

ABSTRAK

Evaluasi pelatihan ICT di UIN Sunan Gunung Djati Bandung dilakukan melalui pengisian survei yang menghasilkan data rating dan data teks. Meskipun tersedia 81.056 *record* data survei dari 8.276 responden, data teks belum dimanfaatkan secara optimal untuk mengidentifikasi tingkat kepuasan mahasiswa. Oleh karena itu, penelitian ini bertujuan untuk mengintegrasikan data rating dan teks serta menerapkan model IndoBERT untuk mengklasifikasikan tingkat kepuasan mahasiswa terhadap pelatihan ICT. Data rating digunakan sebagai dasar pembentukan label, sedangkan data teks digunakan sebagai masukan klasifikasi pada model. Penelitian dilakukan menggunakan metodologi CRISP-DM dengan model *indobenchmark/indobert-base-pl* yang melalui proses *fine-tuning* untuk mengklasifikasikan tiga kategori kepuasan, yaitu Tidak Puas, Cukup Puas, dan Puas. Ketidakseimbangan data ditangani menggunakan *Random Oversampling*, sedangkan *fine-tuning* dilakukan pada tiga skenario pembagian data, yaitu 60%:20%:20%, 70%:15%:15%, dan 80%:10%:10%. Hasil pengujian menunjukkan bahwa Skenario 2 dengan pembagian data 70%:15%:15% menghasilkan performa terbaik dengan *accuracy* 62,95%, *precision* 56,86%, *recall* 62,95%, dan *weighted F1-score* 57,97%. Analisis manual terhadap 150 sampel menunjukkan 92 data mengalami ketidaksesuaian antara label berdasarkan rating dan informasi kepuasan pada teks. Temuan tersebut mengindikasikan adanya ketidaksesuaian antara label dan karakteristik teks yang dapat memengaruhi kemampuan model dalam mengenali setiap kelas, terutama kelas Tidak Puas. Dengan demikian, IndoBERT telah mampu melakukan klasifikasi tingkat kepuasan mahasiswa, tetapi performanya masih belum optimal akibat ketidakseimbangan data dan ketidaksesuaian antara label rating dengan informasi pada teks.

Kata Kunci: IndoBERT, Klasifikasi Teks, *Natural Language Processing*, Pelatihan ICT, Tingkat Kepuasan Mahasiswa.

ABSTRACT

The evaluation of ICT training at UIN Sunan Gunung Djati Bandung is conducted through a survey that generates rating and textual data. Although the survey contains 81,056 records from 8,276 respondents, the textual data have not been optimally utilized to identify students' satisfaction levels. Therefore, this study aims to integrate rating and textual data and apply the IndoBERT model to classify students' satisfaction levels toward ICT training. The rating data are used as the basis for label formation, while the textual data are used as input for classification. The study employs the CRISP-DM methodology using the indobenchmark/indobert-base-pl model, which undergoes fine-tuning to classify three satisfaction categories: Dissatisfied, Moderately Satisfied, and Satisfied. Class imbalance is addressed using Random Oversampling, while fine-tuning is conducted using three data-splitting scenarios: 60%:20%:20%, 70%:15%:15%, and 80%:10%:10%. The results show that Scenario 2, with a 70%:15%:15% data split, achieves the best performance, with an accuracy of 62.95%, precision of 56.86%, recall of 62.95%, and weighted F1-score of 57.97%. Manual analysis of 150 samples identified 92 instances with a mismatch between the rating-based labels and the satisfaction information expressed in the text. This finding indicates that inconsistencies between labels and textual characteristics may affect the model's ability to recognize each class, particularly the Dissatisfied class. Thus, IndoBERT is capable of classifying students' satisfaction levels, although its performance remains suboptimal due to class imbalance and inconsistencies between rating-based labels and the information conveyed in the text.

Keywords: IndoBERT, ICT Training, Natural Language Processing, Student Satisfaction Level, Text Classification.