

Model Analisis Kinerja Sumber Daya Manusia (SDM) di Perguruan Tinggi Berbasis *Multi-modal Data* Menggunakan *Machine Learning*

¹Reni Utami, ²Ari Hidayatullah, ³Anita Ratnasari

^{1,2,3}Fakultas Teknik dan Informatika, Universitas Dian Nusantara, Jakarta, Indonesia

¹reni.utami@dosen.undira.ac.id, ²ari.hidayatullah@dosen.undira.ac.id, ³anita.ratnasari@undira.ac.id

Article Info

Article history:

Received, 2025-11-01

Revised, 2025-11-09

Accepted, 2025-11-11

Kata Kunci:

Analisis kinerja,
Random Forest,
CNN-BiLSTM,
Data statis,
Data teks,

Keywords:

Performance analysis,
Random Forest,
CNN-BiLSTM,
Static data,
Text data,

ABSTRAK

Penelitian ini bertujuan untuk mengembangkan model berbasis *machine learning* untuk analisis kinerja sumber daya manusia (SDM) di perguruan tinggi menggunakan pendekatan *multi-modal* berdasarkan analisis data statis dan data teks. Untuk menganalisis data statis, digunakan algoritma *random forest* (RF) yang digunakan untuk analisis kinerja SDM berdasarkan atribut seperti lama kerja, pelatihan yang diikuti, dan evaluasi kinerja. Dataset untuk eksperimen ini terdiri dari 250 data, yang dibagi menjadi 70% untuk pelatihan, 10% untuk validasi, dan 20% untuk pengujian. Hasil eksperimen dengan model RF menunjukkan akurasi tinggi pada pelatihan (90.4%), namun ada penurunan kinerja pada eksperimen validasi dan pengujian dengan akurasi masing-masing 85.7% dan 82.5%. Dari hasil eksperimen pada data teks yang berupa *feedback* dengan sentimen negatif, positif, dan netral menggunakan model CNN-BiLSTM mendapatkan akurasi 92.6% pada pelatihan, meskipun ada penurunan pada validasi (87.4%) dan pengujian (84.4%). Dataset teks terdiri dari 1.000 data, yang dibagi menjadi 70% untuk pelatihan, 10% untuk validasi, dan 20% untuk pengujian. Penelitian merekomendasikan penerapan pendekatan multi-model dalam menilai kinerja SDM menggunakan algoritma *random forest* (RF) dan model CNN-BiLSTM untuk data yang lebih kompleks pada penelitian selanjutnya.

ABSTRACT

This research aimed to develop a machine learning-based model for human resource (HR) performance analysis in higher education institutions using a multi-modal approach, combining static data and text data analysis. For static data analysis, the random forest (RF) algorithm was employed to assess HR performance based on attributes such as years of service, training attended, and performance evaluations. The dataset for this experiment consisted of 250 data points, which were divided into 70% for training, 10% for validation, and 20% for testing. The experimental results with the RF model showed high accuracy in training (90.4%), although there was a performance drop during validation and testing, with accuracies of 85.7% and 82.5%, respectively. For the text data, which contained feedback with negative, positive, and neutral sentiments, the CNN-BiLSTM model achieved an accuracy of 92.6% in training, despite a decrease in validation (87.4%) and testing (84.4%) accuracies. The text dataset comprised 1,000 data points, divided into 70% for training, 10% for validation, and 20% for testing. The study recommends the application of a multi-model approach to assess HR performance using the RF algorithm for static data and the CNN-BiLSTM model for more complex data in future research.

This is an open access article under the CC BY-SA license.



Penulis Korespondensi:

Anita Ratnasari,
Fakultas Teknik dan Informatika
Universitas Dian Nusantara, Jakarta, Indonesia
Email: anita.ratnasari@undira.ac.id

1. PENDAHULUAN

Perkembangan institusi pendidikan tinggi tidak hanya ditentukan oleh capaian akademik, tetapi juga oleh kualitas kesejahteraan dan motivasi dari sumber daya manusia yang bekerja dalam institusi [1], [2]. Kepuasan kerja memainkan peran krusial dalam membentuk semangat kerja, loyalitas pegawai, dan iklim organisasi yang sehat [3], [4]. Sistem manajemen sumber daya manusia (SDM) sangat penting untuk mengevaluasi dan menganalisis terhadap produktivitas institusi dalam menjalankan program strategis [5]. Hasil evaluasi dapat digunakan untuk memprediksi kepuasan kerja untuk merancang kebijakan sehingga penempatan sumber daya dapat sesuai dengan kebutuhan unit kerja dalam institusi pendidikan tinggi [6], [7].

Dalam analisis kinerja SDM dapat menganalisis data statis dan data teks. Beberapa penelitian menggunakan data statis untuk analisis kinerja karyawan, seperti penelitian Dong (2024) menerapkan algoritma Random Forest untuk mengklasifikasikan dan mengevaluasi efektivitas data sumber daya manusia (SDM) dengan fokus untuk mendukung pengambilan keputusan dan meningkatkan efisiensi organisasi. Model yang dikembangkan mengotomatiskan kategorisasi data SDM, termasuk catatan karyawan, evaluasi kinerja, dan kegiatan pelatihan, serta memvalidasi keakuratan model analisis [8]. Penelitian Alsaadi et al., (2022) mengidentifikasi elemen-elemen penting yang berkontribusi terhadap kinerja karyawan, yang merupakan masalah utama dalam sudut pandang sumber daya manusia (SDM). Penelitian ini menemukan bahwa fitur karyawan seperti lembur, jumlah proyek, dan level pekerjaan memiliki dampak signifikan terhadap attrition. Beberapa algoritma klasifikasi seperti *decision trees* (DT), *logistic regression* (LR), *random forests* (RF), dan *K-means clustering* digunakan untuk analisis kinerja karyawan, dengan analisis komparatif untuk menentukan model dengan akurasi tertinggi [9]. Penelitian Xu (2021) menggunakan algoritma *random forest* (RF) untuk membangun model evaluasi kualitas pengelolaan SDM di perusahaan, berdasarkan data historis. Sebelum melatih *classifier*, algoritma RF sorting digunakan untuk mengukur pentingnya fitur dan melakukan reduksi dimensi dengan memilih 75% fitur untuk mengatasi ketidakseimbangan sampel pelatihan. Penelitian ini menekankan pentingnya penggunaan SDM yang rasional agar perusahaan dapat memaksimalkan potensi karyawan dan menciptakan manfaat jangka panjang melalui penempatan yang tepat di posisi yang sesuai [10].

Beberapa penelitian menggunakan data teks untuk analisis kinerja karyawan seperti Lee & Song (2024) mengembangkan model konseptual pengalaman positif karyawan dengan menggunakan analisis sentimen dalam strategi sumber daya manusia (SDM) berbasis algoritma untuk meningkatkan pemahaman profesional HR mengenai pengalaman karyawan dan pengambilan keputusan berbasis data untuk menciptakan lingkungan kerja yang positif. Melalui analisis sentimen, peneliti mengidentifikasi fitur yang menggambarkan pengalaman karyawan, yang dikelompokkan dalam empat kluster yang mempengaruhi pengalaman karyawan berupa pekerjaan, hubungan, sistem organisasi, dan budaya organisasi [11]. Penelitian Zhang & Qi (2022) mengembangkan dua model prediksi stres karyawan menggunakan pendekatan *deep learning* (DL) untuk mengatasi masalah stres berlebihan yang memengaruhi produktivitas, keselamatan, dan kesehatan karyawan. Model ini dirancang untuk analisis kinerja karyawan berdasarkan data seperti gaji, waktu kerja dengan akurasi 71,2% untuk model klasifikasi untuk model CNN. Penelitian ini menunjukkan bahwa pendekatan *deep learning* menjadi rekomendasi model untuk perusahaan dengan akurasi prediksi yang lebih baik [12].

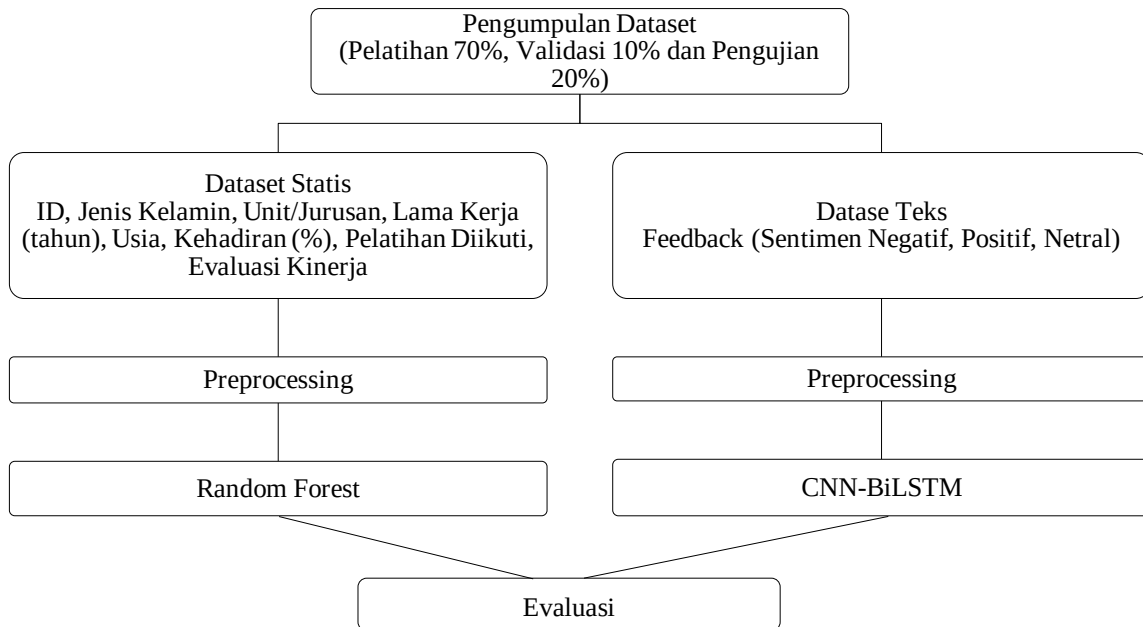
Dalam berbagai penelitian sebelumnya, analisis kinerja sumber daya manusia (SDM) telah dilakukan menggunakan data statis dan data teks secara terpisah. Penelitian seperti yang dilakukan oleh Dong (2024) dan Xu (2021) lebih fokus pada pemanfaatan data statis, seperti catatan karyawan, evaluasi kinerja, dan data historis, untuk mengembangkan model klasifikasi menggunakan algoritma seperti *random forest* (RF). Di sisi lain, penelitian yang dilakukan oleh Lee & Song (2024) dan Zhang & Qi (2022) mengandalkan data teks, seperti sentimen atau data terkait stres karyawan, dengan memanfaatkan model berbasis *deep learning* (DL), seperti *Convolutional Neural Networks* (CNN) dan *Bi-directional Long Short-Term Memory* (Bi-LSTM), untuk menganalisis dan memprediksi kinerja karyawan.

Namun, sebagian besar penelitian yang ada belum menggabungkan kedua jenis data tersebut (data statis dan data teks) dalam satu analisis komprehensif yang dapat memberikan gambaran mengenai kinerja karyawan. Gap penelitian ini adalah bahwa sebagian besar pendekatan sebelumnya hanya terbatas pada analisis data statis atau teks secara terpisah, sedangkan dalam konteks kinerja SDM, data dari berbagai sumber (baik statis maupun teks) dapat memberikan rekomendasi yang dapat menjadi bahan pertimbangan pengambilan kebijakan.

Penelitian ini bertujuan mengembangkan model analisis kinerja karyawan yang memanfaatkan data multimodal, yaitu gabungan data statis dan teks. Data statis akan dianalisis menggunakan algoritma Random Forest (RF), dengan mengelompokkan dan mengevaluasi data numerik atau kategori dari karyawan. Sementara itu, data teks berisi *feedback* SDM akan dianalisis menggunakan model CNN-Bi-LSTM. Model ini dapat mengidentifikasi pola temporal dalam data teks dan memberikan konteks yang lebih dalam mengenai pengalaman atau kondisi karyawan yang dapat memengaruhi kinerja.

2. METODE PENELITIAN

Eksperimen penelitian menggunakan Laptop HP OMEN 16-AM0999TX dengan Intel Core i9-14900HX, GPU RTX 5070 8GB, RAM 32GB, dan SSD 1TB untuk tahap pelatihan dan pengujian. Eksperimen dilakukan dengan menggunakan bahasa pemrograman Python dan library pendukung scikit-learn, TensorFlow, dan Keras yang sudah banyak digunakan dan diterapkan eksperimen data teks dan data statis [19], [20]. Secara lengkap tahapan penelitian dapat dilihat pada **Gambar 1**.



Gambar 1 Tahapan Penelitian

Berdasarkan tahapan penelitian diatas, data penelitian ini terbagi dua yaitu data statis dan data teks. Data statis merupakan data primer dari hasil pengisian kuesioner kinerja dari 250 sumber daya manusia yang berkerja di perguruan tinggi. Untuk data teks merupakan data gabungan dari data primer dan data sekunder. Atribut data statis terdiri dari demografis, penilaian kinerja, catatan kehadiran, partisipasi pelatihan, dan pertanyaan terbuka mengenai pengalaman kerja.

Data statis berupa data numerik dan kategorikal dilakukan praproses menggunakan dengan *min-max scaling* dan *label encoding* [15]. Entri data yang tidak lengkap atau tidak valid dihapus sebelum analisis. Data numerik dan kategorikal yang sudah difinalisasi akan dilakukan eksperimen dengan menggunakan algoritma *random forest* (RF) [16]. Algoritma RF menganalisis hubungan non-linear dan memberikan *feature importance* berdasarkan atribut *ID*, jenis kelamin, unit/jurusan, lama kerja (tahun), usia, kehadiran (%), pelatihan diikuti, evaluasi kinerja dan menghitung nilai *voting* mayoritas, menggunakan Persamaan (1).

$$\hat{y} = \text{mode}\{h_b(x), b = 1, 2, \dots, B\} \quad (1)$$

Variabel $\hat{y}$ merupakan hasil prediksi kelas, $h_b(x)$ adalah prediksi dari pohon ke- b , dan B adalah jumlah pohon dalam hutan keputusan. Hasil akhir berupa pemilihan kelas mayoritas dari seluruh hasil pohon data yang dibuat berdasarkan data statis kinerja.

Data teks yang berisi feedback kinerja SDM akan dianalisis menggunakan CNN-BiLSTM. Feedback akan dikelompokkan menjadi sentiment negatif, positif dan netral. Sebelum dilakukan melakukan eksperimen menggunakan CNN-BiLSTM, dilakukan penerapan teknik praproses. Teknik pertama Adalah penghapusan tanda baca seperti "efektivitas, kualitas, produktivitas.". Tanda baca seperti koma dan titik dapat dihapus dalam analisis teks jika tidak mempengaruhi makna. Teknik tokenisasi diterapkan dengan cara memotong kalimat "Karyawan yang memiliki motivasi tinggi cenderung lebih produktif" bisa menjadi susunan kata "Karyawan", "Memiliki" dan seterusnya. Teknik *lowercase* diterapkan dengan mengubah kalimat seperti "Kinerja Karyawan yang Baik", setelah diubah menjadi huruf kecil menjadi "kinerja karyawan yang baik".

Teknik penghapusan stopwords dilakukan dengan menghapus kata yang tidak memberikan banyak informasi penting, sebagai contoh dalam kalimat "Karyawan yang bekerja di perusahaan ini mendapatkan penghargaan," kata-kata seperti "yang", "di", "ini" adalah *stopwords* dan dapat dihapus. Teknik stemming diterapkan dengan mengubah kata menjadi menjadi bentuk dasar seperti kata "berkinerja" diproses menjadi bentuk dasar "kerja". Setiap kata yang sudah dalam bentuk potongan kata diproses menggunakan Word2Vec

untuk memetakan kata-kata ke dalam ruang vektor dimensi. Jika teks dalam dataset memiliki panjang yang berbeda, proses padding dilakukan untuk menyamakan panjang teks menggunakan padding di awal atau akhir kalimat, sehingga teks memiliki panjang yang seragam. Hasil dari tahap praproses akan di proses menggunakan model CNN dengan mengekstraksi fitur melalui operasi konvolusi dengan menggunakan Persamaan (2).

$$z_{i,j}^{(k)} = \sum_{m=0}^{M-1} \sum_{n=0}^{N-1} x_{i+m,j+n} \cdot w_{m,n}^{(k)} + b^{(k)} \quad (2)$$

Berdasarkan persamaan diatas dilakukan perhitungan setiap dimensi vector dengan input berupa nilai input x , nilai bobot kernel w dan nilai bias $b^{(k)}$. Layer utama yang digunakan pada CNN berupa Conv1D dan LayerMaxPooling untuk mengurangi dimensi dan menentukan informasi penting hasil proses dari CNN. Hasil dari CNN akan diproses menggunakan BiLSTM menangkap konteks dua arah dengan memproses informasi forward dan backward dengan menggunakan Persamaan (3).

$$h_t^{\rightarrow} = f(Wx_t + Uh_{t-1}^{\rightarrow} + b), h_t^{\leftarrow} = f(Wx_t + Uh_{t+1}^{\leftarrow} + b) \quad (3)$$

Persamaan diatas dihitung berdasarkan nilai $h_t^{\rightarrow}$ dan $h_t^{\leftarrow}$ yang mempresentasikan hidden state forward dan backward dari setiap dimensi vector hasil dari CNN. Setiap kinerja model dievaluasi menggunakan teknik accuracy, precision, recall, F1-score.

3. HASIL DAN ANALISIS

Data statis diambil dari 250 sumber daya manusia yang terdiri atas dosen dan staf dari 11 perguruan tinggi di DKI Jakarta. Komposisi responden berdasarkan jenis kelamin menunjukkan keseimbangan proporsional, meskipun perempuan sedikit lebih banyak (54%) dibanding laki-laki (46%) seperti terlihat pada (Tabel 1). Kategori kelompok usia terbanyak berada pada rentang 41–50 tahun (43,6%), oleh usia 31–40 tahun (38,8%). Jumlah responden pada usia di bawah 30 maupun di atas 50 relatif lebih sedikit dibanding dua kelompok utama seperti yang terlihat pada Tabel 2.

Tabel 1 Deskripsi Sumber Data Primer

Kategori	Sub-Kategori	Jumlah
Jenis Kelamin	Laki	127
	Perempuan	123
	41–50 tahun	109
Kelompok Usia	31–40 tahun	97
	< 30 tahun	23
	> 50 tahun	21
Lama Kerja (tahun)	5–10 tahun	127
	< 5 tahun	123
	> 10 tahun	0
Unit Kerja	Humas, Promosi, dan Kemahasiswaan	39
	Perpustakaan dan Arsip	39
	Sarana dan Prasarana	27
Jumlah Pelatihan	Tata Usaha Fakultas	26
	Administrasi Umum dan SDM	24
	1–2 kali	129
Distribusi Evaluasi Kinerja Responden	≥ 3 kali	87
	Tidak Pernah	34
	Tinggi	138
	Sedang	109
	Rendah	3

Berdasarkan durasi masa kerja, mayoritas responden berada dalam rentang pengalaman kerja 5–10 tahun (50,8%), disusul oleh mereka yang bekerja kurang dari 5 tahun (49,2%) (Tabel 3). Berdasarkan unit atau jurusan, keterlibatan responden paling banyak berasal dari bidang Humas, Promosi, dan Kemahasiswaan serta Perpustakaan dan Arsip (masing-masing 15,6%), diikuti Sarana dan Prasarana (10,8%), Tata Usaha Fakultas (10,4%), dan Administrasi Umum serta SDM (9,6%). Pada kategori kegiatan pelatihan menunjukkan bahwa mayoritas mengikuti pelatihan sebanyak 1–2 kali (51,6%), sementara 34,8% mengikuti tiga kali atau lebih, dan 13,6% belum pernah mengikuti pelatihan. Sebagian besar responden memperoleh skor kinerja dalam kategori tinggi, yakni sebesar 55,2%, berdasarkan hasil distribusi data yang ditampilkan, sedangkan kategori sedang mencapai 43,6%, dan hanya sebagian kecil (1,2%) pada kategori rendah. Detail dataset eksperimen dapat dilihat pada **Tabel 2**.

Tabel 2 Deskripsi Dataset Eksperimen

Eksperimen	Model	Data Primer	Data Sekunder	Pelatihan (70%)	Validasi (10%)	Pengujian (20%)
Data Statis	Random Forest	250	0	175	25	50
Data Teks	CNN-BiLSTM	250	1.000	875	125	250

Dalam tahap eksperimen menggunakan data statis, eksperimen menggunakan algoritma Random Forest (RF) untuk memprediksi kinerja SDM. Atribut yang dianalisis berdasarkan lama kerja dengan interval lama kerja lebih besar atau sama dengan 5 tahun dan lama kerja kurang dari 5 tahun. atribut selanjutnya dengan interval pelatihan diikuti lebih dari 3 dan pelatihan diikuti kurang dari atau sama dengan 3. Hasil eksperimen untuk data statis dapat dilihat pada **Tabel 3**.

Tabel 3 Hasil Kinerja *Random Forest* (RF)

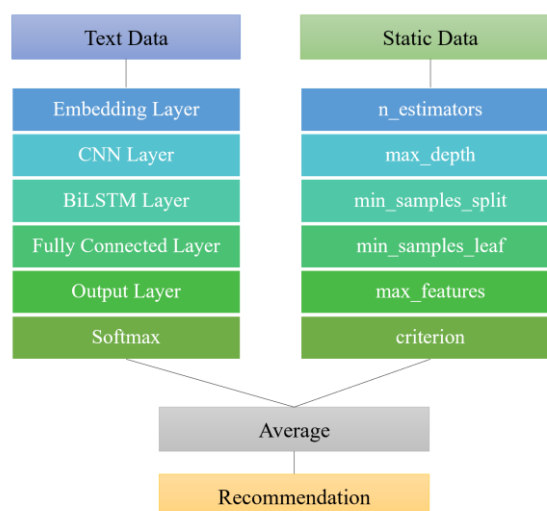
Tahap	Accuracy	Precision	Recall	F1-Score
Eksperimen Pelatihan	90.4%	89.6%	91.5%	92.5%
Eksperimen Validasi	85.7%	83.9%	84.5%	85.2%
Eksperimen Pengujian	82.5%	83.5%	80.2%	81.5%

Hasil eksperimen untuk data teks yang menggunakan model CNN-BiLSTM mendapatkan kinerja yang baik, meskipun ada sedikit penurunan akurasi saat beralih dari pelatihan ke validasi dan pengujian. Pada eksperimen pelatihan, model menghasilkan akurasi 92.6%, *precision* 90.2%, *recall* 90.4%, dan F1-score 91.4%, yang menunjukkan bahwa model dapat mengklasifikasikan data dan memiliki keseimbangan antara *precision* dan *recall*. Namun, pada eksperimen validasi, terdapat penurunan performa dengan akurasi 87.4%, *precision* 85.7%, *recall* 86.5%, dan F1-score 86.2%, yang menunjukkan bahwa model mengalami sedikit penurunan ketika diuji pada data pelatihan. Pada eksperimen pengujian, akurasi model menurun lebih jauh menjadi 84.4%, dengan *precision* 83.6%, *recall* 82.7%, dan F1-score 84.6%. Hasil eksperimen dapat dilihat pada **Tabel 4**.

Tabel 4 Hasil Kinerja CNN-BiLSTM

Tahap	Accuracy	Precision	Recall	F1-Score
Eksperimen Pelatihan	92.6%	90.2%	90.4%	91.4%
Eksperimen Validasi	87.4%	85.7%	86.5%	86.2%
Eksperimen Pengujian	84.4%	83.6%	82.7%	84.6%

Penelitian ini mengembangkan model analisis kinerja Sumber Daya Manusia (SDM) di perguruan tinggi dengan menggunakan data statis dan data teks berbasis multi-modal. Untuk data statis, algoritma Random Forest diterapkan dengan parameter yaitu *n_estimators*, *max_depth*, *min_samples_split*, *min_samples_leaf*, *max_features*, dan *criterion*. Untuk data teks, digunakan model yang menggabungkan convolutional neural network (CNN) dan *Bidirectional Long Short-Term Memory* (BiLSTM) untuk memproses data teks. Data teks kemudian diproses melalui *fully connected layer* dan *output layer*, dengan softmax digunakan untuk klasifikasi. Arsitektur *multi-modal* yang diusulkan dapat dilihat pada **Gambar 2**.



Gambar 2 Arsitektur Multi-modal Data Untuk Analisis Kinerja SDM

4. KESIMPULAN

Berdasarkan hasil eksperimen, penelitian ini mendapatkan hasil bahwa model *random forest* dan CNN-BiLSTM memiliki kinerja yang baik dalam analisis kinerja Sumber Daya Manusia (SDM) di perguruan tinggi, dengan hasil yang cukup memadai meskipun ada penurunan performa pada eksperimen validasi dan pengujian. Pada eksperimen pelatihan, *random forest* (RF) mencapai akurasi 90.4%, precision 89.6%, recall 91.5%, dan F1-score 92.5%, yang menunjukkan kemampuan model dalam memproses data statis dengan baik. Namun, pada eksperimen validasi dan pengujian, hasil RF menurun dengan akurasi 85.7% dan 82.5%, serta F1-score 85.2% dan 81.5%. Jika dibandingkan dengan model CNN-BiLSTM untuk data teks, kinerja model yang lebih tinggi pada eksperimen pelatihan dengan akurasi 92.6%, precision 90.2%, recall 90.4%, dan F1-score 91.4%. Namun, seperti halnya Random Forest, model CNN-BiLSTM juga mengalami penurunan pada eksperimen validasi dan pengujian, dengan akurasi 87.4% dan 84.4%, serta F1-score 86.2% dan 84.6%. Penurunan ini menunjukkan bahwa meskipun kedua model dapat membangun model dengan baik dalam pelatihan, hasil menunjukkan adanya kesalahan dalam prediksi data validasi dan pengujian.

UCAPAN TERIMA KASIH

Ucapan terima kasih diberikan Direktorat Penelitian Dan Pengabdian Kepada Masyarakat Kementerian Pendidikan Tinggi, Sains, dan Teknologi, LLDIKTI Wilayah III dan Universitas Dian Nusantara yang telah mendukung pelaksanaan program penelitian ini melalui hibah tahun anggaran 2025 dengan nomor kontrak 0971/LL3/AL.04/2025 dan No.124/C3/DT.05.00/PM/2025.

REFERENSI

- [1] A. Alam, "Impact of university's human resources practices on professors' occupational performance: empirical evidence from India's higher education sector," in *Inclusive businesses in developing economies: Converging people, profit, and corporate citizenship*, Springer, 2022, pp. 107–131.
- [2] R. S. Hassan, H. M. G. Amin, and H. Ghoneim, "Decent work and innovative work behavior of academic staff in higher education institutions: the mediating role of work engagement and job self-efficacy," *Humanit. Soc. Sci. Commun.*, vol. 11, no. 1, pp. 1–19, 2024.
- [3] M. Alenezi, "Digital learning and digital institution in higher education," *Educ. Sci.*, vol. 13, no. 1, p. 88, 2023.
- [4] J. Chaanine, "Strengthening Lebanese health care: exploring the impact of organizational culture on employee loyalty through trust and job satisfaction," *Leadersh. Heal. Serv.*, vol. 38, no. 2, pp. 263–279, 2025.
- [5] H. A. Obeng, R. Arhinful, L. Mensah, and J. S. Owusu-Sarfo, "Assessing the influence of the knowledge management cycle on job satisfaction and organizational culture considering the interplay of employee engagement," *Sustainability*, vol. 16, no. 20, p. 8728, 2024.
- [6] M. S. A. Basalamah and A. As'ad, "The role of work motivation and work environment in improving job satisfaction," *Golden Ratio Hum. Resour. Manag.*, vol. 1, no. 2, pp. 94–103, 2021.
- [7] M. Mohiuddin, E. Hosseini, S. B. Faradonbeh, and M. Sabokro, "Achieving human resource management sustainability in universities," *Int. J. Environ. Res. Public Health*, vol. 19, no. 2, p. 928, 2022.
- [8] F. Dong, "Random Forest Algorithm for HR Data Classification and Performance Analysis in Cloud Environments," *Int. J. Adv. Comput. Sci. Appl.*, vol. 15, no. 11, 2024.
- [9] E. M. T. A. Alsaadi, S. F. Khlebus, and A. Alabaichi, "Identification of human resource analytics using machine learning algorithms," *Telkomnika (Telecommunication Comput. Electron. Control.)*, vol. 20, no. 5, pp. 1004–1015, 2022.
- [10] Y. Xu, "Design of human resource allocation algorithm based on improved random forest," in *2021 International Conference on Aviation Safety and Information Technology*, 2021, pp. 656–661.
- [11] J. Lee and J. H. Song, "How does algorithm-based HR predict employees' sentiment? Developing an employee experience model through sentiment analysis," *Ind. Commer. Train.*, vol. 56, no. 4, pp. 273–289, 2024.
- [12] Y. Zhang and E. Qi, "Happy work: Improving enterprise human resource management by predicting workers' stress using deep learning," *PLoS One*, vol. 17, no. 4, p. e0266373, 2022.
- [13] P. P. Bhatt, J. V. Nasriwala, and R. R. Savant, "Template-Based Thinning Method for Handwritten Gujarati Character's Strokes and its Classification for Writer-Dependent Gujarati Font Synthesis," in *Advanced Machine Intelligence and Signal Processing*, Springer, 2022, pp. 203–216.
- [14] S. Cui, "Context-Aware TinyML Model for Korean Handwriting Recognition Under Online Assisted

- Learning Scenes,” *Internet Technol. Lett.*, p. e636, 2025.
- [15] N. Kumar, A. Chikmath, G. Y. BR, H. R. Naidu, and D. Acharya, “Interpreting doctor notes using handwriting recognition and deep learning techniques: a survey,” in *2023 International Conference on Advances in Electronics, Communication, Computing and Intelligent Information Systems (ICAECIS)*, IEEE, 2023, pp. 703–708.
- [16] N. M. Tahir, A. N. Ausat, U. I. Bature, K. A. Abubakar, and I. Gambo, “Off-line handwritten signature verification system: Artificial neural network approach,” *Int. J. Intell. Syst. Appl.*, vol. 10, no. 1, p. 45, 2021.
- [17] J. Miao, P. Liu, C. Chen, and Y. Qiao, “Normal Template Mapping: An Association-Inspired Handwritten Character Recognition Model,” *Cognit. Comput.*, vol. 16, no. 3, pp. 1103–1112, 2024.
- [18] S. K. Singh and A. Chaturvedi, “An efficient multi-modal sensors feature fusion approach for handwritten characters recognition using Shapley values and deep autoencoder,” *Eng. Appl. Artif. Intell.*, vol. 138, p. 109225, 2024.