metodologis, penelitian ini juga memberikan bukti empiris kinerja model pralatih bahasa Indonesia (IndoBERT dan IndoBERTweet) untuk klasifikasi *stance* pada data tekstual komentar YouTube.

METODE PENELITIAN

Berdasarkan tujuan penelitian yang telah didefinisikan, pada Gambar 1 dapat terlihat alur yang akan digunakan dalam penelitian ini.



Gambar 1. Alur Penelitian

Pengumpulan Data

Pengumpulan data dilakukan menggunakan YouTube API v3 melalui pencarian video yang membahas potensi lonjakan/risiko dan anjuran kewaspadaan COVID-19 dan HMPV dengan kata kunci “covid indonesia” dan “hmpv indonesia” pada video yang diunggah pada periode Januari - Oktober 2025. Proses pencarian video dilakukan pada 2 tahap. Tahap pertama dilakukan untuk mengidentifikasi kanal media resmi yang mengunggah video dengan topik pembahasan yang relevan. Kemudian dari hasil tahap pertama akan dilakukan pemilihan kanal media resmi berdasarkan data dari Database Dewan Pers (Dewan Pers, 2025) untuk memastikan konten berasal dari media yang terverifikasi.

Dari hasil kanal video terverifikasi pada tahap pertama, kemudian akan dilakukan pencarian tahap kedua dengan kata kunci dan periode waktu yang sama pada masing-masing kanal agar diperoleh hasil yang lebih komprehensif. Hasil video pada tahap kedua akan diverifikasi secara manual pada judul video dengan kriteria inklusi yang membahas COVID-19 maupun HMPV serta mengecualikan judul yang mengandung kata COVID/HMPV namun tidak relevan dengan kriteria inklusi yang misalnya video dengan judul “Mantan Relawan Covid-19 Gelar Reuni”, “Tanggapi Tuduhan Trump, Tiongkok Duga COVID-19 Berasal dari AS” dan sebagainya. Pada tahap ini juga akan dilakukan deduplikasi video berdasarkan video_id.

Berdasarkan daftar video yang dihasilkan pada Tahap Kedua selanjutnya akan dilakukan pengunduhan komentar pada masing-masing video. Untuk keperluan analisis, penelitian ini hanya akan menggunakan data komentar yang bersifat langsung (*parent comment*) karena mencerminkan respons langsung terhadap narasi yang disampaikan dalam video.

Pemilahan Data Anotasi

Dari korpus komentar langsung yang terkumpul, dilakukan pengambilan sampel untuk anotasi manual sebesar 50% dari data yang telah dilakukan pembersihan awal (deduplikasi dan penghapusan komentar yang hanya terdiri dari emoji atau karakter non alfanumerik), dengan menjaga proporsi per penyakit (COVID-19 vs HMPV). Sisa 50% sampel yang belum dianotasi akan diprediksi dengan menggunakan model yang memiliki performa *F1-macro* terbaik pada eksperimen pemodelan relevansi maupun *stance*.

Pelabelan Data

Dalam *stance detection*, posisi seseorang dalam data tekstual dianalisis terhadap suatu target tertentu (Küçük & Can, 2020). Pada penelitian ini, narasi risiko didefinisikan/dioperasionalisasikan sebagai target *stance* tunggal yaitu “Perlu meningkatkan kewaspadaan terhadap risiko/potensi lonjakan penyakit menular”. Pernyataan ini diterapkan pada komentar dari video bertopik COVID-19 maupun HMPV. Oleh karena itu, kedua penyakit diperlakukan sebagai konteks topik dalam korpus, bukan sebagai target penilaian yang berbeda. Definisi target tunggal ini dimaksudkan menjaga fokus analisis pada penerimaan maupun penolakan kewaspadaan terhadap ancaman penyakit sehingga memungkinkan perbandingan yang konsisten pada seluruh komentar dari kedua penyakit.

Kebijakan atau intervensi spesifik (misalnya vaksinasi, masker, *lockdown*, atau obat/herbal) digunakan sebagai indikator kontekstual untuk menginferensikan posisi seseorang

terhadap target dan tidak diperlakukan sebagai target *stance* terpisah. Dengan demikian, komentar yang membahas intervensi spesifik tetap dinilai sejauh implikasinya terhadap penerimaan/penolakan klaim kewaspadaan risiko.

Pelabelan data dilakukan melalui anotasi manual untuk mengklasifikasikan komentar ke dalam empat kelas berdasarkan sikap terhadap target *stance* yang telah didefinisikan, yaitu *In Favor*, *Against*, *Neutral*, dan *Irrelevant*. Kelas *Irrelevant* digunakan untuk menampung komentar yang tidak relevan dengan klaim kewaspadaan terhadap COVID-19/HMPV (misalnya spam, promosi, percakapan di luar topik). Kriteria operasional masing-masing kelas dan contoh komentar disajikan pada Tabel 1 dan Tabel 2.

Proses anotasi dilakukan secara independen terhadap seluruh data sampel oleh dua anotator (A1 dan A2) untuk meminimalkan bias individual. Perbedaan label antara A1 dan A2 diselesaikan melalui mekanisme adjudikasi oleh anotator ketiga (A3) dengan memilih salah satu label dari A1 atau A2 sesuai pedoman anotasi, tanpa menambahkan label baru. Untuk memastikan kualitas label, reliabilitas anotasi dievaluasi menggunakan *Cohen's Kappa* pada hasil anotasi A1 dan A2 sebelum adjudikasi, sehingga dataset berlabel yang digunakan pada tahap pemodelan memiliki konsistensi yang terukur.

Tabel 1. Kriteria Operasional Anotasi Data

Kelas	Kriteria Operasional
<i>In Favor</i>	Komentar menerima/menguatkan narasi risiko atau mendukung kewaspadaan (termasuk mengusulkan tindakan protektif karena menganggap ancaman relevan).
<i>Against</i>	Komentar menolak/meremehkan narasi risiko (hoax, lebay, berlebihan), atau mengajak tidak perlu waspada. Penolakan secara tidak langsung juga termasuk misal dengan merujuk ke fenomena lainnya untuk menolak/meremehkan narasi risiko.
<i>Neutral</i>	Komentar masih relevan membahas seputar narasi risiko, namun tidak dapat ditentukan sikapnya apakah dia termasuk <i>In Favor</i> atau <i>Against</i> .
<i>Irrelevant</i>	Komentar yang tidak relevan karena tidak membahas narasi risiko termasuk spam dan percakapan di luar topik.

Tabel 2. Contoh Komentar

Kelas	Contoh Komentar
<i>In Favor</i>	yang namanya lagi sakit yaa masa keluar rumah dan tentu aja ke dokter...tp yaa tetap berlaku waspada 😊
<i>Against</i>	Lebay ah, dari dulu ga pernah vaksin dan pake masker alhamdulillah sampe sekarang masih bisa hidup dan komentar ke video ini.
<i>Neutral</i>	Jadi apakah puskesmas udh menyediakan alat test gratis kayak waktu pandemi dulu untuk HMPV ini??
<i>Irrelevant</i>	Maaf saya komen soal siaran di tv digital. Kenapa harus banyak mengiklankan acaranya sendiri? Lebih bagus dan bermanfaat kalau durasi iklan buat menayangkan acara pokok. Juga nggak membosankan penonton.

Prapemrosesan Data

Tahap prapemrosesan data dilakukan untuk membersihkan dan menyiapkan teks komentar sebelum dianalisis lebih lanjut. Proses ini mencakup lima tahapan antara lain:

1. *Cleaning*

Data media sosial umumnya mengandung derau (misalnya tanda baca, angka, simbol, *username*, dan URL) yang dapat menurunkan kinerja klasifikasi. Oleh karena

itu elemen-elemen tersebut dihapus pada tahap *cleaning* (Kumar et al., 2025). Pada tahap ini juga dilakukan konversi emoji menjadi teks serta konversi teks menjadi huruf kecil (*lowercasing*) untuk menyeragamkan token.

2. *Normalization*

Normalization adalah konversi kata dari bentuk tidak baku menjadi baku untuk menjaga konsistensi makna teks.

3. *Stopwords Removal*

Proses penghapusan *stopwords* dilakukan untuk menghilangkan kata-kata yang sering muncul namun tidak memiliki pengaruh yang signifikan terhadap kalimat. Menghapus jenis kata ini dapat membantu mengurangi dimensi dan biaya komputasi sehingga meningkatkan kinerja model secara keseluruhan (Duong & Nguyen-Thi, 2021).

4. *Stemming*

Proses ini mengubah kata ke bentuk dasarnya, sehingga meminimalkan kompleksitas istilah yang ditokenisasi (Yulita et al., 2023). Pada tahap ini, seluruh afiks, termasuk prefiks, infiks, dan sufiks, dihilangkan dari kata.

5. *Tokenization*

Merupakan proses memecah kalimat menjadi kumpulan *token*. Proses ini bertujuan untuk mengeksplorasi dan mengidentifikasi kata kunci dalam sebuah kalimat.

Kelima tahapan di atas dilakukan untuk pemodelan topik menggunakan LDA. Untuk klasifikasi dengan model prelatih BERT, prapemrosesan data hanya mencakup *cleansing* untuk mempertahankan informasi leksikal dan konteks yang dibutuhkan oleh *tokenizer* dan representasi *subword* pada BERT.

Stance Detection

Untuk menangkap dinamika penerimaan dan penolakan publik terhadap narasi risiko, penelitian ini memodelkan komentar sebagai respons sikap terhadap target *stance* tunggal yaitu “Perlu meningkatkan kewaspadaan terhadap risiko/potensi lonjakan penyakit menular” yang secara konsisten diterapkan pada komentar dari video bertopik COVID-19 maupun HMPV. Dengan target tunggal ini, target tidak dimasukkan sebagai *input* model klasifikasi (model dilatih hanya dari teks komentar), sementara target tetap menjadi acuan anotasi dan evaluasi sehingga label *stance* selalu ditafsirkan relatif pada target yang sama.

Penelitian ini menerapkan skema klasifikasi dua tahap. Pada tahap pertama, dilakukan klasifikasi relevansi untuk memisahkan komentar *relevant* terhadap target *stance* (berisi *in favor*, *against*, dan *neutral*) serta *irrelevant*. Pada tahap kedua, hanya komentar pada kategori *relevant* yang diproses lebih lanjut menggunakan klasifikasi *stance* pada tiga kelas, yaitu *in favor*, *against*, dan *neutral*. Eksperimen dilakukan dengan membandingkan dua model prelatih BERT, yaitu IndoBERT (Wilie et al., 2020) dan IndoBERTweet (Koto et al., 2021). Pelatihan dilakukan dengan melakukan pencarian kombinasi hiperparameter pada *batch size* 32, rentang *epoch* 3-4, dan *learning rate* pada $3e-5$ dan $5e-5$ sesuai rentang hiperparameter yang direkomendasikan oleh (Devlin et al., 2018). Mengingat data yang digunakan dalam penelitian ini memiliki ketidakseimbangan kelas, akan dilakukan eksperimen pemodelan klasifikasi

dengan menggunakan skenario tanpa *resampling* dan *Random Oversampling* (ROS) yang diterapkan pada data pelatihan.

Proses permodelan kemudian akan dilakukan dengan menggunakan mekanisme *Stratified 5 K-Fold Cross Validation* untuk memastikan pengukuran performa yang lebih representatif (Cotfas et al., 2021). Dengan pembagian tersebut, pada saat pemodelan, pada setiap lipatan data akan terbagi dengan rasio 80% untuk set pelatihan dan 20% set pengujian sehingga memungkinkan representasi sampel yang cukup memadai jika dibandingkan dengan penggunaan jumlah lipatan sebanyak 10 buah. Jumlah lipatan tersebut dipilih karena jumlah data kelas minoritas terbatas yang disebabkan oleh ketidakseimbangan kelas yang cukup signifikan. Dengan mekanisme *Stratified 5-Fold Cross Validation*, nilai performa yang dilaporkan merupakan rata-rata metrik evaluasi dari lima lipatan pengujian.

Random Oversampling

ROS adalah teknik prapemrosesan data untuk mengatasi ketidakseimbangan kelas dengan menambah jumlah contoh pada kelas minoritas melalui duplikasi acak atas *instance* yang sudah ada (Padurariu & Breaban, 2019). Tujuan utamanya ialah membuat ukuran kelas minoritas menjadi sebanding sehingga model tidak bias terhadap kelas dengan frekuensi tertinggi. Pada penelitian ini ROS dipilih bekerja karena menduplikasi teks di set pelatihan pada kelas minoritas tanpa memperluas rentang pencarian hiperparameter dalam eksperimen.

Evaluasi Performa Model

Pada penelitian ini, *confusion matrix* (Gambar 2) digunakan sebagai pendekatan evaluasi kinerja model. Metriks tersebut digunakan karena memberikan gambaran rinci mengenai kualitas prediksi model klasifikasi (Fahmy Amin, 2023). Berdasarkan *confusion matrix*, dihitung empat metrik evaluasi turunan yaitu *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. *Accuracy* digunakan untuk menggambarkan proporsi prediksi yang benar secara keseluruhan, sedangkan *Precision* dan *Recall* digunakan untuk menilai ketepatan dan kelengkapan prediksi pada tiap kelas. Sementara itu, *F1-Score* digunakan sebagai ukuran ringkas yang menyeimbangkan *Precision* dan *Recall*.

		Predicted	
		Positive	Negative
Actual	Positive	TP	FN
	Negative	FP	TN

Gambar 2. *Confusion Matrix*

Penjelasan *Confusion Matrix*:

- True Positive* (TP) : nilai aktual dan prediksi sama-sama positif
- False Positive* (FP) : nilai aktual negatif, nilai prediksi positif
- True Negative* (TN) : nilai aktual dan prediksi sama-sama negatif
- False Negative* (FN) : nilai aktual positif, nilai prediksi negatif

Pada penelitian ini, penentuan model dan konfigurasi hiperparameter terbaik dilakukan dengan *F1-Macro*, yaitu rata-rata F1 dari seluruh kelas dengan bobot yang sama. Pemilihan agregasi skor secara *macro* dilakukan karena tugas klasifikasi pada penelitian ini memiliki ketidakseimbangan distribusi kelas, sehingga metrik ini lebih representatif dibandingkan *accuracy* yang dapat bias terhadap kelas mayoritas. Selain itu, penggunaan *F1-macro* membantu memastikan performa model tetap diperhatikan pada kelas minoritas, yang relevan untuk menjaga kualitas prediksi pada seluruh kategori label.

Pemodelan Topik Dengan Menggunakan LDA

Pemodelan topik merupakan salah satu bagian dari metode *unsupervised* untuk menemukan topik atau tema yang ada di dalam suatu korpus (Yan et al., 2022). Teknik ini menggunakan pendekatan statistik untuk menemukan relasi antara kata-kata, dokumen, dan topik (Negara et al., 2019). Salah satu metode pemodelan topik yang umum digunakan adalah *Latent Dirichlet Allocation* (LDA) (Sharma et al., 2022). Oleh karena itu, penelitian ini menggunakan LDA untuk mengidentifikasi topik narasi pada komentar YouTube yang membahas narasi risiko COVID-19 dan HMPV, dengan tujuan memetakan pola narasi pada kategori stance *in favor* dan *against* serta kategori penyakit (COVID-19 dan HMPV). Dalam penentuan jumlah topik yang optimal, metode evaluasi yang umum digunakan adalah *Coherence Score* (c_v) (Han et al., 2023). Skor koherensi mengukur kesamaan semantik antara kata-kata dengan skor tinggi dalam suatu topik. Skor yang lebih tinggi menunjukkan bahwa teks lebih mudah dipahami dan terstruktur dengan baik (Negara et al., 2019). Nilai koherensi yang dianggap layak dalam LDA berkisar antara 0,3 hingga 0,6 untuk menemukan kesamaan semantik.

HASIL DAN PEMBAHASAN

Hasil pengumpulan video dan pengunduhan komentar video menghasilkan statistik data sebagaimana ditampilkan pada Tabel 3 di bawah ini. Dari hasil pengumpulan data tahap kedua dihasilkan sebanyak 833 video yang diunggah oleh kanal media resmi. Selanjutnya dari hasil verifikasi pada judul video, diperoleh 387 video relevan (216 COVID-19; 171 HMPV), sedangkan 446 video dikecualikan karena tidak membahas isu penyakit sesuai kriteria. Sesuai desain penelitian, analisis dilakukan pada komentar langsung (*parent comment*) sehingga terkumpul 19.415 komentar. Setelah deduplikasi dan pembersihan (menghapus komentar yang hanya berisi emoji atau karakter non-alfanumerik) tersisa 19.172 komentar (13.006 COVID-19; 6.166 HMPV). Untuk anotasi manual, diambil sampel acak 50% per penyakit sehingga diperoleh 9.586 komentar (6.503 COVID-19; 3.083 HMPV), sementara sisanya akan diprediksi menggunakan model terbaik dan digabung dengan data hasil anotasi untuk membentuk korpus pemodelan topik.

Tabel 3. Hasil Pengumpulan Data

Kategori	Jumlah Video	Total Komentar	Komentar Parent	Komentar Balasan
Dikecualikan	446	Komentar tidak diunduh		
COVID	216	15.809	13.174	2.635
HMPV	171	8.062	6.241	1.821
Total	833	23.871	19.415	4.456

Hasil Pelabelan Data

Tabel 4. Hasil Pelabelan Data

Relevansi	Stance	Total	Persentase	COVID	HMPV
Relevant	In Favor	955	10.0%	497	458
Relevant	Against	7190	75.0%	5.187	2.003
Relevant	Neutral	936	9.8%	489	447
Irrelevant	-	505	5.3%	330	175
Total		9.586	100%	6.503	3.083

Sampel data yang dianotasi secara manual terdiri dari 9.586 komentar, mencakup 6.503 komentar COVID-19 dan 3.083 komentar HMPV. Kualitas anotasi dievaluasi menggunakan *Cohen's Kappa* sebesar 0.733 yang menunjukkan kesepakatan substansial antara anotator 1 dan anotator 2 sehingga sampel data dianggap reliabel untuk digunakan dalam penelitian. Dari hasil anotasi oleh anotator 1 dan anotator2, terdapat 968 komentar yang mengalami perbedaan label dan diselesaikan melalui adjudikasi oleh anotator 3 dengan memilih label dari salah satu anotator sesuai panduan anotasi yang telah diberikan. Berdasarkan Tabel 4, dapat terlihat bahwa distribusi kelas menunjukkan adanya ketidakseimbangan kelas dengan dominasi kelas *against* oleh karena itu dalam eksperimen pemodelan klasifikasi akan menggunakan *Random Oversampling* (ROS) yang diharapkan dapat meningkatkan kinerja pemodelan klasifikasi.

Hasil Pemodelan Relevansi

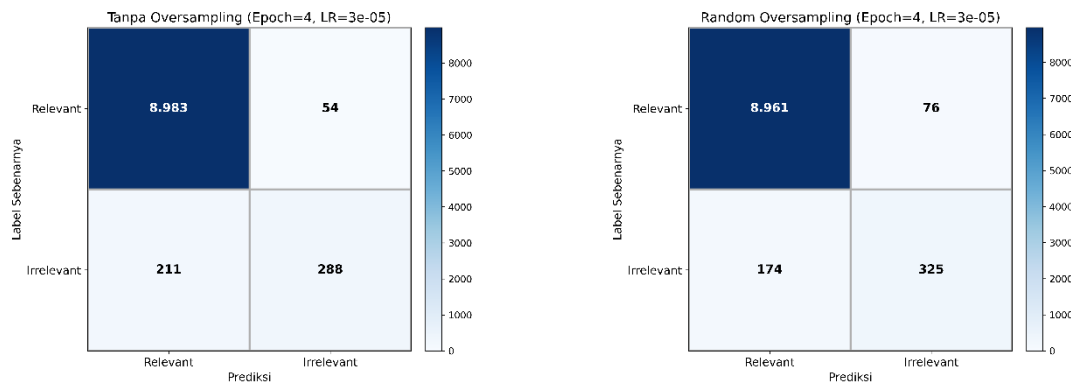
Berdasarkan hasil pengukuran performa pada pemodelan relevansi pada Tabel 5, IndoBERTweet memberikan skor *F1-Macro* tertinggi sebesar 0.8541 pada konfigurasi *epoch* 4, *learning rate* 3e-05, dan penyeimbangan kelas menggunakan *Random Oversampling* (ROS). Kombinasi hiperparameter dan penanganan ketidakseimbangan kelas tersebut unggul dibandingkan konfigurasi terbaik IndoBERT (*F1-Macro* 0.8017).

Tabel 5. Kinerja Terbaik Pemodelan Relevansi

Parameter	Precision	Recall	Accuracy	F1 (Macro)	Sampel
IndoBERTweet					
Epoch 4, LR 5e-05	0.9002	0.7947	0.9716	0.8345	Imbalanced
Epoch 4, LR 3e-05	0.8982	0.8214	0.9738	0.8541	ROS
IndoBERT					
Epoch 4, LR 3e-05	0.8614	0.7623	0.9659	0.8017	Imbalanced
Epoch 4, LR 5e-05	0.8305	0.7701	0.9626	0.7938	ROS

Tabel 6. Perbandingan Dampak ROS pada IndoBERTweet untuk Klasifikasi Relevansi (mean, 5-fold CV)

Model	Kelas	Precision	Recall	F1-Score	Accuracy
IndoBERTweet <i>imbalanced</i>	Relevant	0.9771	0.9940	0.9855	0.9722
	Irrelevant	0.8421	0.5772	0.6849	
Epoch 4, LR 3e-05	Macro	0.9100	0.7856	0.8344	0.9738
IndoBERTweet <i>balanced</i> ROS	Relevant	0.9810	0.9916	0.9862	
	Irrelevant	0.8105	0.6513	0.7222	
Epoch 4, LR 3e-05	Macro	0.8952	0.8214	0.8541	



Gambar 3. Perbandingan Dampak ROS Pada Hiperparameter Terbaik IndoBERTweet Relevansi

Detail per kelas pada model IndoBERTweet dengan kombinasi *epoch* 4 dan LR 3e-05 dengan performa terbaik dapat dilihat pada Tabel 6. Berdasarkan tabel tersebut, penerapan ROS memberikan dampak positif pada sampel data yang mengalami ketidakseimbangan kelas pada penelitian ini. Pada kombinasi hiperparameter tersebut, dengan *oversampling*, terdapat peningkatan *recall* kelas minoritas *irrelevant* dari 0.5772 menjadi 0.6513 (+0.0741 poin), dengan konsekuensi penurunan *precision* dari 0.8421 menjadi 0.8105 (-0.0316 poin). Pola ini konsisten dengan *confusion matrix* yang dapat dilihat pada Gambar 3 dimana lebih banyak komentar *irrelevant* terdeteksi benar, namun terdapat peningkatan kecil pada komentar *relevant* yang diklasifikasikan menjadi *irrelevant*. Secara keseluruhan, ROS memperbaiki keseimbangan deteksi kelas minoritas dan dipilih sebagai konfigurasi untuk prediksi relevansi pada data yang belum berlabel.

Hasil Pemodelan *Stance*

Berdasarkan hasil pengukuran performa pada pemodelan *stance* yang dapat dilihat pada Tabel 7, IndoBERTweet memberikan skor *F1-Macro* tertinggi sebesar 0.8474 pada konfigurasi *epoch* 4, *learning rate* 3e-05, dan penyeimbangan kelas menggunakan *Random Oversampling* (ROS). Kombinasi hiperparameter dan penanganan ketidakseimbangan kelas tersebut unggul dibandingkan konfigurasi terbaik IndoBERT (*F1-macro* 0.7978).

Tabel 7. Kinerja Terbaik Pemodelan *Stance* (mean, 5-fold CV)

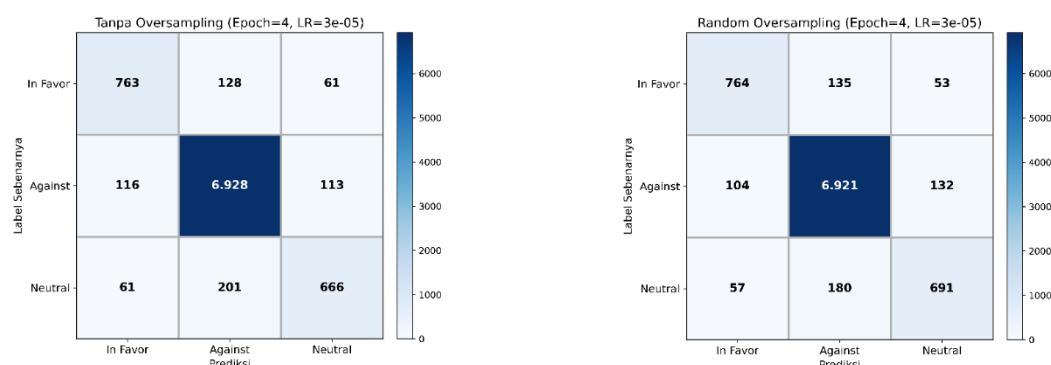
Parameter	<i>Precision</i>	<i>Recall</i>	<i>Accuracy</i>	<i>F1 (Macro)</i>	Sampel
IndoBERTweet					
Epoch 4, LR 5e-05	0.8543	0.8359	0.9262	0.8443	Imbalance
Epoch 4, LR 3e-05	0.8581	0.8381	0.9269	0.8474	ROS
IndoBERT					
Epoch 3, LR 5e-05	0.8147	0.7834	0.9046	0.7978	imbalance
Epoch 3, LR 3e-05	0.7964	0.7581	0.8918	0.7964	ROS

Pada detail per kelas pada konfigurasi IndoBERTweet *epoch* 4, *learning rate* 3e-05 dengan performa terbaik yang dapat dilihat pada Tabel 8, penerapan ROS memberikan dampak positif pada sampel data yang mengalami ketidakseimbangan kelas pada penelitian ini. Pada kombinasi hiperparameter tersebut, dengan *oversampling*, *F1-macro* meningkat dari 0.8405

menjadi 0.8474 dan *recall macro* dari 0.8291 menjadi 0.8381. Peningkatan terutama terlihat pada kelas *Neutral*, dengan *recall* naik dari 0.7177 menjadi 0.7446 dan *F1-macro* naik dari 0.7534 menjadi 0.7661, sementara kelas *Against* relatif stabil pada nilai yang sangat tinggi (F1 sekitar 0.961). Pola ini konsisten dengan *confusion matrix* pada Gambar 4, ROS meningkatkan jumlah prediksi benar pada kelas *Neutral* dan mengurangi sebagian kesalahan klasifikasi terhadap kelas lain, meskipun terdapat *trade-off* kecil berupa perubahan pada beberapa kesalahan antar kelas. Secara keseluruhan, ROS membantu memperbaiki keseimbangan performa antar kelas pada pemodelan *stance* sehingga IndoBERTweet pada kombinasi hiperparameter ini dipilih sebagai konfigurasi untuk prediksi *stance* pada data yang belum berlabel.

Tabel 8. Perbandingan Dampak ROS pada IndoBERTweet untuk klasifikasi *stance* (mean, 5-fold CV)

Model	Kelas	Precision	Recall	F1-Score	Accuracy
IndoBERTweet <i>imbalanced</i> Epoch 4, LR 3e-05	<i>In Favor</i>	0.8117	0.8015	0.8066	0.9248
	<i>Against</i>	0.9547	0.9680	0.9613	
	<i>Neutral</i>	0.7929	0.7177	0.7534	
	<i>Macro</i>	0.8549	0.8291	0.8405	
IndoBERTweet <i>balanced</i> ROS Epoch 4, LR 3e-05	<i>In Favor</i>	0.8259	0.8025	0.8141	0.9269
	<i>Against</i>	0.9565	0.9670	0.9617	
	<i>Neutral</i>	0.7888	0.7446	0.7661	
	<i>Macro</i>	0.8581	0.8381	0.8474	



Gambar 4. Perbandingan Dampak ROS Pada Hiperparameter Terbaik IndoBERTweet *Stance*

Hasil Prediksi Data Sisa

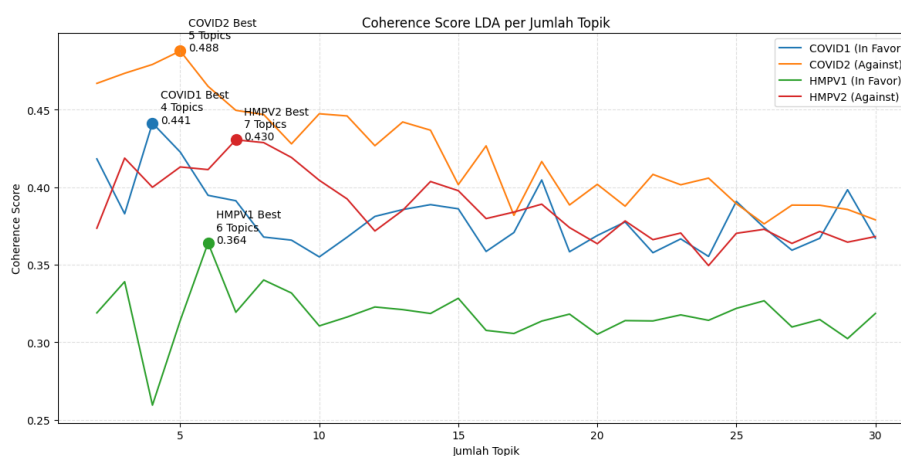
Sesuai dengan hasil pengukuran performa model dalam klasifikasi relevansi maupun *stance*, model yang digunakan untuk melabeli sisa data adalah IndoBERTweet dengan kombinasi *epoch* 4 dan *learning rate* 3e-05. Keduanya dilatih dengan menerapkan *oversampling* menggunakan ROS. Berdasarkan hal tersebut, hasil prediksi data sisa yang tidak dianotasi (50% Sampel) dapat dilihat pada Tabel 9 dimana proporsinya relatif serupa dengan data yang telah dianotasi secara manual. Proporsi kelas *Against* tetap menjadi yang kelas mayoritas dengan 7427 sampel (sekitar 77.5%). Kelas *In Favor*, *Neutral*, dan *Irrelevant* tetap menjadi minoritas.

Tabel 9. Proporsi Hasil Prediksi Data Sisa

Relevansi	Stance	Total	Persentase	COVID	HMPV
Relevant	In Favor	875	9.1%	475	400
Relevant	Against	7427	77.5%	5.296	2.160
Relevant	Neutral	870	9.1%	492	378
Irrelevant	-	414	4.3%	269	145
Total		9.586	100%	6.503	3.083

Hasil Pemodelan Topik

Gambar 5 menampilkan hasil pemilihan jumlah topik (k) pada pemodelan LDA menggunakan *coherence score* (rentang $k=1-30$) pada data gabungan sampel anotasi dan hasil prediksi. Pemodelan dilakukan terpisah untuk masing-masing penyakit (COVID-19 dan HMPV) serta untuk kelompok stance *in favor* dan *against* guna menangkap perbedaan struktur narasi penerimaan maupun penolakan terhadap ancaman risiko. Kurva menunjukkan bahwa skor *coherence* cenderung mencapai puncak pada k yang relatif kecil hingga menengah, kemudian menurun ketika k diperbesar, yang mengindikasikan bahwa penambahan topik setelah titik optimum lebih banyak memecah tema menjadi subtopik yang semakin spesifik namun kurang koheren. Berdasarkan *coherence score* tertinggi, jumlah topik optimal yang digunakan adalah: COVID *in favor* $k=4$ (0,441), COVID *against* $k=5$ (0,488), HMPV *in favor* $k=6$ (0,364), dan HMPV *against* $k=7$ (0,430).



Gambar 5. Hasil Grid Search Coherence Score LDA

Kata Kunci dan Interpretasi Topik Kategori Stance In Favor

Tabel 10. Kata Kunci dan Interpretasi Kategori In Favor COVID-19

No	Kata Kunci	Interpretasi Topik
1	covid, virus, kena, minum, sakit, makan, sembuh, obat, banyak, hari, batuk, baru, bulan, demam, tubuh	Gejala dan Penyembuhan
2	indonesia, covid, allah, orang, jangan, sehat, jaga, sakit, hati, tangan, masuk, takut, keluar, moga, amin	Respons Spiritual dan Doa
3	sehat, masker, pakai, negara, jaga, covid, vaksin, moga, indonesia, tetap, aman, tutup, semua, penting, kasus	Perilaku Protektif Pencegahan dan Penyembuhan
4	jadi, jangan, indonesia, waspada, covid, dulu, siap, masuk, banyak, perintah, luar, negeri, kalau, orang, salah	Kewaspadaan Publik dan Kebijakan Perbatasan

Tabel 11. Kata Kunci dan Interpretasi Kategori *In Favor* HMPV

No	Kata Kunci	Interpretasi Topik
1	covid, lebih, hari, china, obat, kayak, tahun, kena, sembuh, bulan, siap, baik, tinggal, minum, datang	Gejala dan Penyembuhan
2	hati, bandara, jaga, aman, indonesia, tutup, lalu, sehat, kalau, tiongkok, wabah, banyak, baru, semua, masuk	Kewaspadaan Publik dan Kebijakan Perbatasan
3	makan, sehat, minum, jangan, jaga, banyak, tetap, tubuh, obat, putih, virus, imun, vaksin, lockdown, rakyat	Perilaku Protektif Pencegahan dan Penyembuhan
4	sakit, batuk, kena, sembuh, rasa, demam, minggu, gejala, tenggorok, biasa, pilek, hmpv, sama, panas, lama	Testimoni Keparahan Penyakit
5	negara, virus, orang, jadi, moga, masker, covid, pakai, allah, indonesia, jangan, pandemi, amin, sekarang, masyarakat	Respons Spiritual dan Doa
6	virus, indonesia, masuk, dulu, covid, jangan, baru, perintah, cegah, china, bilang, kalau, negeri, awal, luar	Kewaspadaan Publik dan Kebijakan Perbatasan

Berdasarkan interpretasi topik pada kategori *In Favor* pada kedua penyakit (Tabel 10–11), tema narasi yang terbentuk pada COVID-19 dan HMPV relatif serupa dan merefleksikan penerimaan narasi risiko melalui pengakuan risiko melalui testimoni keparahan penyakit, tindakan protektif pencegahan dan penyembuhan, respons spiritual, serta kewaspadaan publik bernuansa kebijakan perbatasan. Analisis tema narasi yang muncul dalam kategori ini antara lain:

1. Gejala dan Penyembuhan

Pada topik ini, kata kunci yang muncul pada kedua penyakit mengindikasikan pengakuan risiko yang (covid) berupa pengalaman gejala (sakit, batuk, tenggorok, panas, demam) disertai indikator durasi (hari, minggu, lama) dan pemulihan (sembuh) yang dikaitkan dengan penanganan mandiri (makan, minum, obat). Pada HMPV, pengakuan risiko juga muncul dalam bentuk rujukan pengalaman sebelumnya (covid, tahun, kayak, sama).

2. Perilaku Protektif Pencegahan dan Penyembuhan

Pada topik ini, kata kunci yang muncul pada kedua penyakit mengindikasikan dukungan tindakan protektif (masker, pakai, dan jaga), intervensi pencegahan (vaksin), serta perilaku hidup sehat dan penggunaan ramuan herbal untuk meningkatkan imunitas (makan, sehat, minum, tubuh, imun, dan putih). Pada topik ini juga muncul dorongan untuk pencegahan mobilitas (indonesia, tutup, *lockdown*).

3. Respons Spiritual dan Doa

Pada topik ini, kata kunci berbasis spiritualitas muncul pada kedua penyakit (allah, moga, amin) sebagai respons terhadap ancaman (takut) disertai dengan ajakan kehati-hatian (hati, masker, sehat, jaga, sakit, tangan) dalam konteks mobilitas (negara, masuk, keluar) agar ancaman tidak tereskalasi menjadi pandemi seperti sebelumnya (covid, pandemi).

4. Kewaspadaan Publik dan Kebijakan Perbatasan

Pada topik ini, kata kunci yang muncul pada kedua penyakit mengindikasikan dorongan kebijakan (waspada, siap, perintah) dalam konteks (jangan, masuk, luar, negeri, bandara, tutup, china, tiongkok) merujuk pada pengalaman penanganan pandemi sebelumnya (awal, dulu, bilang, aman, baru).

Kata Kunci dan Interpretasi Topik Kategori *Stance Against*

Tabel 12. Kata Kunci dan Interpretasi Kategori *Against* COVID-19

No	Kata Kunci	Interpretasi Topik
1	vaksin, jual, covid, bill, coba, bisnis, gates, obat, ujung, indonesia, jadi, datang, suntik, kemarin, kali	Dugaan Motif Bisnis dan Ekonomi
2	covid, bikin, berita, rakyat, jangan, kayak, negara, ekonomi, biasa, usah, enggak, susah, dulu, begini, orang	Ketidakpercayaan pada Media
3	covid, sakit, mati, berita, sehat, jangan, semua, orang, percaya, media, virus, buat, dunia, manusia, banyak	Ketidakpercayaan pada layanan Medis
4	rakyat, takut, virus, covid, indonesia, cari, duit, buat, kerja, uang, varian, jadi, lebih, perintah, makan	Perbandingan Prioritas Risiko
5	covid, bisnis, percaya, jangan, rakyat, besar, mulai, proyek, cuan, bohong, korupsi, kasus, berita, koruptor, akal	Dugaan Motif Bisnis dan Ekonomi

Tabel 13. Kata Kunci dan Interpretasi Kategori *Against* HMPV

No	Kata Kunci	Interpretasi Topik
1	kata, sama, dharma, benar, pongrekun, percaya, darma, ingat, covid, orang, lama, bilang, pandemi, uang, jadi	Rujukan Figur Publik untuk Mendelegitimasi Risiko
2	virus, duit, korupsi, cari, bahaya, biasa, muncul, lebih, bisnis, wabah, akal, ladang, bulan, koruptor, lebaran	Perbandingan Prioritas Risiko
3	berita, jangan, virus, bikin, panik, usah, takut, indonesia, begini, media, percaya, baik, rakyat, sehat, sebar	Ketidakpercayaan pada Media
4	besar, takut, berita, media, buat, china, virus, makin, indonesia, terus, bodoh, kerja, ikut, jangan, heboh	Ketidakpercayaan pada Media
5	virus, rakyat, global, indonesia, dunia, biar, negara, perintah, ekonomi, buat, susah, baru, main, mati, elit	Dugaan Motif Bisnis dan Ekonomi
6	vaksin, bisnis, jual, obat, cuan, ujung, mulai, siap, masker, baru, drama, sehat, cair, jilid, beli	Dugaan Motif Bisnis dan Ekonomi
7	covid, sakit, jadi, sekarang, musim, dulu, indonesia, biasa, batuk, banyak, rumah, tahun, orang, waktu, dokter	Normalisasi Risiko dan Ketidakpercayaan pada Layanan Medis

Berdasarkan interpretasi topik pada kategori *Against* pada kedua penyakit (Tabel 12–13), narasi yang terbentuk pada COVID-19 dan HMPV cenderung mendelegitimasi ancaman penyakit dan menolak ajakan kewaspadaan. Pola umum yang muncul meliputi rujukan figur publik, ketidakpercayaan pada media dan layanan medis, tuduhan motif ekonomi/bisnis, serta perbandingan prioritas risiko. Pada HMPV terdapat kekhasan berupa normalisasi risiko. Analisis tema narasi yang muncul dalam kategori ini antara lain:

1. Rujukan Figur Publik untuk Mendelegitimasi Risiko:

Pada topik yang spesifik muncul dalam korpus HMPV, kata kunci yang muncul mengindikasikan rujukan figur (benar, percaya, kata, dharma, pongrekun) muncul bersama token (covid, pandemi, uang), yang membingkai risiko sebagai kepentingan pihak tertentu.

2. Ketidakpercayaan pada Media:

Pada topik ini, kata kunci yang muncul pada kedua penyakit (media, berita) muncul bersama penolakan (jangan, percaya) dan framing hiperbola/panik (bikin, panik, heboh, bodoh) dalam konteks penyakit (virus, covid, china), mengindikasikan

ketidakpercayaan terhadap pemberitaan yang dianggap membesar-besarkan untuk menciptakan ketakutan publik secara berlebihan.

3. Ketidakpercayaan pada Layanan Medis

Pada topik ini, kata kunci yang muncul (rumah, sakit, dokter) dan (sakit, mati) mengarah pada kecurigaan terhadap aspek penanganan medis, dikaitkan dengan pengalaman pandemi sebelumnya melalui kemunculan kata kunci (covid, pandemi).

4. Dugaan Motif Bisnis dan Ekonomi

Pada topik ini, kata kunci yang muncul pada kedua penyakit (proyek, cuan, bisnis, jual, obat, vaksin, uang/duit) muncul bersama delegitimasi narasi risiko (bohong, akal, drama) yang mengindikasikan delegitimasi risiko yang ditunjukkan dengan dugaan/kecurigaan terhadap motif ekonomi/bisnis di belakangnya.

5. Normalisasi Risiko

Topik ini spesifik muncul dalam korpus HMPV yang mengindikasikan normalisasi risiko terhadap penyakit dengan framing “flu biasa musiman” melalui kemunculan kata kunci (musim, biasa, batuk, dan dulu).

6. Perbandingan Prioritas Risiko

Pada topik ini, kata kunci yang muncul pada kedua jenis penyakit mengindikasikan perbandingan risiko (rakyat, lebih, takut, bahaya) dengan fenomena politik (korupsi) pada COVID-19 dan ekonomi (cari, duit, kerja, uang, makan) pada HMPV sehingga fokus ancaman penyakit dialihkan ke risiko yang dianggap lebih mendesak.

KESIMPULAN

Penelitian ini memetakan respons publik terhadap narasi risiko penyakit menular COVID-19 dan HMPV pada komentar YouTube melalui pendekatan *stance detection* dan pemodelan topik LDA. Secara metodologis, model IndoBERTweet cenderung lebih sesuai dibandingkan IndoBERT dengan karakteristik komentar YouTube yang cenderung tidak formal. Dengan penanganan ketidakseimbangan kelas menggunakan ROS, IndoBERTweet mencapai *F1-macro* 0.8541 pada klasifikasi relevansi dan 0.8474 pada klasifikasi *stance*. Secara substantif, kedua penyakit memiliki distribusi proporsi kelas *stance* yang cenderung serupa dan didominasi oleh kelas *against*. Hasil LDA menunjukkan bahwa pada kelas *in favor* narasi berkaitan dengan pengakuan risiko yang diekspresikan dengan gejala dan penyembuhan, dukungan perilaku protektif dan pencegahan, ekspresi spiritualitas melalui doa, dorongan kewaspadaan bernuansa kebijakan pembatasan. Sebaliknya, pada kelas *against* muncul narasi delegitimasi berupa ketidakpercayaan terhadap media dan layanan medis, tuduhan motif bisnis/ekonomi, perbandingan risiko, rujukan figur tertentu untuk delegitimasi risiko, serta normalisasi risiko yang hanya muncul pada korpus HMPV. Temuan ini menegaskan potensi kombinasi *stance detection* dan LDA untuk mendukung praktik *social listening* dalam mengidentifikasi pola penerimaan dan penolakan yang relevan digunakan dalam perumusan strategi komunikasi risiko berbasis data. Penelitian ini memiliki keterbatasan dalam cakupan jenis platform media sosial dan jenis kanal media resmi, keterbatasan analisis pada *parent comment*, keterbatasan rentang hiperparameter serta metode penanganan ketidakseimbangan

data. Penelitian selanjutnya dapat memperluas cakupan jenis media sosial dan kanal, menganalisis komentar balasan, serta mengevaluasi variasi strategi penanganan ketidakseimbangan data dan konfigurasi hiperparameter yang lebih luas. Selain itu, penelitian ini menggunakan target tunggal sehingga tidak dapat digeneralisasi langsung ke isu/klaim lain tanpa penyesuaian target dan pelatihan ulang model.

Ucapan terima kasih disampaikan kepada Lembaga Pengelola Dana Pendidikan (LPDP) atas dukungan pendanaan melalui program beasiswa yang memungkinkan penelitian ini dilaksanakan.

REFERENSI

- Alrowais, R. K. (2024). *Arabic stance detection of COVID-19 vaccination using transformer-based approaches : a comparison study*. 42(4), 1319–1339. <https://doi.org/10.1108/AGJSR-01-2023-0001>
- Berg, S. H., O'Hara, J. K., Shortt, M. T., Thune, H., Brønnick, K. K., Lungu, D. A., Røislien, J., & Wiig, S. (2021). Health authorities' health risk communication with the public during pandemics: a rapid scoping review. *BMC Public Health*, 21(1), 1401. <https://doi.org/10.1186/s12889-021-11468-3>
- Bimantara, I. M. S., Irdyanti, M., Nisa, C., & Purwitasari, D. (2025). Content and Network Feature in Attention-based Neural Network for Stance Detection on COVID-19 Vaccination Tweets. *JOIV: International Journal on Informatics Visualization*, 9(1), 286–294. <https://doi.org/10.62527/joiv.9.1.2671>
- Burgess, A. (2019). Environmental risk narratives in historical perspective: from early warnings to 'risk society' blame. *Journal of Risk Research*, 22(9), 1128–1142. <https://doi.org/10.1080/13669877.2018.1517383>
- Burnham, M. (2023). *Stance Detection: A Practical Guide to Classifying Political Beliefs in Text*. <https://doi.org/10.1017/psrm.2024.35>
- Charles-Smith, L. E., Reynolds, T. L., Cameron, M. A., Conway, M., Lau, E. H. Y., Olsen, J. M., Pavlin, J. A., Shigematsu, M., Streichert, L. C., Suda, K. J., & Corley, C. D. (2015). Using Social Media for Actionable Disease Surveillance and Outbreak Management: A Systematic Literature Review. *PloS One*, 10(10), e0139701. <https://doi.org/10.1371/journal.pone.0139701>
- Cotfas, L.-A., Delcea, C., Gherai, R., & Roxin, I. (2021). Unmasking People's Opinions behind Mask-Wearing during COVID-19 Pandemic—A Twitter Stance Analysis. *Symmetry*, 13(11). <https://doi.org/10.3390/sym13111995>
- Damanik, F. J., & Setyohadi, D. B. (2021). Analysis of public sentiment about COVID-19 in Indonesia on Twitter using multinomial naive bayes and support vector machine. *IOP Conference Series: Earth and Environmental Science*, 704(1), 1–12. <https://doi.org/10.1088/1755-1315/704/1/012027>
- Delcea, C., Cotfas, L.-A., Crăciun, L., & Molănescu, A. G. (2022). New Wave of COVID-19 Vaccine Opinions in the Month the 3rd Booster Dose Arrived. *Vaccines*, 10(6). <https://doi.org/10.3390/vaccines10060881>
- Delussu, F., Id, M. T., & Id, L. G. (2022). *Evidence of pandemic fatigue associated with stricter tiered COVID-19 restrictions*. November 2020, 1–14. <https://doi.org/10.1371/journal.pdig.0000035>
- Devlin, J., Chang, M.-W., Lee, K., Google, K. T., & Language, A. I. (2018). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. *Naacl-Hlt*

2019, *Mlm*, 4171–4186.

Dewan Pers. (2025). *Data Perusahaan Pers*. <https://dewanpers.or.id/data>

Duong, H. T., & Nguyen-Thi, T. A. (2021). A review: preprocessing techniques and data augmentation for sentiment analysis. *Computational Social Networks*, 8(1), 1–16. <https://doi.org/10.1186/s40649-020-00080-x>

Fahmy Amin, M. (2023). Confusion Matrix in Three-class Classification Problems: A Step-by-Step Tutorial. *Journal of Engineering Research*, 7(1), 0–0. <https://doi.org/10.21608/erjeng.2023.296718>

Glandt, K., Khanal, S., Li, Y., Caragea, D., & Caragea, C. (2021). Stance detection in COVID-19 tweets. *ACL-IJCNLP 2021 - 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing, Proceedings of the Conference*, 1596–1611. <https://doi.org/10.18653/v1/2021.acl-long.127>

Gurler, D., & Buyukceran, I. (2022). Assessment of the Medical Reliability of Videos on Social Media: Detailed Analysis of the Quality and Usability of Four Social Media Platforms (Facebook, Instagram, Twitter, and YouTube). *Healthcare*, 10(10). <https://doi.org/10.3390/healthcare10101836>

Han, Y. Q., Jeong, W. J., Oh, G. S., Oh, S. H., & Whangbo, T. K. (2023). Using LDA Topic Modeling to Understand Regrowth Factors of the Chinese Gaming Industry in the COVID-19 Era: Current Situation, Future and Predicament. *Journal of Web Engineering*, 22(3), 433–463. <https://doi.org/10.13052/jwe1540-9589.2233>

Harapan, B. N., Harapan, T., Theodora, L., & Anantama, N. A. (2023). From Archipelago to Pandemic Battleground: Unveiling Indonesia's COVID-19 Crisis. *Journal of Epidemiology and Global Health*, 13(4), 591–603. <https://doi.org/10.1007/s44197-023-00148-7>

Kang, S., Choi, Y.-H., & Cheon, M. (2024). Unraveling YouTube Stances on Global Warming: An In-Depth Analysis of Skeptics and Believers. *Environmental Science and Engineering*, 207 – 220. https://doi.org/10.1007/978-981-97-3320-0_16

Kemenkes. (2021). *Pedoman Komunikasi Risiko untuk Penanggulangan Krisis Kesehatan*. Biro Komunikasi dan Pelayanan Masyarakat Kementerian Kesehatan RI.

Kementerian Kesehatan Republik Indonesia. (2021). *Rencana Operasi Penanggulangan COVID-19 Bidang Kesehatan di Indonesia* (R. Djupuri, Supatmi, R. Ramadhanjaya, Wijayanti, A. Subur, & V. Roza (eds.); Revisi 1). Kementerian Kesehatan Republik Indonesia.

Koto, F., Lau, J. H., & Baldwin, T. (2021). {I}ndo{BERT}weet: A Pretrained Language Model for {I}ndonesian {T}witter with Effective Domain-Specific Vocabulary Initialization. In M.-F. Moens, X. Huang, L. Specia, & S. W. Yih (Eds.), *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing* (pp. 10660–10668). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2021.emnlp-main.833>

Küçük, D., & Can, F. (2020). Stance Detection: A Survey. *ACM Comput. Surv.*, 53(1). <https://doi.org/10.1145/3369026>

Kumar, M., Khan, L., & Chang, H. T. (2025). Evolving techniques in sentiment analysis: a comprehensive review. *PeerJ Computer Science*, 11, 1–61. <https://doi.org/10.7717/PEERJ-CS.2592>

Manuaba, I. B. K. (2023). A Sentiment Analysis Model for the COVID-19 Vaccine in

- Indonesia Using Twitter API v2, TextBlob, and Googletrans. *Procedia Computer Science*, 227, 1101–1110. <https://doi.org/10.1016/j.procs.2023.10.621>
- Miao, L., Last, M., & Litvak, M. (2022). Tracking social media during the COVID-19 pandemic: The case study of lockdown in New York State. *Expert Systems with Applications*, 187, 115797. <https://doi.org/10.1016/j.eswa.2021.115797>
- Negara, E. S., Triadi, D., & Andryani, R. (2019). Topic Modelling Twitter Data with Latent Dirichlet Allocation Method. *ICECOS 2019 - 3rd International Conference on Electrical Engineering and Computer Science, Proceeding*, 386–390. <https://doi.org/10.1109/ICECOS47637.2019.8984523>
- Ong, Y. X., Kim, H. K., Pelzer, B. O., Tan, Y. Y., Lim, W. P., Chua, A. K. L., & Koh, B. Y. (2023). Profile identification and characterization of risk perceptions and preventive behaviors during the COVID-19 pandemic: A latent profile analysis. *Frontiers in Psychology*, 14(February), 1–15. <https://doi.org/10.3389/fpsyg.2023.1085208>
- Padurariu, C., & Breaban, M. E. (2019). Dealing with data imbalance in text classification. *Procedia Computer Science*, 159, 736–745. <https://doi.org/10.1016/j.procs.2019.09.229>
- Purba, H., & Panjaitan, E. S. (2025). Sentiment Analysis to Detect Public Anxiety About the HMPV Virus On Social Media. *Brilliance: Research of Artificial Intelligence*, 5(1), 372–383. <https://doi.org/10.47709/brilliance.v5i1.6329>
- Purwitasari, D., Putra, C. B. P., & Raharjo, A. B. (2023). A stance dataset with aspect-based sentiment information from Indonesian COVID-19 vaccination-related tweets. *Data in Brief*, 47, 108951. <https://doi.org/https://doi.org/10.1016/j.dib.2023.108951>
- Sarker, R., Roknuzzaman, A. S. M., Hossain, M. J., Bhuiyan, M. A., & Islam, M. R. (2023). The WHO declares COVID-19 is no longer a public health emergency of international concern: benefits, challenges, and necessary precautions to come back to normal life. In *International journal of surgery (London, England)* (Vol. 109, Issue 9, pp. 2851–2852). <https://doi.org/10.1097/JS9.0000000000000513>
- Sharma, C., Batra, I., Sharma, S., Malik, A., Sanwar Hosen, A. S. M., & Ra, I. H. (2022). Predicting Trends and Research Patterns of Smart Cities: A Semi-Automatic Review Using Latent Dirichlet Allocation (LDA). *IEEE Access*, 10(August), 121080–121095. <https://doi.org/10.1109/ACCESS.2022.3214310>
- Shortt, M. T., Smeets, I., Wiig, S., Berg, S. H., Lungu, D. A., Thune, H., & Røislien, J. (2021). Shortcomings in Public Health Authorities' Videos on COVID-19: Limited Reach and a Creative Gap. *Frontiers in Communication*, Volume 6-. <https://doi.org/10.3389/fcomm.2021.764220>
- Stechemesser, A., Kotz, M., Auffhammer, M., & Wenz, L. (2023). Prolonged exposure weakens risk perception and behavioral mobility response: Empirical evidence from Covid-19. *Transportation Research Interdisciplinary Perspectives*, 22, 100906. <https://doi.org/https://doi.org/10.1016/j.trip.2023.100906>
- Tejamaya, M., Putri, A. A., Widanarko, B., Wirawan, I. M. A., Kurniawan, B., & Thamrin, Y. (2025). Dynamics of COVID-19 Risk Perception in Indonesia: A Consecutive Cross-sectional Study from 2020 to 2022. *Safety and Health at Work*. <https://doi.org/https://doi.org/10.1016/j.shaw.2025.07.005>
- Wang, X., Lin, L., Xuan, Z., Xu, J., Wan, Y., & Zhou, X. (2020). Risk communication on behavioral responses during COVID-19 among general population in China: A rapid national study. *The Journal of Infection*, 81(6), 911–922. <https://doi.org/10.1016/j.jinf.2020.10.031>
- WHO. (2023). *Strategic preparedness and response plan: April 2023-April 2025*. From

- emergency response to long-term COVID-19 disease management: sustaining gains made during the COVID-19 pandemic* (Issue WHO/WHE/SPP/2023.1).
- WHO. (2024). *COVID-19 policy briefs*. <https://www.who.int/emergencies/diseases/novel-coronavirus-2019/covid-19-policy-briefs>
- WHO. (2025). *Social listening in infodemic management for public health emergencies: guidance on ethical considerations*. World Health Organization.
- Wilie, B., Vincentio, K., Winata, G. I., Cahyawijaya, S., Li, X., Lim, Z. Y., Soleman, S., Mahendra, R., Fung, P., Bahar, S., & Purwarianti, A. (2020). {I}ndo{NLU}: Benchmark and Resources for Evaluating {I}ndonesian Natural Language Understanding. In K.-F. Wong, K. Knight, & H. Wu (Eds.), *Proceedings of the 1st Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics and the 10th International Joint Conference on Natural Language Processing* (pp. 843–857). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2020.aacl-main.85>
- World Population Review. (2025). *YouTube Users by Country*. YouTube Country Ranking; World Population Review.
- Yan, H., Ma, M., Wu, Y., Fan, H., & Dong, C. (2022). Overview and analysis of the text mining applications in the construction industry. *Heliyon*, 8(12), e12088. <https://doi.org/https://doi.org/10.1016/j.heliyon.2022.e12088>
- Yulita, I. N., Wijaya, V., Rosadi, R., Sarathan, I., Djuyandi, Y., & Prabuwno, A. S. (2023). Analysis of Government Policy Sentiment Regarding Vacation during the COVID-19 Pandemic Using the Bidirectional Encoder Representation from Transformers (BERT). *Data*, 8(3). <https://doi.org/10.3390/data8030046>