



Implementasi FastText untuk Klasifikasi Bahasa Multilingual: Pendekatan Efisien dalam Sistem Pemrosesan Bahasa Alami (NLP)

Devin Slamet¹, Heri Purnomo², Muhammad Listanto³, Reni Anggariani⁴, Perani Rosyani⁵

¹⁻⁵ Fakultas Ilmu Komputer, Universitas Pamulang, Jl. Raya Puspipetek No. 46, Serpong, Kota Tangerang Selatan, Banten

Email: ¹devinslamet52@gmail.com, ²heripurnomo2041@gmail.com, ³mlistanto29@gmail.com,
⁴renianggariani8@gmail.com, ⁵dosen00837@unpam.ac.id

Abstrak—Klasifikasi bahasa merupakan salah satu tugas dasar dalam pemrosesan bahasa alami (Natural Language Processing/NLP) yang berfungsi untuk mendeteksi bahasa dari suatu teks. Tugas ini menjadi semakin penting seiring meningkatnya penggunaan aplikasi multilingual seperti chatbot, platform digital global, dan sistem moderasi konten. Penelitian ini mengimplementasikan FastText sebagai model klasifikasi untuk mengidentifikasi bahasa menggunakan dataset multilingual dari Kaggle. FastText dipilih karena kemampuannya dalam memanfaatkan karakter n-gram dan subword representation, sehingga tetap mampu memberikan prediksi yang presisi meskipun teks yang digunakan sangat pendek atau memiliki variasi penulisan. Proses penelitian meliputi tahapan preprocessing teks, konversi label ke format FastText, pelatihan model, pengujian, serta evaluasi menggunakan metrik akurasi, precision, recall, F1-score, dan confusion matrix. Hasil evaluasi menunjukkan bahwa model FastText mampu mencapai akurasi lebih dari 90% dengan waktu pelatihan yang ringkas dan efisien. Temuan ini menegaskan bahwa FastText merupakan solusi praktis dan ringan untuk kebutuhan deteksi bahasa otomatis dalam skala besar. Selain itu, penelitian ini memberikan gambaran mengenai bagaimana metode klasifikasi berbasis embedding seperti FastText dapat diterapkan pada dataset multilingual dengan struktur yang beragam. Implementasi FastText dalam proyek ini menunjukkan bahwa pendekatan yang sederhana namun efektif tetap mampu memberikan performa kompetitif untuk tugas-tugas dasar NLP, sehingga dapat dijadikan referensi bagi pengembangan sistem NLP lain yang memerlukan tahap deteksi bahasa sebagai bagian awal pemrosesan. Di samping hasil performa yang baik, penelitian ini juga menunjukkan bahwa proses implementasi FastText relatif mudah dilakukan, sehingga cocok digunakan baik oleh pemula maupun pengembang yang membutuhkan solusi cepat untuk klasifikasi bahasa. Pendekatan ini dapat menjadi dasar bagi pengembangan model NLP yang lebih kompleks dan mendalam di masa mendatang.

Kata kunci: *FastText, NLP, klasifikasi bahasa, machine learning, multilingual.*

Abstract—Language classification is one of the fundamental tasks in Natural Language Processing (NLP), used to automatically detect the language of a given text. This task has become increasingly important with the rapid growth of multilingual applications such as chatbots, global digital platforms, and content moderation systems. This study implements FastText as the classification model for identifying languages using a multilingual dataset obtained from Kaggle. FastText is chosen due to its ability to utilize character n-grams and subword representations, enabling it to produce accurate predictions even when dealing with short texts or inconsistent writing patterns. The research process includes text preprocessing, converting labels into FastText format, model training, testing, and evaluation using metrics such as accuracy, precision, recall, F1-score, and the confusion matrix. The evaluation results show that the FastText model achieves an accuracy of over 90%, with efficient and relatively fast training time. These findings confirm that FastText is a practical and lightweight solution for large-scale automatic language detection tasks. Furthermore, this study demonstrates how embedding-based classification methods like FastText can be applied to multilingual datasets with varying structures. The implementation of FastText in this project shows that a simple yet effective approach can still provide competitive performance for basic NLP tasks, making it a useful reference for developing other NLP systems that require language detection as an initial processing stage. In addition to its strong performance, the study also highlights that implementing FastText is relatively straightforward, making it suitable for both beginners and developers who require a fast solution for language classification. This approach can also serve as a foundation for building more advanced and complex NLP models in the future.

Keywords: *FastText, NLP, language classification, machine learning, multilingual.*

1. PENDAHULUAN

Perkembangan teknologi digital menyebabkan meningkatnya penggunaan data teks dari berbagai negara dan bahasa. Sistem seperti chatbot, layanan pelanggan otomatis, media sosial, serta platform global menghadapi tantangan dalam memproses teks multilingual. Agar sistem dapat



memberikan respons yang tepat, dibutuhkan proses *language identification* sebagai tahap awal dalam NLP.

Namun, proses ini tidak mudah karena teks sering kali pendek, tidak baku, mengandung typo, singkatan, bahasa gaul, atau campuran bahasa. FastText, model berbasis subword dan karakter n-gram yang dikembangkan oleh Meta AI, menawarkan solusi efektif untuk permasalahan tersebut.

FastText mampu membaca representasi kata berdasarkan potongan subword, sehingga dapat mengenali pola linguistik meskipun data memiliki variasi penulisan. Selain itu, FastText dikenal cepat dan efisien dalam pelatihan maupun prediksi, menjadikannya populer untuk aplikasi real-time.

Penelitian ini bertujuan menerapkan FastText untuk deteksi bahasa, melakukan preprocessing data, melatih model menggunakan parameter tertentu, serta mengevaluasi performanya menggunakan metrik NLP standar.

2. METODE

2.1 Rancangan Penelitian

Penelitian dilakukan melalui tahapan utama:

1. Pengumpulan dataset multilingual dari Kaggle,
2. Preprocessing teks,
3. Mengonversi dataset ke format fasttext,
4. Pelatihan model,
5. Pengujian, dan
6. Evaluasi performa.

2.2 Dataset

Dataset berisi ribuan teks pendek dari beberapa bahasa seperti Inggris, Indonesia, Prancis, Arab, Rusia, Spanyol, dll. Tiap baris terdiri dari teks dan label bahasa.

2.3 Preprocessing

Tahapan preprocessing meliputi:

1. Mengubah teks menjadi lowercase,
2. Menghapus tanda baca,
3. Menghapus angka dan simbol tidak relevan,
4. Memastikan data bersih dan siap diformat.

2.4 Format FastText

Data diformat sebagai:

`__label__en hello how are you`

`__label__id apa kabar hari ini`

Setiap baris mewakili satu sampel pelatihan.

2.5 Pelatihan Model

Parameter pelatihan:

- learning rate (lr): 1.0
- epoch: 25
- wordNgrams: 2
- dimensi embedding: 100
- loss function: softmax

Model ini menghasilkan file *model.bin* yang digunakan untuk prediksi.

2.6 Evaluasi

Evaluasi dilakukan dengan metrik:



- akurasi
- precision
- recall
- F1-score
- confusion matrix

3. HASIL DAN PEMBAHASAN

3.1 Hasil Pelatihan

Model FastText berhasil dilatih dengan cepat dan efisien. Proses pelatihan tidak memerlukan sumber daya komputasi besar, sesuai karakteristik FastText sebagai model ringan.

3.2 Analisis Model

FastText menunjukkan performa sangat baik terutama pada bahasa yang memiliki alfabet unik seperti Arab dan Rusia. Pada bahasa yang serumpun seperti Spanyol, Portugis, atau Italia, model tetap stabil karena penggunaan n-gram membantu membedakan pola karakter. FastText juga sangat efektif pada teks pendek atau informal.

3.3 Hasil Evaluasi

Ringkasan hasil evaluasi:

Metrik	Nilai
Akurasi	93.8%
Precision rata-rata	92%
Recall rata-rata	91%
F1-score	91.5%

Language-level sample (misalnya disederhanakan):

Bahasa seperti Inggris, Indonesia, Prancis, dan Spanyol mendapat F1-score **1.00** karena struktur dataset yang rapi.

3.4 Pembahasan

Hasil membuktikan bahwa FastText efektif untuk deteksi bahasa pada teks multilingual. Representasi subword membuat model tetap mampu mengenali pola bahasa meski teks tidak memiliki konteks panjang. FastText juga dapat digunakan pada aplikasi real-time berkat kecepatannya.

4. KESIMPULAN DAN SARAN

4.1 Kesimpulan

Penelitian menunjukkan bahwa FastText merupakan metode yang kuat untuk klasifikasi bahasa pada dataset multilingual. Dengan memanfaatkan subword embedding dan n-gram, model mampu memberikan hasil akurat dan stabil, bahkan pada teks pendek. Hasil evaluasi menunjukkan performa tinggi dengan akurasi 93.8% dan metrik lain yang konsisten. FastText cocok digunakan sebagai komponen awal dalam sistem NLP modern.

4.2 Saran

Untuk penelitian selanjutnya:

1. Perlu menambahkan lebih banyak bahasa untuk meningkatkan generalisasi model.
2. Dapat dilakukan perbandingan FastText dengan model berbasis Transformer seperti mBERT atau XLM-R.



JRIIN : Jurnal Riset Informatika dan Inovasi
Volume 3, No. 9, Februari Tahun 2026
ISSN 3025-0919 (media online)
Hal 2494-2497

3. Perlu diuji pada data real-world seperti komentar media sosial atau pesan pengguna.

REFERENCES

- Ahmed, M., & Rahman, T. (2020). *A Review of Subword Embedding Techniques for Multilingual NLP*. International Journal of Computer Science Research.
- Alshaiba, R., & Mahmoud, B. (2024). *Comparative Performance of FastText and Transformer Models in Language Detection Tasks*. International Journal of Artificial Intelligence Systems.
- Chandra, S., & Iyer, M. (2020). *A Comprehensive Survey of Language Detection Algorithms*. International Journal of AI Research.
- Kaggle. (2023). *Multilingual Text Classification Dataset*. Kaggle.com.
- Khan, S., & Gupta, R. (2022). *Lightweight Approaches for Multilingual Language Identification Using Subword Embeddings*. Journal of Natural Language Engineering.
- Kim, J., & Park, H. (2022). *Improving Language Identification Using N-gram Enhanced Embedding Models*. Journal of Digital Language Processing.
- Lestari, D., & Haris, A. (2022). *Evaluasi Model FastText untuk Identifikasi Bahasa pada Media Sosial*. Jurnal Sains Komputer Indonesia.
- Meta AI Research. (2021). *FastText: Efficient Learning of Word Representations*. Meta Platforms, Inc.
- Meta AI. (2022). *FastText Documentation*. fasttext.cc.
- Suryana, F., & Putra, R. (2024). *Penerapan FastText dan Analisis Confusion Matrix untuk Deteksi Bahasa Otomatis*. Jurnal Teknologi dan Sistem Informasi.