

OPTIMASI PROSES KLASIFIKASI TOPIK BERITA BERBAHASA INDONESIA MENGGUNAKAN MODEL NLP BERBASIS ONNX RUNTIME

Tio Risnanto, Edhy Poerwandono

Teknik Informatika Stikom CKI (Cipta Karya Informatika)
Jl. Radin Inten II No.8 5, RT.5/RW.14, Duren Sawit, Kec. Duren Sawit,
Kota Jakarta Timur, Daerah Khusus Ibukota Jakarta 13440 Telp (021) 21285900
tiolast63@gmail.com

ABSTRAK

Perkembangan berita digital yang pesat memunculkan tantangan baru dalam klasifikasi topik berita secara akurat, terutama ketika satu artikel dapat memiliki lebih dari satu kategori (multi-label). Proses klasifikasi manual terbukti tidak efisien dan rawan kesalahan. Penelitian ini mengusulkan sistem klasifikasi multi-label berita berbasis Term Frequency–Inverse Document Frequency (TF-IDF) dan Logistic Regression yang dioptimalkan menggunakan ONNX Runtime untuk meningkatkan kecepatan inferensi. Dataset penelitian terdiri dari 10.058 artikel berita dari Kompas.com yang telah melalui tahap pra-pemrosesan dan pelabelan multi-label. Model dievaluasi menggunakan metrik F1-score mikro dan makro, serta perbandingan waktu inferensi antara model Scikit-learn (.pkl) dan model yang telah dioptimasi ONNX Runtime (.onnx). Hasil menunjukkan bahwa model ONNX mempertahankan akurasi yang sama (F1-mikro = 0,648; F1-makro = 0,402) dengan pengurangan waktu inferensi dari 0,88 detik menjadi 0,44 detik. Hal ini membuktikan bahwa ONNX Runtime dapat secara efektif meningkatkan efisiensi model NLP ringan untuk implementasi skala besar.

Kata kunci : Klasifikasi Teks, Multi-Label, TF-IDF, Logistic Regression, ONNX Runtime, Optimasi NLP.

1. PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi yang berlangsung pesat telah membawa perubahan signifikan pada pola konsumsi informasi masyarakat. Jika dahulu masyarakat mengandalkan surat kabar atau siaran berita televisi, kini media digital menjadi pilihan utama karena mampu menyajikan informasi secara instan dan waktu nyata [1].

Pemanfaatan media digital yang semakin meluas memunculkan tantangan baru, salah satunya adalah lonjakan jumlah berita digital yang dipublikasikan setiap hari oleh berbagai portal berita. Banyak di antaranya memiliki keterkaitan dengan lebih dari satu kategori topik [2].

Masuknya ribuan berita online setiap hari sering kali menimbulkan ketidaktepatan dalam proses pengkategorian. Ketidaksesuaian ini dapat menghambat efisiensi pencarian informasi dan menurunkan kualitas pengalaman pengguna, sebab berita yang ditampilkan tidak selalu sesuai dengan minat mereka. Umumnya, kesalahan ini disebabkan oleh keterlibatan manusia dalam proses klasifikasi, yang rawan terhadap subjektivitas dan ketidakkonsistenan. Oleh karena itu, diperlukan sistem klasifikasi otomatis yang mampu mengelompokkan berita secara akurat dan konsisten [3].

Natural Language Processing (NLP) menjadi salah satu pendekatan utama dalam mengatasi permasalahan tersebut. Model deep learning seperti Convolutional Neural Network (CNN), Recurrent Neural Network (RNN), dan Long Short-Term Memory (LSTM) telah diterapkan untuk klasifikasi teks berita berbahasa Indonesia.

Salah satu metode untuk meningkatkan efisiensi inferensi model klasifikasi teks adalah dengan memanfaatkan ONNX Runtime, sebuah framework eksekusi lintas-platform yang dirancang untuk menjalankan model machine learning secara optimal. ONNX Runtime dapat mempercepat eksekusi model berbasis scikit-learn, seperti kombinasi Term Frequency–Inverse Document Frequency (TF-IDF) dan Logistic Regression, dengan waktu inferensi lebih singkat, konsumsi memori rendah, serta dukungan lintas-perangkat (CPU maupun GPU).

Framework ini menyediakan berbagai teknik optimasi, termasuk *graph simplification*, *node fusion*, hingga *runtime acceleration*, yang dapat mempercepat proses prediksi bahkan pada model ringan. Meskipun sebagian besar penelitian terdahulu fokus pada penerapan ONNX Runtime pada model besar seperti BERT, kajian terkait efisiensi ONNX Runtime pada model berbasis TF-IDF untuk klasifikasi berita berbahasa Indonesia masih jarang dilakukan.

Penelitian ini bertujuan untuk mengisi kekosongan tersebut dengan merancang sistem klasifikasi topik berita digital multi-label berbasis TF-IDF + Logistic Regression yang dioptimalkan menggunakan ONNX Runtime. Pendekatan ini diharapkan mampu mempertahankan akurasi yang layak sekaligus memberikan peningkatan signifikan pada efisiensi waktu inferensi, sehingga dapat diimplementasikan secara praktis pada berbagai platform, termasuk portal berita digital, dashboard redaksi, dan aplikasi monitoring informasi berbasis teks.

2. TINJAUAN PUSTAKA

Pada penelitian dengan judul “Klasifikasi Hoax Berita Politik Menggunakan Algoritma Long Short-Term Memory (LSTM) dengan Penambahan Fitur Embedding Global Vector (GloVe)”, hasil penelitian tersebut dijelaskan bahwa penambahan fitur embedding GloVe ke model LSTM signifikan meningkatkan performa klasifikasi berita hoax politik. Pendekatan ini efektif dalam membedakan berita asli dan berita hoax dengan tingkat keakuratan yang tinggi, sehingga dapat membantu mengurangi penyebaran berita palsu di media sosial dan platform berita [4].

Pada penelitian dengan judul “Penerapan Metode Multinomial Naïve Bayes untuk Klasifikasi Judul Berita Clickbait dengan Term Frequency - Inverse Document Frequency”, hasil penelitian tersebut dijelaskan bahwa kendala utama adalah algoritma ini kurang akurat karena berdasarkan frekuensi kata saja, sehingga tidak selalu mampu menangkap konteks sebenarnya dari berita. Kesimpulan bahwa metode ini cukup baik tetapi perlu dikembangkan lagi agar akurasi menjadi lebih tinggi, misalnya dengan menambahkan analisis isi berita secara lengkap atau menggunakan teknik yang lebih kompleks [1].

Pada penelitian dengan judul “Kategorisasi Berita Multi-Label Berbahasa Indonesia Menggunakan Algoritma Random Forest”, hasil penelitian tersebut dijelaskan bahwa Random Forest terbukti efektif untuk klasifikasi multi-label berita berbahasa Indonesia. Proses stemming meningkatkan performa klasifikasi. Nilai Hamming Loss 0,126 menunjukkan peningkatan performa dibandingkan penelitian sebelumnya (yang mencapai 0,1472) [2].

Pada penelitian dengan judul “Topic Classification of Online News Articles Using Optimized Machine Learning Models”, hasil penelitian tersebut dijelaskan bahwa Hyperparameter tuning signifikan meningkatkan kinerja model, khususnya SVM, yang awalnya paling buruk tanpa tuning tetapi menjadi paling unggul setelah dioptimasi. Pendekatan ini dapat diterapkan dalam aplikasi dunia nyata untuk klasifikasi berita otomatis dengan tingkat akurasi yang dapat diterima [5].

Pada penelitian dengan judul “Comparative Analysis of TF-IDF and Hashing Vectorizer for Fake News Detection in Sindhi: A Machine Learning and Deep Learning Approach”, hasil penelitian tersebut dijelaskan bahwa teknik TF-IDF secara umum sering memberikan performa yang lebih baik daripada hashing vectorizer dalam dataset Sindhi tersebut. Model neural network (RNN) dan BERT menunjukkan performa terbaik dalam pengujian, menegaskan pentingnya pemilihan teknik feature dan model yang tepat dalam pengolahan teks berbahasa Sindhi [6].

Pada penelitian dengan judul “Fake News Classification Based on Content Level Features”, hasil penelitian dijelaskan bahwa berhasil mengimplementasikan metode ekstraksi fitur dari isi teks berita seperti tokenisasi dan penilaian

menggunakan TF-IDF, serta pemanfaatan teknik NLP dan model machine learning serta neural network, termasuk CNN dan LSTM. Hasilnya, semua model menghasilkan akurasi lebih dari 85%, dengan neural network mencapai di atas 90%, dan CNN dengan Global Maxpool sebagai model terbaik. Solusinya adalah menggabungkan teknik NLP dan pembelajaran mesin/neural network untuk meningkatkan akurasi deteksi berita palsu secara otomatis berbasis konten [7].

Pada penelitian dengan judul “News Classification using Natural Language Processing with TF-IDF and Multinomial Naïve Bayes”, hasil penelitian dijelaskan bahwa Sistem ini berhasil mencapai akurasi validasi sebesar 83% dan akurasi pengujian sebesar 90%, yang secara signifikan dapat mengurangi beban kerja manusia dan meningkatkan kecepatan serta ketepatan pengklasifikasian berita [8].

Pada penelitian dengan judul “Fake News Detection Using Optimized CNN and LSTM”, hasil penelitian dijelaskan bahwa model hybrid CNN-LSTM yang dioptimalkan mampu meningkatkan kecepatan dan akurasi deteksi berita palsu. Pendekatan ini dapat diterapkan di berbagai platform media sosial dan berita online untuk mengurangi penyebaran hoaks secara efektif [9].

Pada penelitian dengan judul “Klasifikasi Judul Berita Bahasa Indonesia Menggunakan Support Vector Machine dan Seleksi Fitur Mutual Information”, hasil penelitian dijelaskan bahwa menggunakan metode SVM dengan fitur yang dipilih melalui mutual information dan kernel RBF dengan threshold sebesar 85% mampu mencapai akurasi tertinggi, yaitu F1-score sekitar 86,15%. Solusinya adalah mengaplikasikan seleksi fitur mutual information untuk mengurangi dimensi fitur dan waktu proses, serta meningkatkan efisiensi model klasifikasi berita online dalam bahasa Indonesia [3].

Pada penelitian dengan judul “Hoax Analyzer for Indonesian News Using Deep Learning Models”, hasil penelitian dijelaskan bahwa menggunakan model deep learning yang dilatih dengan data berita berbahasa Indonesia untuk meningkatkan akurasi deteksi hoax. Kesimpulannya, model DNN mampu mengungguli classifier tradisional dalam tugas klasifikasi teks berbahasa Indonesia dan menunjukkan potensi besar dalam pengembangan alat otomatis untuk mendeteksi berita palsu di Indonesia [10].

Pada penelitian dengan judul “Performance of CNN and LSTM Model on COVID-19 News Headline Data Classification”, hasil penelitian dijelaskan bahwa CNN terbukti lebih unggul dalam mengatasi data tidak seimbang dan menghasilkan hasil klasifikasi yang lebih akurat dan efisien dibandingkan LSTM. Disarankan penggabungan kedua model (CNN dan LSTM) untuk menggabungkan kekuatan keduanya dalam pengembangan model yang lebih optimal di masa mendatang. Secara umum, penelitian ini menunjukkan bahwa CNN merupakan model yang lebih efektif untuk klasifikasi berita COVID-19 dalam

skenario ketidakseimbangan data, dan metode ini dapat diterapkan pada klasifikasi berita tentang pandemi atau penyakit lain [11].

Pada penelitian dengan judul “Combination of BERT and Hybrid CNN-LSTM Models for Indonesia Dengue Tweets Classification”, hasil penelitian dijelaskan bahwa Indo-BERT + Hybrid CNN-LSTM efektif untuk klasifikasi multi-kelas tweet berbahasa Indonesia terkait DBD, meskipun data tidak seimbang. Model ini memiliki sensitivitas (recall) 89% dan presisi 93%, mengungguli pendekatan sebelumnya. Performa dapat ditingkatkan dengan penyeimbangan data (data balancing) pada penelitian selanjutnya agar akurasi di kelas minoritas meningkat[12].

Pada penelitian dengan judul “Hoax analyzer for Indonesian news using RNNs with fastText and GloVe embeddings”, hasil penelitian dijelaskan bahwa GRU lebih baik dibandingkan LSTM pada dataset dengan kemunculan data yang jarang dan teks pendek. Model bidirectional menghasilkan skor metrik yang lebih tinggi daripada model unidirectional. fastText lebih unggul daripada GloVe karena mampu menangani out-of-vocabulary (OOV) words dan kata langka. Kombinasi fastText + Bi-GRU adalah yang terbaik untuk klasifikasi berita hoaks Bahasa Indonesia pada penelitian ini. Saran untuk penelitian selanjutnya adalah menggunakan dataset asli Bahasa Indonesia tanpa terjemahan agar hasil lebih akurat[13].

3. METODE PENELITIAN

Metodologi penelitian ini mencakup serangkaian langkah yang dirancang untuk membangun, melatih, dan menguji model klasifikasi topik berita multi-label berbahasa Indonesia yang dioptimalkan menggunakan ONNX Runtime. Tahapan ini dimulai dari pengumpulan dan pemrosesan data, pelabelan multi-label, pelatihan model, konversi model ke format ONNX, hingga evaluasi kinerja model. Setiap tahap dirancang untuk memastikan kualitas data yang digunakan, ketepatan metode pelatihan, serta validitas hasil yang diperoleh.

3.1. Dataset

```
=====
1. MEMUAT DATA BERITA
=====
• Membaca file: data/berita_kompas_lengkap.json
• Jumlah data awal: 10058
```

Gambar 1. Dataset kompas.com terdiri dari 10.058 artikel

Pada Gambar 1. Diatas dataset yang digunakan dalam penelitian ini berasal dari portal berita Kompas.com dan tersimpan dalam berkas data/berita_kompas_lengkap.json. Dataset tersebut berisi 10.058 artikel berita yang masing-masing memiliki atribut utama berupa judul berita, isi berita, dan topik. Pada beberapa kasus, atribut topik berisi lebih dari satu label sehingga mendukung skenario

klasifikasi multi-label. Sebelum digunakan untuk pelatihan model, dataset ini melalui tahap pemeriksaan kualitas, meliputi deteksi nilai kosong, duplikasi, dan ketidakseimbangan distribusi label. Selanjutnya dilakukan pembersihan teks dengan menghapus tag HTML dan karakter khusus, normalisasi huruf, penghapusan stopword, dan stemming menggunakan pustaka Bahasa Indonesia. Judul dan isi berita digabungkan menjadi satu kolom teks agar informasi konteks lebih lengkap. Dataset bersih ini kemudian diubah menjadi representasi vektor menggunakan metode TF-IDF dan labelnya dikodekan dengan MultiLabelBinarizer sebelum memasuki tahap pelatihan model.

3.2. Pra-Pemrosesan Data

```
=====
2. PEMBERSIHAN DATA
=====
• Setelah drop NA: 10058 baris
• Kolom 'text' = title + content berhasil dibuat

3. STATISTIK DATA
=====
• Jumlah topik unik: 31
• Rata-rata panjang teks: 2366.31 karakter

title topic
9656 Ilmuwan Temukan Warna Baru yang Diberi Nama Ol... TREN
2941 Uji Coba Rute Sawangan-Lebak Bulus, Bus Transj... NEWS
8396 Cerita EAP Ajukan Pembatalan Nikah, Suami Meng... PROV
1498 Skor Arsenal Vs PSG 0-1: Sengatan Paris Memati... BOLA
7229 Pemprov Jakarta Sediakan Kuliner Gratis Saat P... NEWS

4. VISUALISASI DISTRIBUSI TOPIK
=====
• Grafik distribusi disimpan ke: data/distribusi_topik.png

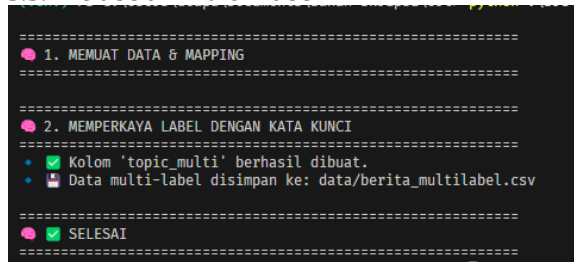
5. DETEKSI ANOMALI DATA
=====
• Teks pendek (<50 karakter):
  title topic
7 No Title No Topic
53 Kendaraan Listrik FEATURED
55 Anak Muda FEATURED
57 Dear, Pak Presiden FEATURED
• Label terlalu panjang (kemungkinan error delimiter):
Empty DataFrame
Columns: [title, topic]
Index: []
• Data disimpan ke: data/berita_cleaned.csv

=====
SELESAI DALAM 2.51 DETIK
=====
```

Gambar 2. Dataset kompas.com terdiri dari 10.058 artikel

Pada Gambar 2. Diatas Pada tahap pembersihan data, berkas mentah data/berita_kompas_lengkap.json diproses sehingga tersisa 10.058 baris setelah penghapusan nilai kosong (NA) dan pembuatan kolom baru text yang merupakan penggabungan title dan content. Analisis statistik awal mengungkap 31 topik unik dengan panjang teks rata-rata sekitar 2.366,31 karakter per artikel; distribusi topik divisualisasikan dan disimpan pada data/distribusi_topik.png. Selama proses deteksi anomali ditemukan beberapa kasus yang perlu ditangani, antara lain entri berisi teks sangat pendek (< 50 karakter) yang kemungkinan tidak representatif (mis. No Title, No Topic), label tidak sesuai atau tersisa tag seperti FEATURED, serta indikasi kesalahan pemisah label pada beberapa entri. Semua anomali tersebut diproses (penghapusan atau normalisasi) sesuai kebijakan pra-pemrosesan, kemudian dataset bersih diekspor ke data/berita_cleaned.csv untuk tahap berikutnya (vektorisasi TF-IDF dan pelabelan multi-label); keseluruhan pipeline pembersihan ini dieksekusi secara efisien dalam waktu sekitar 2,51 detik

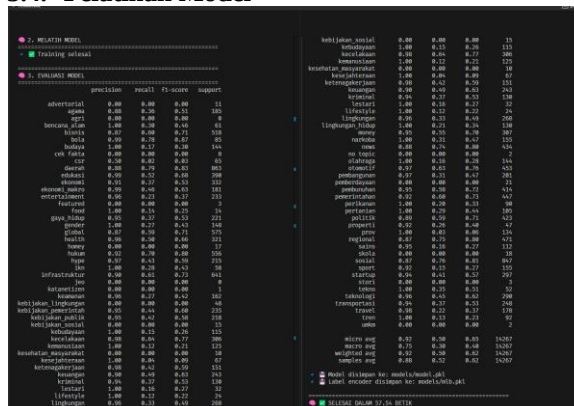
3.3. Pelabelan Multi-Label



Gambar 3. Memperkaya label dengan topic keyword

Pada Gambar 3. Diatas pada tahap pelabelan, data berita yang telah dibersihkan dimuat kembali beserta pemetaan topiknya. Proses ini dilanjutkan dengan memperkaya label menggunakan pendekatan berbasis kata kunci untuk mengidentifikasi dan menambahkan topik tambahan pada artikel yang relevan. Hasilnya, terbentuk kolom baru `topic_multi` yang merepresentasikan daftar topik pada setiap artikel, sehingga memungkinkan skenario klasifikasi multi-label. Untuk mempersiapkan data bagi algoritma pembelajaran mesin, label multi-topik ini kemudian diubah menjadi representasi vektor biner menggunakan `MultiLabelBinarizer`, di mana setiap kolom biner menunjukkan keberadaan atau ketiadaan suatu topik pada artikel tertentu. Dataset akhir multi-label ini disimpan dalam berkas `data/berita_multilabel.csv` sebagai masukan untuk tahap pelatihan model.

3.4. Pelatihan Model

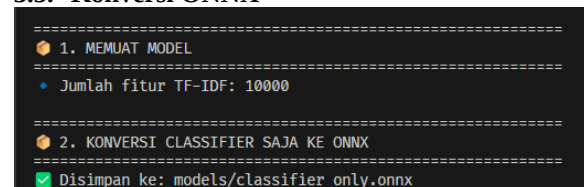


Gambar 4. Pelatihan model

Pada Gambar 4. Diatas pelatihan model dilakukan dengan pipeline yang menggabungkan representasi TF-IDF untuk fitur teks dan classifier `OneVsRestClassifier` dengan Logistic Regression untuk menangani skenario multi-label. Sebelum pelatihan, label dikonversi menjadi vektor biner menggunakan `MultiLabelBinarizer` dan data dibagi secara stratified menjadi 80% untuk pelatihan dan 20% untuk pengujian. Setelah proses fitting, model dievaluasi secara komprehensif dengan metrik per-label (precision, recall, f1-score dan support) serta metrik agregat seperti F1-micro dan F1-macro untuk menilai performa keseluruhan. Objek penting pada akhir pipeline—termasuk model terlatih, TF-IDF

vectorizer, dan label binarizer—disimpan untuk keperluan inferensi dan reproduksibilitas (mis. `models/model.pkl` dan `models/lb.pkl`). Selain itu, hasil per-label digunakan untuk analisis kesalahan dan identifikasi topik yang memerlukan penanganan ketidakseimbangan atau penyetelan threshold. Seluruh proses pelatihan dan evaluasi pada lingkungan pengujian berlangsung dalam waktu sekitar 57,54 detik.

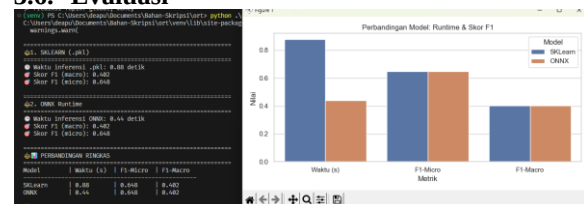
3.5. Konversi ONNX



Gambar 5. Konversi ke ONNX

Pada Gambar 5. Diatas Proses konversi dilakukan dengan `skl2onnx` untuk mengubah model classifier hasil pelatihan (Logistic Regression yang menggunakan representasi TF-IDF dengan jumlah fitur = 10.000) ke format ONNX, dan menyimpan berkas hasil konversi sebagai `models/classifier_only.onnx`. Pada implementasi ini hanya komponen classifier yang dikonversi—vektorisasi TF-IDF tetap disimpan sebagai objek Scikit-learn—sehingga alur inferensi di lingkungan produksi adalah: (1) transformasi teks menjadi vektor TF-IDF menggunakan vectorizer asli, lalu (2) passing vektor tersebut ke ONNX Runtime untuk prediksi. Konversi mencakup pendefinisian `initial_types` (tipe input) dan penentuan opset yang kompatibel, serta dilakukan validasi kesetaraan keluaran dengan model .pkl (memeriksa bahwa prediksi dan metrik seperti F1 tetap konsisten). Penggunaan ONNX Runtime memungkinkan pemanfaatan optimasi graph (graph simplification, node fusion) dan pengaturan session (mis. optimization level, execution provider) untuk mempercepat inferensi; sebagai langkah lanjutan tersedia opsi quantization jika diperlukan untuk memperkecil ukuran model dan mempercepat eksekusi pada perangkat terbatas. Penting untuk memastikan kompatibilitas versi `skl2onnx`, `onnx`, dan `onnxruntime`, serta menyimpan artefak pendukung (vectorizer dan label binarizer) agar pipeline inference dapat direproduksi dengan benar.

3.6. Evaluasi



Gambar 6. Evaluasi SKL vs ONNX Runtime

Pada Gambar 6. Diatas Setelah model Logistic Regression berbasis TF-IDF dikonversi dari format

.pkl ke .onnx menggunakan skl2onnx, dilakukan evaluasi komparatif antara eksekusi menggunakan Scikit-learn dan ONNX Runtime. Pengujian menggunakan dataset uji yang sama menunjukkan bahwa ONNX Runtime mampu memangkas waktu inferensi dari 0,88 detik menjadi 0,44 detik (peningkatan kecepatan sekitar 50%) tanpa penurunan performa prediksi. Nilai F1-Micro tetap sebesar 0,648 dan F1-Macro tetap sebesar 0,402 pada kedua model, menunjukkan bahwa konversi tidak mempengaruhi akurasi. Hasil ini divisualisasikan dalam bentuk diagram batang yang membandingkan waktu eksekusi, F1-Micro, dan F1-Macro untuk masing-masing model, dengan jelas memperlihatkan efisiensi ONNX Runtime dalam mempercepat proses inferensi sambil mempertahankan kualitas prediksi yang konsisten.

4. HASIL DAN PEMBAHASAN

Penelitian ini menunjukkan bahwa ONNX Runtime mampu mengoptimalkan kecepatan inferensi model TF-IDF + Logistic Regression untuk klasifikasi multi-label berita Indonesia tanpa mengurangi akurasi. Ke depan, kombinasi model ringan dengan teknik deep learning dapat dieksplorasi untuk meningkatkan pemahaman konteks berita.

4.1. Pra-pemrosesan dan Statistik

Dataset awal berisi 10.058 artikel berita dari Kompas.com yang masing-masing memiliki atribut judul, isi, dan topik. Setelah proses pembersihan nilai kosong, penghapusan duplikasi, dan penggabungan judul serta isi menjadi kolom text, diperoleh rata-rata panjang teks sebesar 2.366,31 karakter per artikel. Analisis menunjukkan terdapat 31 topik unik, dengan beberapa topik dominan seperti news, bola, dan bisnis. Distribusi topik divisualisasikan untuk memantau keseimbangan label, dan ditemukan adanya ketidakseimbangan di mana topik populer memiliki jumlah artikel yang jauh lebih banyak dibanding topik minor.

4.2. Pelabelan Multi-label

Menggunakan pendekatan berbasis kata kunci, label berita diperkaya sehingga memungkinkan satu artikel memiliki lebih dari satu topik (multi-label). Kolom topic_multi kemudian diubah menjadi vektor biner menggunakan MultiLabelBinarizer, menghasilkan representasi label yang siap diproses oleh algoritma pembelajaran mesin. Dataset multi-label ini disimpan untuk keperluan pelatihan dan evaluasi.

4.3. Pelatihan Model

Menggunakan pendekatan berbasis kata kunci, label berita diperkaya sehingga memungkinkan satu artikel memiliki lebih dari satu topik (multi-label). Kolom topic_multi kemudian diubah menjadi vektor biner menggunakan MultiLabelBinarizer, menghasilkan representasi label yang siap diproses oleh algoritma pembelajaran mesin. Dataset multi-

label ini disimpan untuk keperluan pelatihan dan evaluasi.

4.4. Evaluasi Model Scikit-learn

Hasil evaluasi pada dataset uji menunjukkan variasi performa antar topik. Beberapa topik seperti iklan layanan masyarakat dan gaya hidup memperoleh nilai F1-score di atas 0,80, sedangkan topik dengan jumlah data sedikit menunjukkan skor yang lebih rendah. Secara keseluruhan, model menghasilkan F1-Micro = 0,648 dan F1-Macro = 0,402, menunjukkan bahwa meskipun distribusi label tidak seimbang, model masih dapat memberikan prediksi yang cukup baik pada sebagian besar topik.

4.5. Konversi Model ke ONNX

Model Logistic Regression hasil pelatihan dikonversi ke format .onnx menggunakan pustaka skl2onnx. Pada tahap ini, hanya komponen classifier yang dikonversi, sedangkan TF-IDF vectorizer tetap dipertahankan dalam format Scikit-learn. Hal ini memungkinkan alur inferensi yang efisien, di mana transformasi teks ke vektor dilakukan terlebih dahulu, kemudian hasil vektor dikirim ke model ONNX untuk prediksi.

4.6. Perbandingan Scikit-learn vs ONNX Runtime

Pengujian komparatif menunjukkan bahwa ONNX Runtime secara konsisten menghasilkan prediksi yang identik dengan model Scikit-learn, terbukti dari nilai F1-Micro dan F1-Macro yang sama (0,648 dan 0,402). Namun, dari sisi waktu inferensi, ONNX Runtime menunjukkan peningkatan signifikan, memproses seluruh data uji dalam 0,44 detik dibandingkan Scikit-learn yang membutuhkan 0,88 detik — peningkatan efisiensi sekitar 50%.

4.7. Implikasi dan Potensi Pengembangan

Hasil ini membuktikan bahwa ONNX Runtime dapat menjadi solusi optimasi inferensi untuk model klasifikasi teks ringan berbasis TF-IDF + Logistic Regression, terutama pada aplikasi yang membutuhkan respons cepat seperti portal berita digital atau sistem monitoring informasi real-time. Ke depan, pendekatan serupa dapat diperluas dengan mengintegrasikan model ringan berbasis deep learning seperti DistilBERT atau fastText yang juga dikonversi ke ONNX, guna meningkatkan pemahaman konteks dan akurasi prediksi tanpa mengorbankan kecepatan.

5. KESIMPULAN DAN SARAN

Penelitian ini berhasil mengimplementasikan sistem klasifikasi multi-label berita berbahasa Indonesia menggunakan kombinasi TF-IDF dan Logistic Regression yang kemudian dioptimalkan melalui konversi ke format ONNX untuk dieksekusi dengan ONNX Runtime. Berdasarkan pengujian pada dataset berita Kompas.com yang berjumlah 10.058 artikel dan terdiri dari 31 topik unik, model menghasilkan kinerja yang konsisten baik sebelum

maupun sesudah konversi, dengan skor F1-Micro sebesar 0,648 dan F1-Macro sebesar 0,402 pada kedua versi model.

Konversi model ke ONNX terbukti mampu mengurangi waktu inferensi hingga 50% (dari 0,88 detik menjadi 0,44 detik) tanpa menurunkan akurasi. Hal ini menunjukkan bahwa ONNX Runtime dapat menjadi solusi efektif untuk meningkatkan efisiensi eksekusi model pembelajaran mesin, khususnya pada skenario yang membutuhkan pemrosesan cepat seperti portal berita digital dan sistem pemantauan informasi real-time.

Selain itu, penelitian ini memberikan gambaran bahwa pendekatan multi-label dengan model ringan dapat diintegrasikan dengan pipeline inferensi yang efisien, sehingga layak dipertimbangkan untuk implementasi di lingkungan produksi. Ke depannya, penelitian dapat dikembangkan dengan mengeksplorasi model lightweight deep learning seperti DistilBERT atau fastText yang juga dikonversi ke ONNX, untuk meningkatkan pemahaman konteks dan akurasi tanpa mengorbankan kecepatan. Penelitian selanjutnya juga dapat menitikberatkan pada penanganan ketidakseimbangan data antar-topik, misalnya melalui teknik data augmentation atau class-weight adjustment untuk memperbaiki kinerja pada topik minor.

DAFTAR PUSTAKA

- [1] F. A. Ramadhan, S. H. Sitorus, and T. Rismawan, "Penerapan Metode Multinomial Naïve Bayes untuk Klasifikasi Judul Berita Clickbait dengan Term Frequency - Inverse Document Frequency," *Jurnal Sistem dan Teknologi Informasi (JustIN)*, vol. 11, no. 1, p. 70, Jan. 2023, doi: 10.26418/justin.v11i1.57452.
- [2] B. Hendra Mahendra and U. Novia Wisesty, "Kategorisasi Berita Multi-Label Berbahasa Indonesia Menggunakan Algoritma Random Forest." [Online]. Available: <https://www.jawapos.com>.
- [3] I. Putu Gede Hendra Suputra, I. Gede Sukadarmika, N. Putra Sastra, I. Gusti Bgs Darmika Putra, and I. Wayan Trisna Wahyudi, "KLASIFIKASI JUDUL BERITA BAHASA INDONESIA MENGGUNAKAN SUPPORT VECTOR MACHINE DAN SELEKSI FITUR MUTUAL INFORMATION," *Jurnal Pendidikan Teknologi dan Kejuruan*, vol. 22, no. 1, 2025, [Online]. Available: www.kompas.com.
- [4] R. Aji, S. #1, H. Fachrizal, E. K. #2, C. S. Kusuma, and A. #3, "JEPIN (Jurnal Edukasi dan Penelitian Informatika) Klasifikasi Hoax Berita Politik Menggunakan Algoritma Long Short-Term Memory (LSTM) dengan Penambahan Fitur Embedding Global Vector (GloVe)," [Online]. Available: <https://www.kominfo.go.id>
- [5] S. Daud, M. Ullah, A. Rehman, T. Saba, R. Damaševičius, and A. Sattar, "Topic Classification of Online News Articles Using Optimized Machine Learning Models," *Computers*, vol. 12, no. 1, Jan. 2023, doi: 10.3390/computers12010016.
- [6] R. Roshan, I. A. Bhacho, and S. Zai, "Comparative Analysis of TF-IDF and Hashing Vectorizer for Fake News Detection in Sindhi: A Machine Learning and Deep Learning Approach †," *Engineering Proceedings*, vol. 46, no. 1, 2023, doi: 10.3390/engproc2023046005.
- [7] C. M. Lai, M. H. Chen, E. Kristiani, V. K. Verma, and C. T. Yang, "Fake News Classification Based on Content Level Features," *Applied Sciences (Switzerland)*, vol. 12, no. 3, Feb. 2022, doi: 10.3390/app12031116.
- [8] Nadira Alifia Ionendri, Feri Candra, and Afdi Rizal, "News Classification using Natural Language Processing with TF-IDF and Multinomial Naïve Bayes," *Journal of Applied Computer Science and Technology*, vol. 6, no. 1, pp. 37–45, Jun. 2025, doi: 10.52158/jacost.v6i1.1099.
- [9] E. D. Ajik, G. N. Obunadike, and F. O. Echobu, "Fake News Detection Using Optimized CNN and LSTM Techniques," *Journal of Information Systems and Informatics*, vol. 5, no. 3, pp. 1044–1057, Aug. 2023, doi: 10.51519/journalisi.v5i3.548.
- [10] B. P. Nayoga, R. Adipradana, R. Suryadi, and D. Suhartono, "Hoax Analyzer for Indonesian News Using Deep Learning Models," in *Procedia Computer Science*, Elsevier B.V., 2021, pp. 704–712. doi: 10.1016/j.procs.2021.01.059.
- [11] D. Kurniasari, L. N. Azizah, P. H. Khotimah, and Warsono, "Performance of CNN and LSTM Model on COVID-19 News Headline Data Classification," *Baghdad Science Journal*, vol. 22, no. 7, pp. 2458–2468, Jul. 2025, doi: 10.21123/2411-7986.5009.
- [12] W. Anggraeni, M. F. A. Kusuma, E. Riksakomara, R. P. Wibowo, Pujiadi, and S. Sumpeno, "Combination of BERT and Hybrid CNN-LSTM Models for Indonesia Dengue Tweets Classification," *International Journal of Intelligent Engineering and Systems*, vol. 17, no. 1, pp. 813–826, 2024, doi: 10.22266/ijies2024.0229.68.
- [13] R. Adipradana, B. P. Nayoga, R. Suryadi, and D. Suhartono, "Hoax analyzer for indonesian news using rnns with fasttext and glove embeddings," *Bulletin of Electrical Engineering and Informatics*, vol. 10, no. 4, pp. 2130–2136, Aug. 2021, doi: 10.11591/eei.v10i4.2956.