

ABSTRAK

Pemanfaatan *deep learning* saat ini sudah banyak diadaptasikan dalam berbagai proses bisnis karena dapat memberikan banyak manfaat dalam berbagai permasalahan. Terjadi permasalahan pada sistem repository widyatama dimana lamanya proses pembaharuan data karya tulis ilmiah dikarenakan semakin banyaknya data karya tulis ilmiah baru. Dengan semakin bertambahnya jumlah karya tulis ilmiah yang harus dimasukan setiap tahunnya menyebabkan “*bottleneck*” dalam proses pembaharuan data tersebut, permasalahan tersebut terjadi karena bertambahnya beban proses yang harus dikerjakan tidak diiringi dengan peningkatan kapasitas dalam pembaharuan data. Hal tersebut menyebabkan ketersediaan data baru pada sistem menjadi tidak *up to date*.

Pengembangan model *deep learning* dengan algoritma LSTM (*Long Short-Term Memory*) untuk melakukan klasifikasi teks abstrak karya tulis ilmiah berdasarkan program studi akan dipilih sebagai solusi yang diusulkan untuk mengatasi permasalahan tersebut. Model yang sudah berhasil dikembangkan akan diintegrasikan kedalam API sehingga pemanfaatan model akan lebih mudah untuk diintegrasikan dengan sistem repository.

Hasil pelatihan model LSTM untuk klasifikasi teks abstrak karya tulis ilmiah memperoleh akurasi tertinggi sebesar 93,31% dengan nilai rata-rata presisi, recall, dan F1-score adalah 95,66%. Dengan adanya model tersebut yang kemudian diintegrasikan dengan API yang berhasil dikembangkan untuk memenuhi fungsionalitas model, diharapkan dapat membantu meningkatkan proses pembaharuan data karya tulis ilmiah dalam sistem repository widyatama. Proses klasifikasi karya tulis ilmiah yang selama ini dilakukan secara manual dapat dilakukan secara otomatis dengan adanya API dan model LSTM tersebut.

Kata Kunci: *Deep Learning, Long Short-Term Memory, API, klasifikasi teks multi label, machine learning*

ABSTRACT

The utilization of deep learning is now widely adapted in various business processes because it offers numerous benefits for different problems. There is an issue with the Widyatama repository system where the updating process of scientific papers is delayed due to the increasing number of new scientific papers. The growing number of scientific papers that need to be added each year causes a 'bottleneck' in the data updating process. This problem arises because the increasing workload is not accompanied by an increase in capacity for updating data. As a result, the availability of new data in the system is not up to date.

The development of a deep learning model using the LSTM (Long Short-Term Memory) algorithm for classifying the abstract texts of scientific papers based on study programs will be chosen as the proposed solution to address this issue. The successfully developed model will be incorporated into an API, making it easier to integrate the model with the repository system.

The training results of the LSTM model for abstract text classification of scientific papers achieved the highest accuracy of 93.31%, with an average precision, recall, and F1-score of 95.66%. By integrating this model with an API where the LSTM model operates, it is expected to help improve the data updating process of scientific papers in the Widyatama repository system. The classification process of scientific papers, which has been done manually, can now be automated with the implementation of this LSTM model.

Keywords: *Deep Learning, Long Short-Term Memory, API, multi-label text classification, machine learning*