

Klasifikasi Dukungan Politik Berdasarkan Analisis Heatmap Aktivitas Komentar di Facebook Menggunakan CNN 1D

Wahyudin Ahmadi^{1*}, Bramasto Daru Satrio², Suparno³

¹ Pendidikan Bahasa Inggris, Universitas Panca Sakti Bekasi

² Sistem Informasi, Universitas Panca Sakti Bekasi

³ Teknik Informatika, Universitas Panca Sakti Bekasi

^{1*} majnunahmadi@gmail.com, ² bramasto56964@gmail.com, ³ parno9054@gmail.com

Abstrak

Analisis sentimen opini publik di media sosial menghadapi tantangan keterbatasan data (*small dataset*) dan ketidakseimbangan kelas (*class imbalance*) yang dapat menyebabkan *overfitting* pada model *deep learning*. Penelitian ini mengkaji performa arsitektur *Convolutional Neural Network* 1-Dimensi (CNN-1D) untuk klasifikasi sentimen opini publik terhadap kinerja pemerintahan Presiden Prabowo Subianto ke dalam tiga kategori: Netral, Keluhan/Aduan, dan Apresiasi. Dataset terdiri dari 434 komentar media sosial dengan dominasi kategori keluhan. Untuk mengatasi *overfitting* dan bias kelas, penelitian ini menerapkan reduksi dimensi *embedding*, regularisasi *dropout* (0.5), serta pembobotan kelas (*class weight*) berdasarkan *inverse class frequency*. Pengujian dilakukan dengan pembagian data 80:20 dan *early stopping* untuk mencegah *overfitting*. Hasil evaluasi menunjukkan akurasi 70.11% dengan presisi tertinggi pada kategori keluhan (81.35%), sementara kategori netral masih mengalami kesalahan klasifikasi akibat keterbatasan sampel dan tumpang tindih semantik. Pembahasan menyoroti perbandingan teoretis CNN-1D dengan Support Vector Machine (SVM) pada dataset berskala kecil serta analisis kesalahan prediksi melalui grafik *heatmap confusion matrix*.

Kata Kunci: Analisis Sentimen, CNN 1D, *Class weight*, Small Dataset, Opini Publik.

PENDAHULUAN

Opini publik di media sosial merupakan instrumen penting bagi evaluasi kebijakan pemerintahan. Pada era pemerintahan Presiden Prabowo Subianto, berbagai komentar publik yang muncul di laman Facebook resmi beliau (@PrabowoSubianto) merefleksikan beragam sentimen mulai dari keluhan administratif, apresiasi terhadap program kerja, hingga respon netral berupa diskusi kebijakan umum. Sebagai salah satu platform dengan jangkauan luas di Indonesia, akun Facebook resmi Presiden menjadi ruang publik digital bagi masyarakat untuk menyampaikan tanggapan, aspirasi, maupun kritik terhadap berbagai kebijakan yang dijalankan pemerintah. Analisis terhadap opini publik yang terekam di laman resmi ini menjadi penting karena dapat memberikan gambaran persepsi masyarakat secara luas terhadap kinerja pemerintah, sebagaimana telah diteliti dalam berbagai studi analisis sentimen kebijakan publik di Indonesia (Nugraha & al., 2025; Riwaldi & Aripin, 2026).

Analisis sentimen merupakan bagian dari *opinion mining* yang bertujuan untuk mengidentifikasi dan mengklasifikasikan polaritas opini dalam teks (Pang & Lee, 2008). Dalam konteks media sosial, volume data yang besar membuat proses analisis secara manual menjadi tidak efisien, sehingga diperlukan pendekatan otomatis berbasis *Machine Learning* maupun *Deep Learning*. Perkembangan terkini menunjukkan bahwa pendekatan *Deep Learning* seperti *Convolutional Neural Network* 1-Dimensi (CNN-1D) mulai banyak digunakan untuk klasifikasi teks di media sosial. Talai & Kherici (2023) membuktikan bahwa arsitektur CNN-1D efektif untuk berbagai tugas klasifikasi teks, meskipun menghadapi tantangan *overfitting* pada dataset tertentu. Studi lain di Indonesia juga mengimplementasikan CNN untuk analisis sentimen *hashtag* viral di media sosial (Fahrezy et al., 2025), serta menggunakan arsitektur hybrid CNN-BiLSTM untuk menganalisis opini publik terhadap program pemerintah pada *dataset multi-platform* (Riwaldi & Aripin, 2026).

Keunggulan utama CNN dalam klasifikasi teks terletak pada kemampuannya mengekstrak fitur lokal (n-gram) secara otomatis melalui operasi konvolusi, tanpa memerlukan konstruksi fitur manual yang rumit. Penelitian oleh Tan et al. (2022) memperkenalkan *Dynamic Embedding Projection-Gated Convolutional Neural Network* (DEP-CNN) yang menunjukkan bahwa pendekatan berbasis CNN dengan mekanisme *gating* mampu meningkatkan akurasi klasifikasi teks secara signifikan pada berbagai dataset benchmark. Kemampuan ini diperkuat dengan penggunaan representasi kata berbasis *embedding* yang memungkinkan model memahami hubungan semantik antar kata secara lebih mendalam Asudani et al. (2023) dalam tinjauan komprehensifnya menyatakan bahwa pemilihan model *word embedding* yang tepat sangat krusial dalam menentukan kinerja analisis teks berbasis *deep learning*. Penelitian lain oleh Huana et al. (2022) juga menunjukkan bahwa kombinasi CNN dengan mekanisme *attention* pada lapisan *embedding* mampu mengekstrak fitur semantik lokal dan global secara simultan untuk meningkatkan akurasi klasifikasi teks. Di Indonesia, penerapan CNN untuk analisis sentimen opini publik terhadap layanan publik telah menunjukkan hasil yang menjanjikan dengan akurasi mencapai 87% pada klasifikasi multi-kelas (Rohalia & Rahmika, 2026).

Meskipun demikian, penggunaan CNN pada *dataset* berukuran kecil memiliki tantangan tersendiri, terutama terkait risiko *overfitting*. Penelitian tentang penanganan *overfitting* pada klasifikasi teks berbasis *neural networks* dengan dataset terbatas (250-650 sampel) menunjukkan bahwa teknik regularisasi dan *dropout* efektif dalam mengurangi *overfitting* (Ilyas et al., 2024). Teknik *dropout* yang diperkenalkan oleh Srivastava et al. (2014) bekerja dengan menonaktifkan *neuron* secara acak selama proses pelatihan, sehingga mencegah ketergantungan antar koneksi saraf (*co-adaptation*) yang memicu *overfitting*.

Perkembangan lebih lanjut dari teknik *dropout* seperti AutoDropout oleh Pham & al. (2021) yang mempelajari pola *dropout* secara adaptif menunjukkan hasil yang lebih baik dibandingkan *dropout* konvensional pada berbagai tugas, termasuk *language modeling* dan klasifikasi teks. Penelitian oleh Wen & al. (2021) mengusulkan algoritma *maximum pooling dropout* yang dikombinasikan dengan *weight attenuation* untuk mengatasi *overfitting* pada model CNN. Metode ini bekerja dengan menonaktifkan sebagian *neuron* pada lapisan *pooling* (dikenal sebagai *maximum pooling dropout*) serta menerapkan regularisasi berbasis *weight attenuation* saat jaringan belajar dari kesalahan prediksi (*backpropagation*). Hasil eksperimen menunjukkan bahwa pendekatan ini mampu menurunkan *error rate* klasifikasi data set lebih dari 10% dibandingkan metode lainnya. Hal ini menunjukkan bahwa metode regularisasi terus berkembang untuk meningkatkan generalisasi model pada berbagai arsitektur jaringan saraf.

Selain *overfitting*, tantangan lain yang sering dihadapi dalam analisis sentimen media sosial adalah ketidakseimbangan kelas (*class imbalance*), di mana satu kategori sentimen mendominasi secara signifikan dibandingkan kategori lainnya. Yilmaz et al. (2021) memperkenalkan metode *dynamic weighting* untuk mengatasi label *imbalance* dalam analisis sentimen multi-bahasa, sementara pendekatan lain menggunakan SMOTE dan ROS untuk menyeimbangkan data sebelum pelatihan model (Boulesnane et al., 2022). Penanganan ketidakseimbangan data menjadi krusial karena model cenderung bias terhadap kelas mayoritas jika tidak diberikan perlakuan khusus.

Berdasarkan latar belakang tersebut, penelitian ini mengkaji performa algoritma CNN-1D untuk mengklasifikasikan sentimen opini publik terhadap kinerja pemerintahan Presiden Prabowo Subianto ke dalam tiga kategori: Netral, Keluhan/Aduan, dan Apresiasi. Dataset yang digunakan berjumlah 434 komentar media sosial yang didominasi oleh kategori keluhan. Penelitian ini secara khusus meneliti bagaimana teknik optimasi berupa penyesuaian dimensi embedding, regularisasi dropout, dan pembobotan kelas (*class weight*) dapat diterapkan agar model CNN-1D tetap optimal dan tidak mengalami *overfitting* ekstrem pada dataset berukuran kecil.

METODE

a. Preprocessing Data Teks

Data teks mentah dari media sosial memiliki tingkat derau (noise) yang tinggi. Preprocessing yang dilakukan meliputi:

1. *Cleaning*: Menghapus URL, karakter non-ASCII (*emoji*), tanda baca, dan angka.
2. *Case Folding*: Mengubah seluruh teks menjadi huruf kecil (*lowercase*).
3. *Normalisasi Slang*: Mengganti kata-kata tidak baku (contoh: *gk* → tidak, *pln* → perusahaan listrik negara, *yth* → yang terhormat) menggunakan kamus *slang* terstruktur.
4. *Stopwords Removal*: Menghapus kata umum yang tidak memiliki bobot semantik penting, dengan tetap mempertahankan kata negasi kritis seperti tidak, bukan, belum, jangan untuk menjaga kualitas sentimen negatif.

b. Tokenisasi dan Padding

Teks bersih dikonversi menjadi urutan angka (*sequence*) menggunakan *Tokenizer* dengan batas kosakata (*vocabulary size*) sebesar 5.000 kata unik. Kalimat kemudian disamakan panjangnya (*padding*) dengan batas maksimum panjang kata (*MAX_LEN*) sebesar 100 kata secara *post-padding*.

c. Arsitektur Model CNN-1D

CNN-1D dipilih karena kemampuannya dalam mengekstrak fitur lokal (*n-gram*) secara otomatis melalui operasi konvolusi pada satu dimensi sekuensial (Kiranyaz et al., 2021; Talai & Kherici, 2023). Berbeda dengan CNN-2D yang umum digunakan untuk data spasial seperti gambar, CNN-1D lebih sesuai untuk data teks karena kernel konvolusi bergerak sepanjang dimensi waktu (urutan kata) dan mencakup seluruh dimensi saluran (*channel*) embedding (Kiranyaz et al., 2021). Struktur ini memungkinkan *kernel* menangkap pola hubungan antarkata yang berdekatan secara lebih ekspresif dibandingkan CNN-2D yang hanya mencakup sebagian dimensi saluran (Chen et al., 2019; Habib et al., 2022).

Model CNN-1D dirancang secara sekuensial dengan struktur sebagai berikut:

1. *Embedding Layer*: Memetakan indeks kata ke dalam ruang vektor berdimensi 50.
2. *Conv1D Layer*: Terdiri atas 64 filter dengan ukuran kernel 5 dan fungsi aktivasi ReLU untuk mengekstrak fitur lokal kalimat.
3. *Global MaxPooling1D*: Mereduksi dimensi dengan mengambil nilai aktivasi terbesar dari setiap filter hasil konvolusi.
4. *Dense Layer (Hidden)*: Memiliki 32 unit dengan fungsi aktivasi ReLU.
5. *Dropout Layer*: Diterapkan dengan tingkat 0.5 (50%) sebelum lapisan klasifikasi untuk mencegah ketergantungan antarkoneksi saraf (*co-adaptation*) yang memicu *overfitting*.
6. *Dense Layer (Output)*: Memiliki 3 unit kelas (Netral, Keluhan/Aduan, Apresiasi) dengan fungsi aktivasi Softmax.

c. Penanganan Imbalance Data (Class weight)

Distribusi kelas dalam dataset sangat timpang di mana kategori Keluhan/Aduan mendominasi secara signifikan. Formula pembobotan kelas yang digunakan adalah;

$$w_j = \frac{N}{C \times n_j}$$

Dimana:

w_j adalah bobot untuk kelas j .

N adalah total sampel data training.

C adalah jumlah jumlah kelas yaitu 3.

n_j adalah jumlah sampel pada kelas j .

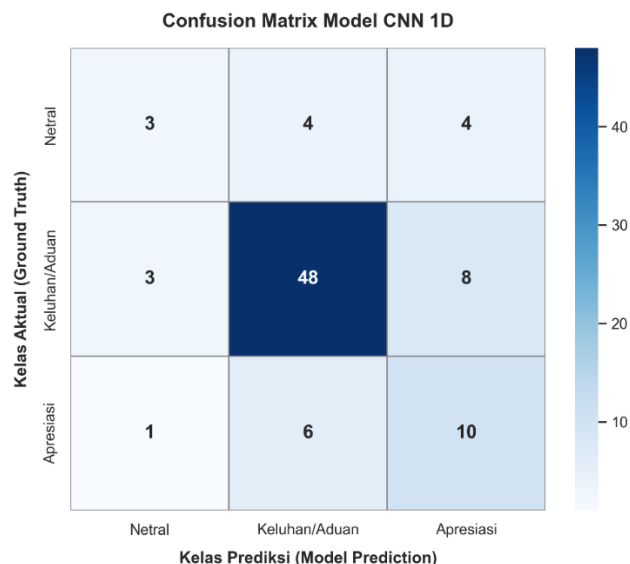
Penggunaan pembobotan kelas (*class weight*) merupakan strategi penting untuk mengatasi ketidakseimbangan data, terutama pada dataset berukuran kecil, karena memaksa model untuk memberikan perhatian lebih pada kelas minoritas.

HASIL DAN PEMBAHASAN

A. Confusion Matrix

Proses *training* dilakukan dengan ukuran *batch* 16 dan batas maksimum 30 *epoch*, dibantu oleh mekanisme *Early Stopping* berbasis *validation loss* (*patience* = 3). Model berhenti berlatih saat *validation loss* tidak lagi menunjukkan penurunan untuk menghindari *overfitting* yang lebih dalam.

Pengujian model menggunakan 20% data validasi (87 sampel) menghasilkan akurasi keseluruhan sebesar 70.11%. Kinerja model dievaluasi lebih lanjut menggunakan *Confusion Matrix* yang tersaji pada Gambar di bawah ini:



Gambar 1. Grafik *Heatmap Confusion Matrix* pada Model

A. Pembahasan

Secara teoretis, arsitektur *Deep Learning* seperti CNN-1D membutuhkan ribuan hingga jutaan sampel data untuk dapat menggeneralisasi representasi teks secara optimal. Hal ini berbeda dengan pendekatan klasifikasi tradisional seperti *Support Vector Machine* (SVM) yang dikombinasikan dengan TF-IDF. SVM bekerja dengan memetakan teks ke dalam ruang fitur vektor TF-IDF dan mencari *hyperplane* pemisah terbaik. Pada *dataset* 434 baris, SVM cenderung lebih stabil karena parameter yang dipelajari jauh lebih sedikit dibandingkan parameter bobot pada lapisan *Embedding* dan Conv1D di jaringan saraf dalam.

Namun, terdapat kelemahan pada SVM + TF-IDF, yaitu ketidakmampuannya dalam menangkap informasi semantik berbasis urutan kata (*word order*). TF-IDF menggunakan pendekatan Bag-of-Words yang mengabaikan posisi kata. Sebaliknya, CNN-1D dengan bantuan lapisan *Embedding* dimensi 50 dan filter konvolusi ukuran kernel 5 mampu mendeteksi pola hubungan antarkata yang berdekatan (*local n-grams*), seperti frasa negasi ("tidak setuju", "belum gratis") secara lebih kontekstual.

1. Analisis Faktor Keberhasilan Model CNN-1D pada Sampel Terbatas

Meskipun dilatih dengan *dataset* berukuran kecil (434 komentar), model CNN-1D dalam penelitian ini tidak mengalami *overfitting* yang merusak performa berkat tiga intervensi arsitektur:

- Reduksi Dimensi *Embedding*:** Penggunaan dimensi *embedding* sebesar 50 (turun dari standar umum 100 atau 300) berhasil membatasi kapasitas ruang pencarian parameter bobot, sehingga menekan kecenderungan model untuk menghafal pola data latih secara spesifik. *Embedding* dimensi tinggi pada *dataset* kecil sangat rentan menyebabkan *overfitting* karena model dapat "menghafal" asosiasi yang tidak relevan.
- Penerapan Dropout Tinggi (0.5):** Lapisan dropout memaksa model untuk mematikan setengah *neuron* secara acak pada setiap langkah iterasi training. Hal ini memaksa jaringan mempelajari representasi fitur yang

- redundan dan lebih tangguh (*robust*), bukan mengandalkan jalur neuron spesifik yang rentan mengalami pencocokan berlebih (*overfitting*).
- c. Pembobotan Kelas (*Class weight*): Tanpa *class weight*, model akan bias ke arah kelas mayoritas (Keluhan/Aduan) karena pembaruan *gradien loss* didominasi oleh kelas tersebut. Dengan mengalikan loss dari kelas minoritas (Netral dan Apresiasi) dengan bobot yang lebih besar, gradien terpaksa bergeser untuk mengenali pola-pola khas kelas minoritas.
2. Analisis Kesalahan Prediksi (*Error Analysis*) melalui *Confusion Matrix*
- Berdasarkan Gambar 1, kelas Keluhan/Aduan (1) memiliki akurasi individual tertinggi, di mana 48 dari 59 data aktual berhasil diprediksi dengan benar (81.35%). Keberhasilan ini disebabkan oleh melimpahnya data latihan untuk kelas ini yang kaya akan kosakata negatif yang khas (contoh: kece bong, turunkan, mengganggu, iblis).
- Sebaliknya, performa terburuk terjadi pada kelas Netral (0), di mana hanya 3 dari 11 data aktual yang terklasifikasi dengan benar (27.27%). Sisanya salah diklasifikasikan ke dalam kategori Keluhan/Aduan (4 sampel) dan Apresiasi (4 sampel).
- Setelah dianalisis secara semantik, kesalahan ini terjadi karena:
- a. Komentar netral pada objek evaluasi pemerintah sering kali menggunakan kosakata yang tumpang tindih dengan keluhan atau apresiasi. Sebagai contoh, kalimat "kebijakan ini harus diawasi dengan ketat" secara semantik bernada netral/saran, namun model CNN-1D cenderung mengasosiasikan kata "diawasi" dan "ketat" ke dalam kluster keluhan karena kemiripan konteks penggunaan kata tersebut pada data training kategori keluhan.
 - b. Ukuran sampel kelas netral yang terlampau sedikit membuat model kesulitan mengekstrak fitur pembeda yang unik untuk kelas netral.

Pada kelas Apresiasi (2), model berhasil mengklasifikasikan 10 dari 17 data aktual dengan benar (58.82%), lebih rendah dibandingkan kelas Keluhan (81.35%) namun lebih tinggi dibandingkan kelas Netral (27.27%). Hal ini disebabkan oleh fenomena bahasa berupa penggunaan gaya bahasa sarkasme atau kalimat apresiasi yang diawali dengan menceritakan masalah terlebih dahulu, sehingga filter konvolusi berukuran kernel 5 lebih dominan menangkap kluster kata bermakna negatif di awal kalimat ketimbang pujian di akhir kalimat.

KESIMPULAN

Penelitian ini membuktikan bahwa arsitektur Deep Learning CNN-1D dapat diimplementasikan pada dataset berskala kecil (434 komentar) dengan akurasi validasi sebesar 70.11%. Meskipun secara teoretis algoritma klasik seperti SVM + TF-IDF lebih efisien untuk jumlah sampel terbatas, CNN-1D menawarkan keunggulan dalam ekstraksi fitur lokal (n-gram) secara otomatis tanpa memerlukan konstruksi fitur manual yang rumit.

Penggunaan teknik regularisasi berupa *dropout* 0.5, reduksi dimensi embedding (dari 100/300 menjadi 50), serta penerapan *class weight* terbukti krusial dalam menahan laju *overfitting* dan bias kelas pada data yang sangat tidak seimbang. Ketiga teknik ini berhasil memungkinkan model CNN-1D tetap bekerja optimal meskipun dilatih pada data dengan jumlah terbatas dan dominasi kelas keluhan yang signifikan.

Berdasarkan analisis confusion matrix, model menunjukkan performa terbaik pada kelas Keluhan/Aduan dengan akurasi 81.35%, yang disebabkan oleh melimpahnya data latih untuk kelas tersebut. Sebaliknya, kelas Netral menjadi tantangan terbesar dengan akurasi hanya 27.27%, akibat keterbatasan sampel dan tumpang tindih semantik dengan kelas lain. Kelas Apresiasi berada di posisi tengah dengan akurasi 58.82%, di mana kesalahan klasifikasi banyak terjadi akibat fenomena sarkasme dan struktur kalimat yang diawali dengan narasi negatif sebelum pujian di akhir kalimat.

Penelitian ini juga menyoroti perbedaan pendekatan antara CNN-1D dan SVM+TF-IDF dalam menangani data teks. SVM lebih stabil secara parametrik pada dataset kecil, namun CNN-1D unggul dalam menangkap hubungan semantik antar kata melalui representasi embedding dan operasi konvolusi. Hal ini menunjukkan bahwa deep learning tetap memiliki nilai tambah meskipun pada dataset terbatas, selama didukung oleh teknik regularisasi yang tepat.

DAFTAR PUSTAKA

- Asudani, D. S., Nagwani, N. K., & Singh, P. (2023). Impact of word embedding models on text analytics in deep learning environment: a review. *Artificial Intelligence Review*, 56(9), 10345. <https://doi.org/10.1007/s10462-023-10419-1>
- Boulesnane, A., Meshoul, S., & Aouissi, K. (2022). Influenza-like Illness Detection from Arabic Facebook Posts Based on Sentiment Analysis and 1D Convolutional Neural Network. *Mathematics*, 10(21), 4089. <https://doi.org/10.3390/math10214089>
- Chen, Z., Feng, A., & He, J. (2019). Text sentiment classification based on 1D convolutional hybrid neural network. *Journal of Computer Applications*, 39(7), 1936–1941. <https://doi.org/10.11772/j.issn.1001-9081.2018122477>
- Fahrezy, A. R., Refina, S. A., Chandra, J. G., & Yora, S. (2025). Sentiment Analysis of Indonesian Hashtag #kaburajadulu Using CNN. *2025 International Conference on Information Technology and Computing (ICITCOM)*, 222–227. <https://doi.org/10.1109/ICITCOM64883.2025.11265040>

- Habib, M. A., Akhand, M. A. H., & Kamal, M. A. S. (2022). Emotion Recognition from Microblog Managing Emoticon with Text and Classifying using 1D CNN. *Journal of Computer Science*, 18(12), 1170–1178. <https://doi.org/10.3844/jcssp.2022.1170.1178>
- Huana, H., Guoa, Z., Caib, T., & Heb, Z. (2022). A text classification method based on a convolutional and bidirectional long short-term memory model. *Connection Science*, 34(1), 2108–2124. <https://doi.org/10.1080/09540091.2022.2098926>
- Ilyas, R., Kasyidi, F., & Eriyadi, M. N. (2024). Penanganan Overfitting pada Klasifikasi Berita Hoax berbasis Neural Networks dengan Dropout dan Regularization. *Jurnal Teknologi Dan Komputer*, 10(2). <https://doi.org/10.31294/jtk.v10i2.23121>
- Kiranyaz, S., Avci, O., Abdeljaber, O., Ince, T., Gabbouj, M., & Inman, D. J. (2021). 1D convolutional neural networks and applications: A survey. *Mechanical Systems and Signal Processing*, 151, 107398. <https://doi.org/10.1016/j.ymssp.2020.107398>
- Nugraha, W., & al., et. (2025). Public Sentiment Toward the Indonesian Capital Relocation Policy on X Using a BiLSTM-CNN Model. *Telematika*, 20(2). <https://doi.org/10.61769/telematika.v20i2.796>
- Pang, B., & Lee, L. (2008). Opinion Mining and Sentiment Analysis. *Foundations and Trends in Information Retrieval*, 2(1–2), 1–135.
- Pham, H., & al., et. (2021). AutoDropout: Learning Dropout Patterns to Regularize Deep Networks. *AAAI 2021*. <https://arxiv.org/abs/2101.01761>
- Riwalidi, M. R. F., & Aripin, A. (2026). Hybrid CNN-BiLSTM untuk Analisis Sentimen Multi-Platform terhadap Insiden Keamanan Pangan Program Makan Bergizi Gratis. *BIT (Budi Luhur Information Technology)*, 8(1). <https://doi.org/10.32493/bit.v8i1.9896>
- Rohalia, A., & Rahmika, A. R. (2026). Leveraging Convolutional Neural Networks and Random Forests for Advanced Sentiment Classification of Social Media Responses on Public Services. *JAIC (Journal of Artificial Intelligence and Computing)*, 10(1). <https://doi.org/10.30871/jaic.v10i1.11965>
- Srivastava, N., Hinton, G., Krizhevsky, A., Sutskever, I., & Salakhutdinov, R. (2014). Dropout: A Simple Way to Prevent Neural Networks from Overfitting. *Journal of Machine Learning Research*, 15(1), 1929–1958.
- Talai, Z., & Kherici, N. (2023). Text Classification and Sentiment Analysis using 1D CNN. *Information Processing at the Digital Age Journal*, 27(2), 46–51.
- Tan, Z., Chen, J., Kang, Q., Zhou, M., Abusorrah, A., & Sedraoui, K. (2022). Dynamic Embedding Projection-Gated Convolutional Neural Networks for Text Classification. *IEEE Transactions on Neural Networks and Learning Systems*, 33(3), 973–982. <https://doi.org/10.1109/TNNLS.2020.3036192>
- Wen, L., & al., et. (2021). Maximum Pooling Dropout with Weight Attenuation for Overfitting in CNN. *Wireless Communications & Mobile Computing*. <https://doi.org/10.1155/2021/8493795>
- Yilmaz, S. F., Kaynak, E. B., Koç, A., Dibeklioglu, H., & Kozat, S. S. (2021). Multi-Label Sentiment Analysis on 100 Languages with Dynamic Weighting for Label Imbalance. *IEEE Transactions on Neural Networks and Learning Systems*, 33(1), 331–343. <https://doi.org/10.1109/TNNLS.2021.3094304>