

Pemanfaatan Pemrosesan Bahasa Alami untuk Identifikasi Kecemasan Berbahasa Inggris

Yasri Ridho Pahlevi¹, Kemas Muslim Lhaksana², Moch Arif Bijaksana³

^{1,2,3}Fakultas Informatika, Universitas Telkom, Bandung

¹yasridho@student.telkomuniversity.ac.id, ²kemasmuslim@telkomuniversity.ac.id,

³arifbijaksana@telkomuniversity.ac.id

Abstrak

Deteksi Kecemasan Berbahasa Asing (FLA), sebuah hambatan psikologis signifikan, secara tradisional bergantung pada kuesioner yang subjektif. Meskipun metode deteksi otomatis yang lebih objektif seperti analisis fisiologis telah berkembang, pemanfaatan Pemrosesan Bahasa Alami (NLP) untuk menganalisis esai mahasiswa menawarkan alternatif yang lebih praktis dan dapat diskalakan. Penelitian ini bertujuan mengembangkan dan mengevaluasi sebuah sistem identifikasi FLA berbasis teks yang efisien, dengan fokus pada penggunaan model *machine learning* klasik yang hemat sumber daya komputasi. Dengan dataset yang secara inheren tidak seimbang, enam model klasifikasi dan empat metode representasi teks dievaluasi menggunakan *Stratified 5-Fold Cross-Validation*, dengan performa diukur melalui *macro-average F1-score* untuk menilai stabilitas. Hasil menunjukkan bahwa kombinasi SBERT dan Random Forest, meski mencapai F1-score tertinggi (0.849), gagal mengenali kelas minoritas (*macro-average F1-score* 0.42). Sebaliknya, model TF-IDF dengan BernoulliNB terbukti lebih sederhana namun stabil, dengan *macro-average F1-score* 0.65. Temuan ini menegaskan bahwa untuk pemanfaatan praktis dalam dunia pendidikan, stabilitas model dalam mengenali semua kelompok siswa secara adil jauh lebih krusial daripada sekadar mencapai skor akurasi tertinggi yang bisa menyesatkan.

Kata Kunci: kecemasan berbahasa, klasifikasi teks, pembelajaran mesin, ketidakseimbangan kelas, rekayasa fitur, analisis sentimen

Abstract

Detecting Foreign Language Anxiety (FLA), a significant psychological barrier, has traditionally relied on subjective questionnaires. Although more objective automated detection methods like physiological analysis have emerged, leveraging Natural Language Processing (NLP) to analyze student essays offers a more practical and scalable alternative. This research aims to develop and evaluate an efficient text-based FLA identification system, focusing on the use of computationally efficient classical machine learning models. With an inherently imbalanced dataset, six classification models and four text representation methods were evaluated using *Stratified 5-Fold Cross-Validation*, with performance measured by the *macro-average F1-score* to assess stability. Results show that the combination of SBERT and Random Forest, despite achieving the highest F1-score (0.849), failed to recognize the minority class (*macro-average F1-score* of 0.42). Conversely, the TF-IDF model with BernoulliNB proved to be simpler yet more stable, with a *macro-average F1-score* of 0.65. This finding confirms that for practical application in education, a model's stability in fairly recognizing all student groups is far more crucial than merely achieving the highest, potentially misleading, accuracy score.

Keywords: language anxiety, text classification, machine learning, class imbalance, feature engineering, sentiment analysis