

BAB 1 PENDAHULUAN

1.1 Latar belakang

Media sosial X (dahulu dikenal sebagai Twitter) memegang peran krusial sebagai cerminan digital dari sikap, opini, dan dinamika sosial masyarakat. Berdasarkan penelitian Azmi & Al-Ghadir (2024), platform ini sangat efektif digunakan untuk menangkap respons publik secara *real-time* terhadap berbagai isu kontroversial. Di Indonesia, salah satu fenomena lokal yang belakangan ini memicu perdebatan sengit di media sosial X adalah tren *Sound Horeg*. Istilah ini merujuk pada aktivitas penggunaan sistem tata suara (*sound system*) berskala masif, umumnya dalam acara karnaval atau hajatan yang menghasilkan dentuman suara bervolume ekstrem di ruang publik.

Media sosial menjadi ruang diskusi publik yang berguna untuk menyampaikan berupa kritik, opini, dan respon terhadap isu-isu sosial maupun politik. Dengan adanya keterbukaan ini juga dapat menyebabkan rentan terhadap ujaran yang tidak santun dan intoleransi yang dapat menyebabkan mengganggu kualitas komunikasi publik. Penelitian oleh Oz & Nurumov, (2022) mengkaji dua bentuk ujaran negatif yaitu komentar tidak santun dan komentar intoleran. Komentar tidak santun memiliki nada kasar namun masih memiliki argumen rasional, sedangkan komentar intoleran lebih menyerang ke identitas atau kelompok berdasarkan ras, agama, orientasi politik sehingga dapat berpotensi merusak hubungan komunikasi publik.

Fenomena ini menuai reaksi beragam dari masyarakat beberapa menyebutnya sebagai bentuk hiburan namun tak sedikit pula yang mengecamnya karena dianggap mengganggu kenyamanan dan membahayakan. Beberapa menyebutkan bahwa volume suara yang dihasilkan dapat menyebabkan kerusakan properti seperti kaca rumah yang pecah atau struktur bangunan yang bergetar. Penelitian yang dilakukan oleh Endriks Endrianto, (2023) dari segi kesehatan, batas aman pendengaran standar manusia adalah 85db (*desibel*). Penelitian oleh Ergun

dkk., (2024) menunjukkan bahwa paparan suara di atas 65db dapat menurunkan kemampuan pendengaran, terutama pada frekuensi rendah. Oleh karena itu fenomena sound horeg tidak hanya menjadi kebisingan semata, tetapi juga mencerminkan adanya ketegangan sosial dan sentimen negatif yang muncul dalam ruangan digital. Sentimen sentimen tersebut terutama pada intoleransi dapat menjadi data penting untuk dianalisis guna memahami bagaimana opini publik terbentuk dan berkembang di media sosial.

Analisis sentimen dan analisis intoleransi merupakan dua pendekatan yang sering digunakan dalam kajian teks digital, namun memiliki fokus yang berbeda. Menurut penelitian yang dilakukan Febriyani & Februriyanti, (2023) analisis sentimen adalah mengidentifikasi dan mengklasifikasi emosi atau opini pengguna terhadap sesuatu isu, produk, atau layanan, yang dikategorikan ke dalam sentimen positif, negatif, atau netral. Sementara itu, menurut penelitian Dhiwa Aqsha, (2024) analisis intoleransi lebih spesifik mengarah pada deteksi ujaran kebencian atau konten yang menyerang kelompok berdasarkan identitas seperti agama, ras, orientasi seksual, atau ideologi tertentu. Dalam penelitian yang dilakukan oleh Subramanian dkk., (2023) dijelaskan bahwa meskipun kedua pendekatan ini saling berkaitan, analisis intoleransi memerlukan pemahaman konteks yang lebih mendalam dibandingkan analisis sentimen biasa karena ujaran kebencian bisa bersifat implisit, kontekstual, atau disamarkan dalam bentuk sarkasme. Perbedaan ini penting terutama dalam menganalisis fenomena sosial digital seperti sound horeg, karena komentar yang muncul bisa kritik keras, ekspresi emosi, hingga ujaran intoleran yang menyerang kelompok tertentu. Dengan memahami perbedaan ini, pendekatan analisis menjadi lebih tepat sasaran, dan hasil klasifikasi dapat digunakan secara lebih bertanggung jawab, baik dalam konteks akademik, sosial, maupun kebijakan publik.

Meskipun berbagai penelitian telah dilakukan untuk mendeteksi ujaran kebencian di media sosial, masih terdapat celah metodologi dan evaluasi yang signifikan dari studi terdahulu. Pertama, penelitian oleh Turki & Roy, (2022), klasifikasi biner, sehingga gagal menangkap kompleksitas bahasa seperti sarkasme atau kritik sosial tersamar. Kedua penelitian oleh Kennedy Malanga Ndenga,

(2023), yang menggunakan *embedding* kalimat tingkat lanjut (*Universal Sentence Encoder*) masih bergantung pada daftar kata statis ujaran kebencian dari NCIC yang bersifat statis, sehingga berisiko menghasilkan *false positive* karena kurangnya konteks semantik yang dinamis. ketiga, penelitian terdahulu oleh Putri dkk., (2023) justru terlalu berfokus pada analisis sentimen terhadap brand tanpa mendeteksi hate speech secara eksplisit. Studi tersebut juga mengabaikan validitas *ground truth* karena tidak menggunakan anotasi manual, serta hanya mengevaluasi model menggunakan akurasi tanpa memperhitungkan ketidakseimbangan kelas (*class imbalance*).

Serangkaian keterbatasan tersebut menjadi landasan utama dilakukannya penelitian ini. Menutup keterbatasan studi terdahulu, penelitian ini mengusulkan pendekatan *Ensemble Learning* yang mengombinasikan metode *Support Vector Machine* (SVM) dan *K-Nearest Neighbor* (KNN) melalui strategi seperti *Voting Classifier*. Algoritma SVM dipilih karena efektivitasnya yang teruji dalam menangani data teks berdimensi tinggi pada klasifikasi ujaran kebencian Nadhia Salsabila Azzahra dkk., (2021). Sementara itu, KNN diintegrasikan karena keunggulannya dalam mengenali pola kedekatan lokal antar-data berbasis pembobotan TF-IDF Nova Adi Saputra dkk., (2024). Melalui strategi *ensemble* (seperti *Voting Classifier*), menggabungkan kedua metode ini secara dinamis diharapkan dapat saling mengurangi keterbatasan masing-masing algoritma, sehingga model lebih sensitif dalam memisahkan batas tipis antara kritik sosial yang tajam dengan ujaran intoleransi *Sound Horeg* yang bersifat implisit. Selain itu, guna mengatasi kelemahan evaluasi konvensional, penelitian ini mengintegrasikan metrik *Precision*, *Recall*, dan *F1-Score* untuk memastikan performa model tetap valid dan objektif di tengah tantangan ketidakseimbangan distribusi data (*class imbalance*) di media sosial X.

1.2 Rumusan Masalah

- 1) Bagaimana algoritma *support vector machine* (SVM), *k-nearest neighbor* (KNN), dan model *Ensemble* (Gabungan SVM dan KNN dengan skema *soft*

voting) dapat mengenali ujaran intoleransi pada fenomena *sound* horeg di media X sehingga dapat digunakan sebagai bahan analisis?

- 2) Apa saja karakteristik ujaran yang menjadi sumber kesulitan model dalam mengenali wacana intoleransi pada fenomena *sound* horeg, dan bagaimana kesalahan pola kesalahan klasifikasi tersebut dapat dijelaskan?
- 3) Bagaimana dinamika wacana intoleransi pada fenomena *sound* horeg di media sosial X dapat diuraikan melalui hasil klasifikasi dan ekstraksi kata kunci diskriminatif model sebagai instrumen analisis?

1.3 Tujuan

- 1) Mengetahui kemampuan algoritma *support vector machine* (SVM), *k-nearest neighbor* (KNN), dan model *Ensemble* dalam mengenali ujaran intoleransi pada fenomena *sound* horeg di media sosial X sehingga layak digunakan sebagai instrumen analisis.
- 2) Mengidentifikasi karakteristik ujaran yang menjadi sumber kesulitan model dalam mengenali wacana intoleransi pada fenomena *sound* horeg, serta menjelaskan pola kesalahan klasifikasi yang muncul.
- 3) Menguraikan dinamika wacana intoleransi pada fenomena *sound* horeg di media sosial X melalui hasil klasifikasi dan ekstraksi kata kunci diskriminatif model sebagai instrumen analisis.

1.4 Manfaat

1.4.1. Bagi akademis

- a. Memberikan kontribusi dalam pengembangan kajian analisis sentimen berbasis *machine learning*, khususnya pada isu sosial dan intoleransi di media sosial
- b. Menjadi referensi untuk penelitian selanjutnya yang ingin mengkaji kombinasi metode klasifikasi seperti SVM dan KNN dengan pendekatan *ensembl*.

1.4.2. Bagi praktisi

- a. Dapat digunakan sebagai dasar pengembangan sistem deteksi otomatis komentar atau postingan yang berpotensi mengandung intoleransi.

- b. Membantu masyarakat dalam memahami perbedaan antara kritik sosial dan ujaran kebencian melalui pendekatan analisis data.

1.4.3. Bagi penulis

- a. Menambah pengalaman dan pemahaman penulis dalam menerapkan teknik *machine learning* untuk analisis teks.
- b. Melatih kemampuan penulis dalam merancang dan mengimplementasikan sistem klasifikasi berbasis algoritma SVM dan KNN secara integrasi

1.5 Batasan Masalah

Berdasarkan latar belakang yang sudah di bahas maka batasan masalah dari penelitian ini sebagai berikut:

- a. Data yang digunakan hanya berupa teks, dan tidak termasuk media lain seperti gambar, video, atau audio.
- b. Objek yang dianalisis dibatasi pada komentar atau postingan terkait fenomena *sound horeg* yang berasal dari platform media sosial *X*.
- c. Fokus analisis adalah ujaran yang mengandung unsur intoleransi