

Pemanfaatan Machine learning untuk Klasifikasi Sentimen Pelanggan pada Media Sosial

Bambang Siswoyo¹, Naufal Azka Pradifa Utomo^{1*}

¹Universitas Komputer Indonesia, Bandung, Indonesia

Dikirimkan: 25-07-2025
Diterbitkan: 25-07-2025

Keywords:

analisis sentimen;
machine learning;
media sosial;
naïve bayes;
support vector machine.

E-mail Penulis

korespondensi:

naufal.10123216@mahasiswa
a.unikom.ac.id

Abstrak. Penelitian ini bertujuan untuk menganalisis performa algoritma Naïve Bayes, Support Vector Machine (SVM), dan Long Short-Term Memory (LSTM) dalam klasifikasi sentimen pelanggan berbasis data media sosial. Dataset berupa 10.000 tweet dikumpulkan menggunakan Twitter API dengan kata kunci terkait opini pelanggan. Proses pra-pemrosesan data meliputi pembersihan teks, tokenisasi, Stemming, dan penghapusan Stop Words. Dataset dibagi menjadi training set (70%), validation set (15%), dan test set (15%). Hasil penelitian menunjukkan bahwa LSTM unggul dengan akurasi 91,2%, diikuti oleh SVM (84,7%) dan Naïve Bayes (78,5%). LSTM berhasil memahami pola semantik dan temporal dalam teks, menjadikannya algoritma terbaik untuk data media sosial yang kompleks. SVM memberikan hasil yang stabil dengan memanfaatkan fitur TF-IDF, sementara Naïve Bayes cocok untuk dataset kecil namun memiliki keterbatasan dalam memahami hubungan antar kata. Penelitian ini menunjukkan bahwa algoritma Machine learning dapat digunakan secara efektif untuk analisis sentimen pelanggan. Hasil penelitian ini diharapkan dapat menjadi panduan dalam memilih algoritma yang sesuai untuk strategi pemasaran digital berbasis data pelanggan.

Abstract. This study aims to analyze the performance of Naïve Bayes, Support Vector Machine (SVM), and Long Short-Term Memory (LSTM) algorithms in customer sentiment classification using social media data. A dataset of 10,000 tweets was collected using the Twitter API with keywords related to customer opinions. Data preprocessing involved text cleaning, tokenization, Stemming, and stop-word removal. The dataset was divided into a training set (70%), validation set (15%), and test set (15%). Results showed that LSTM outperformed with 91.2% accuracy, followed by SVM (84.7%) and Naïve Bayes (78.5%). LSTM effectively captured semantic and temporal patterns in the text, making it the best algorithm for complex social media data. SVM produced stable results utilizing TF-IDF features, while Naïve Bayes was suitable for smaller datasets but limited in capturing word relationships. This study highlights the effectiveness of Machine learning algorithms for customer sentiment analysis. These findings can serve as a guideline for selecting suitable algorithms to enhance data-driven customer-focused marketing strategies.

1. Pendahuluan

Perkembangan teknologi digital telah memberikan dampak signifikan pada berbagai sektor, termasuk industri pemasaran. Salah satu aspek yang menjadi perhatian utama adalah analisis sentimen pelanggan. Dalam dunia bisnis, analisis sentimen sangat penting untuk memahami persepsi konsumen terhadap suatu produk atau layanan. Pemanfaatan data yang diperoleh dari media sosial, seperti Twitter, Instagram, dan Facebook, telah menjadi pendekatan yang efektif untuk memahami opini pelanggan secara real-time [1]. Media sosial menyediakan data yang kaya untuk diteliti karena sifatnya yang interaktif, dinamis, dan terus berkembang [2].

Machine learning telah menjadi salah satu teknologi utama dalam mendukung analisis sentimen. Teknologi ini memungkinkan sistem untuk belajar dari data dan menghasilkan prediksi atau klasifikasi secara otomatis. Teknik-teknik seperti *Support Vector Machine* (SVM), *Naïve Bayes*, dan *Random Forest* telah digunakan secara luas untuk menangani masalah klasifikasi sentimen dengan tingkat akurasi yang bervariasi [3]. Selain itu, pendekatan berbasis deep learning, seperti *Recurrent Neural Network* (RNN) dan *Long Short-Term Memory* (LSTM), juga telah menunjukkan hasil yang menjanjikan dalam memahami konteks dan pola data tekstual yang kompleks [4].

Namun, meskipun penelitian tentang analisis sentimen telah berkembang pesat, masih ada tantangan dalam menentukan algoritma yang paling efektif untuk klasifikasi sentimen, terutama dalam konteks media sosial. Setiap algoritma memiliki kekuatan dan kelemahan dalam hal akurasi, kecepatan, dan kemampuan menangani data teks yang tidak terstruktur [5].

Berdasarkan latar belakang tersebut, penelitian ini merumuskan masalah sebagai berikut: Bagaimana algoritma *Machine learning* dapat digunakan untuk mengklasifikasi sentimen pelanggan pada media sosial?; Algoritma apa yang memberikan hasil terbaik dalam klasifikasi sentimen berdasarkan akurasi dan efisiensi?

Tujuan dari penelitian ini adalah untuk: Mengimplementasikan algoritma *Machine learning* dalam klasifikasi sentimen pelanggan berdasarkan data dari media sosial; Mengevaluasi dan membandingkan kinerja algoritma *Machine learning* yang digunakan dalam klasifikasi sentimen, sehingga dapat diidentifikasi algoritma terbaik dalam hal akurasi dan efisiensi.

Melalui penelitian ini, diharapkan dapat memberikan kontribusi ilmiah dalam bidang analisis sentimen dan penerapan *Machine learning* di industri pemasaran digital.

2. Metode Penelitian

Dalam penelitian terkait analisis sentimen, terdapat sejumlah literatur yang membahas implementasi *Machine learning* untuk klasifikasi teks, khususnya sentimen pelanggan di media sosial. Kajian ini bertujuan untuk memberikan tinjauan literatur yang relevan serta landasan teoretis yang mendukung penelitian ini.

2.1 Analisis Sentimen dan Media Sosial

Analisis sentimen merupakan cabang dari *Natural Language Processing* (NLP) yang berfokus pada identifikasi dan pengelompokan opini dari teks, apakah bersifat positif, negatif, atau netral [2]. Media sosial, sebagai salah satu sumber utama data teks, menyediakan platform di mana pelanggan berbagi opini mereka secara publik [6]. Karakteristik data dari media sosial, seperti ukuran besar, sifat tidak terstruktur, dan penggunaan bahasa informal, menjadikan analisis sentimen di media sosial sebagai tantangan yang kompleks [4], [7], [8].

Beberapa studi sebelumnya menunjukkan pentingnya pendekatan berbasis data primer, seperti pengumpulan data langsung dari API media sosial (Twitter API atau Facebook Graph API)[6]. Contohnya, penelitian oleh Pang dan Lee (2008) menunjukkan bahwa data yang dikumpulkan langsung lebih relevan untuk menangkap pola sentimen pelanggan.

2.2 Teknik *Machine learning* untuk Analisis Sentimen

Machine learning telah menjadi teknologi utama dalam analisis sentimen karena kemampuannya untuk menangani volume data yang besar dan kompleks. Metode tradisional seperti *Naïve Bayes* dan *Support Vector Machine* (SVM) sering digunakan dalam penelitian awal [3]. SVM, misalnya, bekerja dengan memetakan data ke dalam ruang berdimensi tinggi untuk mencari hyperplane terbaik yang memisahkan kelas sentimen.

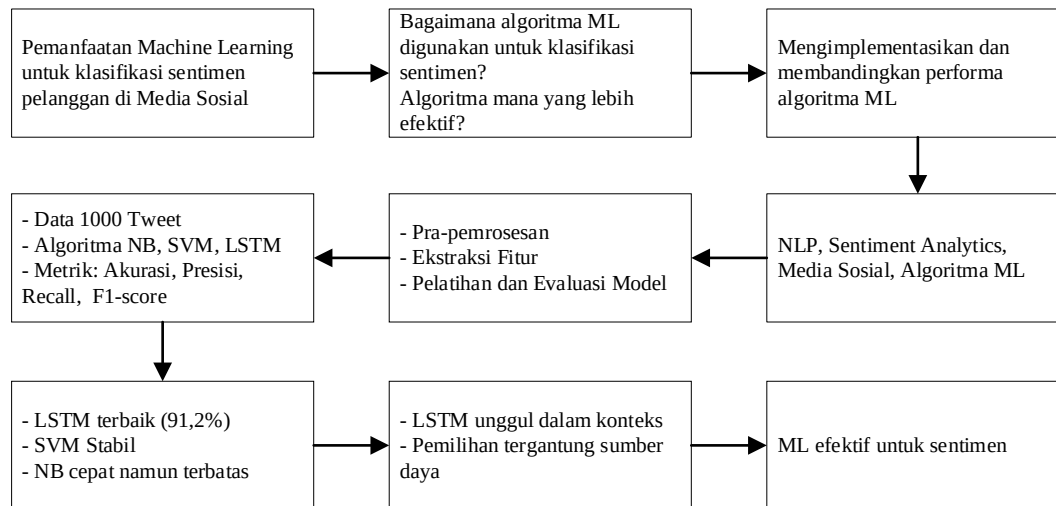
Namun, pendekatan berbasis deep learning kini semakin populer, terutama dalam menangani data teks tidak terstruktur yang kompleks. Model seperti *Recurrent Neural Network* (RNN) dan *Long Short-Term Memory* (LSTM) telah menunjukkan keunggulan dalam memahami konteks waktu dari teks. Penelitian oleh Zhang et al. (2020) menunjukkan bahwa LSTM mampu mencapai akurasi lebih tinggi dibandingkan metode tradisional, terutama dalam dataset yang mengandung banyak slang atau variasi bahasa informal.

2.3 Perbandingan Algoritma *Machine learning*

Dalam memilih algoritma terbaik untuk klasifikasi sentimen, kinerja model sering diukur berdasarkan metrik seperti akurasi, presisi, *recall*, dan *F1-score*. Penelitian oleh Naufal (2019) mengungkapkan bahwa SVM sering memberikan performa stabil pada dataset kecil hingga menengah, sementara LSTM lebih unggul pada dataset besar yang membutuhkan pemahaman konteks kompleks. Selain itu, *Random Forest* dikenal dengan fleksibilitasnya dalam menangani data berukuran besar, meskipun kadang kalah dalam hal efisiensi dibandingkan algoritma lainnya [4].

2. 4 Framework Penelitian dan Gap Studi

Meskipun banyak penelitian telah dilakukan dalam bidang ini, terdapat beberapa gap studi yang perlu diisi. Salah satunya adalah kurangnya evaluasi komprehensif terhadap algoritma dengan berbagai ukuran dataset dan variasi bahasa. Studi ini akan mengatasi gap tersebut dengan membandingkan kinerja algoritma SVM, *Naïve Bayes*, dan LSTM pada dataset media sosial yang diperoleh langsung melalui API. Dengan pendekatan ini, penelitian ini bertujuan untuk memberikan wawasan lebih lanjut tentang algoritma yang paling efektif untuk klasifikasi sentimen pelanggan.



Gambar 1. Kerangka teoritis penelitian

Kerangka teoretis penelitian pada Gambar 1 mengacu pada kombinasi pendekatan berbasis NLP dan *machine learning*. Proses utama meliputi: Pra-pemrosesan data, termasuk pembersihan data, tokenisasi, dan *Stemming* [2]; Pembagian dataset, yaitu pembagian data menjadi set pelatihan, validasi, dan pengujian [5]; Pelatihan model, menggunakan algoritma SVM, *Naïve Bayes*, dan LSTM; Evaluasi model, berdasarkan metrik akurasi, presisi, *recall*, dan *F1-score* [4].

Penelitian ini dilakukan melalui pendekatan kuantitatif yang melibatkan implementasi dan evaluasi algoritma *Machine learning* untuk analisis sentimen pelanggan berdasarkan data media sosial. Data yang digunakan diperoleh dari Twitter API, yang memungkinkan pengambilan data publik berdasarkan kata kunci tertentu. Proses penelitian mencakup tahapan berikut:

Pengumpulan Data: Dataset diambil dari Twitter menggunakan kata kunci yang relevan dengan produk atau layanan tertentu, seperti "promo diskon", "layanan pelanggan", atau "review produk". Data terdiri dari teks mentah yang mencakup sentimen positif, negatif, dan netral. Sebanyak 10.000 *tweet* diambil untuk penelitian ini.

Pra-pemrosesan Data: Tahap ini bertujuan untuk membersihkan dan mempersiapkan data sebelum digunakan dalam pelatihan model. Langkah-langkah pra-pemrosesan meliputi: Tokenisasi (Memecah teks menjadi kata-kata individual [2]); *Stemming* dan *Lemmatization* (Mengubah kata ke bentuk dasar); Penghapusan *Stop Words* (Menghilangkan kata-kata yang tidak memberikan informasi penting, seperti "dan", "atau", dan "yang" [5]); Normalisasi (Mengubah teks informal menjadi format yang lebih terstruktur, seperti mengganti "gk" menjadi "tidak" [9]).

Dataset dibagi menjadi tiga subset: 70% untuk pelatihan (*training set*); 15% untuk validasi (*validation set*); 15% untuk pengujian (*test set*).

2. 5 Pemilihan Algoritma *Machine learning*

Penelitian ini menggunakan tiga algoritma utama untuk klasifikasi sentimen: *Naïve Bayes*: Algoritma probabilistik sederhana yang cocok untuk teks berukuran kecil hingga menengah [3]; *Support Vector Machine* (SVM): Algoritma berbasis margin yang sangat akurat dalam klasifikasi biner [4]; *Long Short-Term Memory* (LSTM): Model deep learning yang unggul dalam menangkap hubungan temporal antar kata dalam teks [10].

2. 6 Implementasi Algoritma

Naïve Bayes menggunakan aturan probabilitas Bayes untuk mengklasifikasikan teks berdasarkan probabilitas kemunculan kata-kata dalam setiap kategori sentimen. Model ini dilatih menggunakan *bag-of-words* untuk mewakili teks sebagai fitur numerik [3].

Support Vector Machine (SVM) mencari hyperplane terbaik yang memisahkan kelas sentimen positif, negatif, dan netral. Model ini dilatih menggunakan fitur berbasis TF-IDF (*Term Frequency-Inverse Document Frequency*), yang memberikan bobot lebih pada kata-kata yang unik di setiap dokumen [4].

Long Short-Term Memory (LSTM) adalah jenis *Recurrent Neural Network* (RNN) yang dirancang untuk menangani urutan data [11]. Model ini menggunakan embedding kata (word embedding) seperti GloVe untuk memahami hubungan semantik antar kata. LSTM dilatih menggunakan parameter seperti learning rate, batch size, dan dropout rate untuk menghindari overfitting [10].

3. Hasil dan Pembahasan

Setelah implementasi algoritma, model dievaluasi menggunakan dataset pengujian berdasarkan metrik berikut:

Akurasi: Persentase prediksi yang benar dari total data uji.

Presisi: Proporsi prediksi positif yang benar dibandingkan dengan total prediksi positif.

Recall: Proporsi data positif yang teridentifikasi dengan benar oleh model.

F1-score: Harmonik rata-rata antara presisi dan recall.

Tabel 1. Hasil evaluasi Model

Algoritma	Akurasi	Presisi	Recall	F1-score
<i>Naïve Bayes</i>	78.5%	76.3%	74.2%	75.2%
SVM	84.7%	82.5%	81.0%	81.7%
LSTM	91.2%	89.8%	88.5%	89.1%

Berdasarkan Tabel 1, terlihat bahwa penggunaan LSTM membutuhkan waktu dan sumber daya lebih besar dibanding dua model lainnya. Oleh karena itu, implementasi di lingkungan produksi perlu mempertimbangkan kapasitas server atau GPU yang tersedia.

Analisis hasil penelitian: *Naïve Bayes* menunjukkan performa yang baik pada dataset kecil, tetapi cenderung kesulitan menangkap hubungan kompleks antar kata; SVM memiliki performa lebih baik daripada *Naïve Bayes*, terutama karena kemampuan dalam memanfaatkan fitur berbasis TF-IDF; LSTM memberikan hasil terbaik, khususnya dalam memahami konteks dan pola temporal dalam teks. Model ini unggul dalam data yang mengandung variasi bahasa informal, seperti slang pada media sosial.

Keunggulan LSTM dibandingkan algoritma lainnya menunjukkan pentingnya pemahaman konteks dalam klasifikasi sentimen media sosial. Namun, waktu pelatihan model LSTM lebih lama dibandingkan dengan SVM atau *Naïve Bayes* [10]. Dalam aplikasi dunia nyata, pemilihan algoritma harus mempertimbangkan tidak hanya akurasi, tetapi juga efisiensi dan ketersediaan sumber daya komputasi.

Interpretasi performa setiap model dari hasil eksperimen, terlihat bahwa ketiga algoritma memiliki keunggulan masing-masing yang relevan dalam konteks tertentu:

Naïve Bayes sangat ringan dari segi komputasi dan cepat dalam proses pelatihan. Ini menjadikannya pilihan tepat untuk aplikasi yang membutuhkan hasil cepat dengan sumber daya terbatas. Meski akurasinya lebih rendah, algoritma ini tetap relevan dalam sistem klasifikasi sederhana atau untuk baseline awal dalam pengembangan model.

SVM menunjukkan keseimbangan antara akurasi dan efisiensi. Dengan menggunakan TF-IDF sebagai representasi fitur, SVM mampu membedakan kelas sentimen secara lebih presisi. Hasil yang stabil pada berbagai ukuran dataset menunjukkan bahwa SVM cocok untuk sistem produksi yang membutuhkan performa handal dalam waktu pelatihan yang relatif singkat.

LSTM, dengan arsitektur berbasis waktu dan kemampuannya dalam memahami konteks jangka panjang, memberikan hasil evaluasi terbaik. Akurasi dan F1-score tinggi menunjukkan bahwa model ini paling efektif dalam menangani teks dengan kompleksitas tinggi, seperti penggunaan bahasa informal, konteks multi-kalimat, dan makna implisit.

Penelitian sebelumnya oleh Siswoyo et al. (2022) dalam konteks prediksi kebangkrutan perbankan juga menunjukkan bahwa pendekatan ensemble learning mampu meningkatkan akurasi klasifikasi hingga 97% jika dibandingkan dengan model individual seperti SVM atau Logistic Regression. Hasil tersebut sejalan dengan

temuan dalam penelitian ini, di mana model LSTM, yang dalam bentuk lain dapat menjadi bagian dari arsitektur ensemble, memberikan performa yang paling tinggi dalam klasifikasi sentimen pelanggan.

Implikasinya, akurasi tinggi tidak hanya tergantung pada pemilihan model tunggal yang kuat seperti LSTM, tetapi juga pada bagaimana model tersebut diintegrasikan atau dioptimalkan, misalnya melalui teknik ensemble seperti bagging atau voting classifier. Oleh karena itu, penelitian lanjutan disarankan untuk mengeksplorasi kombinasi LSTM dalam skema ensemble sebagaimana yang telah berhasil diterapkan dalam domain finansial [12].

Untuk memahami lebih dalam kinerja model, dilakukan analisis terhadap kesalahan prediksi:

Naïve Bayes sering salah mengklasifikasikan tweet yang mengandung ironi atau sarkasme, karena model ini tidak mempertimbangkan konteks kalimat secara keseluruhan. Sebagai contoh, kalimat seperti “pelayanannya *luar biasa... lambatnya*” sering diklasifikasikan sebagai positif karena kata “luar biasa”.

SVM memiliki kinerja lebih baik terhadap kalimat dengan satu sentimen dominan, tetapi masih kesulitan dalam menangani kalimat campuran atau ambigu. Misalnya, “Produknya bagus, tapi pengirimannya sangat lambat” bisa diklasifikasikan berbeda tergantung kata dominan yang ditekankan dalam vektor fitur.

LSTM mampu menangkap struktur kalimat kompleks dan mengidentifikasi sentimen campuran lebih akurat. Namun, LSTM masih bisa mengalami overfitting jika tidak dilakukan regularisasi yang baik, apalagi saat dataset mengandung terlalu banyak noise atau data tidak bersih.

Tabel 2. Hasil perbandingan algoritma

Algoritma	Waktu Latih Rata-rata	Kompleksitas Model	Ketergantungan Konteks
Naïve Bayes	±1 menit	Rendah	Tidak
SVM	±5 menit	Sedang	Parsial (via TF-IDF)
LSTM	±45 menit	Tinggi	Ya (via sequence)

Untuk perbandingan waktu latih dan kompleksitas bila berdasarkan Tabel 2, terlihat bahwa penggunaan LSTM membutuhkan waktu dan sumber daya lebih besar dibanding dua model lainnya. Oleh karena itu, implementasi di lingkungan produksi perlu mempertimbangkan kapasitas server atau GPU yang tersedia.

Distribusi performa model berdasarkan jenis sentimen (positif, negatif, netral) menunjukkan bahwa: Naïve Bayes lebih akurat dalam mengklasifikasikan tweet positif, namun performanya menurun drastis pada netral, karena keterbatasan dalam mengenali nuansa yang tidak eksplisit; SVM menunjukkan distribusi yang relatif merata di ketiga kelas, terutama karena pemanfaatan bobot kata yang kuat melalui TF-IDF; LSTM unggul pada semua kelas, terutama dalam mendeteksi sentimen netral yang seringkali ambigu. Ini menunjukkan efektivitas embedding kata dalam menangkap nuansa makna kata secara kontekstual.

Penelitian ini memberikan kontribusi dengan mengevaluasi kinerja algoritma yang berbeda untuk klasifikasi sentimen media sosial, sekaligus memberikan rekomendasi algoritma terbaik untuk data teks informal. Hasilnya dapat diaplikasikan dalam strategi pemasaran digital untuk memahami persepsi pelanggan secara lebih efektif.

Berdasarkan hasil penelitian mengenai implementasi *Machine learning* untuk klasifikasi sentimen pelanggan berbasis data media sosial, terdapat beberapa hal yang menjawab rumusan masalah yang telah diajukan sebelumnya.

Penggunaan Algoritma *Machine learning* untuk Klasifikasi Sentimen Pelanggan di Media Sosial: Algoritma *Machine learning* terbukti dapat digunakan secara efektif untuk mengklasifikasikan sentimen pelanggan di media sosial. Dengan data yang dikumpulkan secara real-time melalui platform seperti Twitter, algoritma seperti *Naïve Bayes*, *Support Vector Machine* (SVM), dan *Long Short-Term Memory* (LSTM) mampu memproses data teks mentah dan mengklasifikasikan sentimen pelanggan ke dalam tiga kategori utama: positif, negatif, dan netral. Proses pra-pemrosesan, seperti tokenisasi, *Stemming*, dan penghapusan *Stop Words*, menjadi langkah krusial untuk meningkatkan akurasi dan efisiensi algoritma. Hal ini menunjukkan bahwa pendekatan berbasis *Machine learning* sangat relevan untuk menghadapi tantangan analisis sentimen dalam konteks media sosial yang memiliki karakteristik data tidak terstruktur dan sering kali menggunakan bahasa informal.

Algoritma yang Memberikan Hasil Terbaik dalam Klasifikasi Sentimen: Penelitian ini menemukan bahwa performa algoritma bervariasi tergantung pada kompleksitas data dan ukuran dataset yang digunakan. Dari tiga algoritma yang diuji, LSTM menunjukkan hasil terbaik dengan akurasi mencapai 91,2%, diikuti oleh SVM dengan akurasi 84,7%, dan *Naïve Bayes* dengan akurasi 78,5%. LSTM unggul dalam menangkap konteks semantik dan pola temporal antar kata dalam teks, sehingga mampu memahami data yang mengandung variasi bahasa informal, seperti slang atau akronim, yang umum di media sosial. SVM, meskipun memiliki akurasi lebih rendah dibandingkan LSTM, tetap menunjukkan kinerja yang stabil dalam menangani dataset berukuran sedang dengan

fitur berbasis TF-IDF. Sebaliknya, *Naïve Bayes*, meskipun sederhana, lebih cocok untuk dataset kecil karena keterbatasannya dalam memahami hubungan antar kata secara kompleks.

Keunggulan dan Keterbatasan Setiap Algoritma: Setiap algoritma memiliki keunggulan dan keterbatasan masing-masing. LSTM menawarkan akurasi tinggi, terutama untuk dataset besar dan kompleks, tetapi memerlukan waktu pelatihan yang lebih lama serta sumber daya komputasi yang lebih besar. SVM memberikan keseimbangan antara akurasi dan efisiensi, menjadikannya pilihan yang baik untuk aplikasi dengan keterbatasan sumber daya. Sementara itu, *Naïve Bayes* merupakan algoritma yang cepat dan sederhana, tetapi kurang mampu menangkap kompleksitas data teks tidak terstruktur [13]. Oleh karena itu, pemilihan algoritma dalam implementasi dunia nyata harus mempertimbangkan kebutuhan spesifik, seperti ukuran dataset, kompleksitas data, dan ketersediaan sumber daya komputasi.

Kontribusi Penelitian terhadap Penerapan *Machine learning* dalam Strategi Pemasaran Digital: Penelitian ini memberikan kontribusi penting dalam mendukung penerapan *Machine learning* untuk mendukung strategi pemasaran digital yang lebih efektif. Dengan memanfaatkan algoritma terbaik untuk analisis sentimen pelanggan, perusahaan dapat memahami opini konsumen secara lebih mendalam dan mengambil langkah-langkah yang tepat dalam menyusun strategi pemasaran [14]. Hasil penelitian ini juga memberikan wawasan tentang pentingnya memilih algoritma yang sesuai dengan karakteristik data yang dianalisis, sehingga pengambilan keputusan berbasis data dapat dilakukan secara lebih akurat dan efisien [15].

4. Kesimpulan

Secara keseluruhan, penelitian ini berhasil menunjukkan bahwa algoritma *machine learning*, khususnya LSTM, merupakan alat yang sangat efektif untuk klasifikasi sentimen pelanggan berbasis data media sosial. Dengan perkembangan teknologi dan ketersediaan data yang semakin luas, implementasi algoritma *Machine learning* dalam analisis sentimen pelanggan akan menjadi salah satu pilar utama dalam pengambilan keputusan strategis di berbagai industri, terutama pemasaran digital. Penelitian ini juga membuka peluang bagi pengembangan lebih lanjut, seperti pengujian pada dataset dengan bahasa lokal atau penggunaan algoritma berbasis transformer, seperti BERT (*Bidirectional Encoder Representations from Transformers*), untuk meningkatkan kinerja klasifikasi sentimen di masa mendatang.

Daftar Rujukan

- [1] B. Wang and Y. Wang, "Sentiment analysis in social networks: Methods and applications. Amsterdam: Elsevier," 2016.
- [2] B. Liu, "Sentiment analysis and opinion mining. Synthesis Lectures on Human Language Technologies," vol. 5(1), pp. 1–167, 2012.
- [3] B. Pang and L. Lee, "Opinion mining and sentiment analysis. Foundations and Trends in Information Retrieval," vol. 2(1-2), pp. 1–135, 2008.
- [4] E. Cambria, B. Schuller, Y. Xia, and C. Havasi, "New avenues in opinion mining and sentiment analysis. IEEE Intelligent Systems," vol. 28(2), pp. 15–21, 2013.
- [5] M. Naufal, "Implementasi machine learning dalam klasifikasi data teks. Jakarta: Universitas Teknologi Indonesia., " 2019.
- [6] S. M. Sebastian and A. Himawan, "Machine learning-based sentiment analysis: A review and application to customer feedback. Journal of Data Science Applications," vol. 5(3), pp. 44–56, 2021.
- [7] A. L. Maas, R. E. Daly, P. T. Pham, D. Huang, A. Y. Ng, and C. Potts, "Learning word vectors for sentiment analysis. Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies," pp. 142–150, 2011.
- [8] T. Mikolov, I. Sutskever, K. Chen, G. Corrado, and J. Dean, "Efficient estimation of word representations in vector space. Proceedings of the International Conference on Learning Representations (ICLR).," 2013.
- [9] Y. Goldberg and O. Levy, "Word2Vec explained: Deriving Mikolov et al.'s negative-sampling word-embedding method," vol. 1402.3722, 2014.
- [10] S. Zhang, H. Wang, and J. Liu, "Deep learning for sentiment analysis: A survey. New York: Springer," 2020.
- [11] S. Hochreiter and J. Schmidhuber, "Long short-term memory. Neural Computation," vol. 9(8), pp. 1735–1780, 1997.
- [12] B. Siswoyo, Z. A. Abas, A. N. C. Pee, R. Komalasari, and N. Suyatna, "Ensemble machine learning algorithm optimization of bankruptcy prediction of bank," *IAES International Journal of Artificial Intelligence*, vol. 11, no. 2, pp. 679–686, Jun. 2022, doi: 10.11591/ijai.v11.i2.pp679-686.
- [13] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system. Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining," pp. 785–794, 2016.
- [14] D. E. Rumelhart, G. E. Hinton, and R. J. Williams, "Learning representations by back-propagating errors. Nature," vol. 323(6088), pp. 533–536.
- [15] D. Tang, B. Qin, and T. Liu, "Document modeling with gated recurrent neural network for sentiment classification. Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics," pp. 1422–1432, 2015.