

PENERAPAN LONG SHORT-TERM MEMORY UNTUK PREDIKSI TURNOVER KARYAWAN BERBASIS PEARSON CORRELATION COEFFICIENT DAN ANALISIS LINTAS DEPARTEMEN

Sjaeful Irwan¹, Sebrina Zahratusy Syita², Muhammad Rezqi Al Faresi³

¹Teknik Informatika, Universitas Bani Saleh, Bekasi, Indonesia, 17113

²Manajemen, FEB Universitas Brawijaya, Malang, Indonesia, 65145

³Sains Data, Universitas Terbuka, Tangerang Selatan, Indonesia, 15437

Email: sjaeful@ubs.ac.id¹, sebrinazahra@student.ub.ac.id², 058025017@ecampus.ut.ac.id³

ABSTRAK

Penelitian ini bertujuan untuk mengembangkan dan mengevaluasi model prediksi perputaran karyawan dengan mengkombinasikan metode Long Short-Term Memory (LSTM) dan metode statistik Pearson Correlation Coefficient (PCC) untuk seleksi fitur serta analisis lintas unit organisasi. Tantangan utama dalam analisis sumber daya manusia terletak pada keterbatasan algoritma statis (bukan algoritma dinamis/algoritma adaptif) dalam memahami pola ketergantungan terhadap waktu yang berlangsung dalam jangka panjang. Dengan menggunakan data simulasi yang mereplikasikan karakteristik dataset yang mengacu pada IBM HR Analytics yang terdiri atas 1.470 karyawan dan 28 fitur, model LSTM yang diusulkan dapat mencapai tingkat akurasi 71,86%, presisi 69,47%, recall 55,00%, dan F1-score sebesar 61,40%. Hasil evaluasi berdasarkan unit organisasi diperoleh bahwa Departemen Sumber Daya Manusia memiliki kinerja tertinggi dengan akurasi 73,42%, disusul oleh Departemen Penjualan sebesar 71,10%, dan Departemen Riset dan Pengembangan (R&D) sebesar 69,74%. Hasil dari analisis PCC menunjukkan faktor lembur sebagai pemicu paling dominan dengan nilai korelasi 0,3366, diikuti oleh kepuasan kerja dengan korelasi -0,2038 dan keseimbangan kehidupan dan kerja sebesar -0,1407. Dalam penelitian ini, PCC tidak hanya dimanfaatkan untuk mengurangi jumlah dimensi data, juga untuk memetakan korelasi unik dari faktor-faktor penyebab turnover di setiap unit organisasi, sehingga mampu memberikan rekomendasi strategi retensi yang lebih spesifik dan terarah bagi pihak manajemen organisasi.

Kata Kunci: Long Short-Term Memory (LSTM), Perputaran Karyawan, Pearson Correlation Coefficient, Analisis Departemen, Seleksi Fitur.

ABSTRACT

This study aims to develop and evaluate an employee turnover prediction model by combining the Long Short-Term Memory (LSTM) method and the Pearson Correlation Coefficient (PCC) statistical method for feature selection and cross-organizational unit analysis. The main challenge in human resource analytics lies in the limitations of static algorithms (non-dynamic/non-adaptive algorithms) in understanding long-term temporal dependency patterns. Using simulated data that replicate the characteristics of the IBM HR Analytics dataset, consisting of 1,470 employees and 28 features, the proposed LSTM model achieved an accuracy of 71.86%, precision of 69.47%, recall of 55.00%, and an F1-score of 61.40%. The evaluation results based on organizational units showed that the Human Resources Department achieved the highest performance with an accuracy of 73.42%, followed by the Sales Department at 71.10%, and the Research and Development (R&D) Department at 69.74%. The PCC analysis results indicated that overtime was the most dominant triggering factor with a correlation value of 0.3366, followed by job satisfaction with a correlation of -0.2038 and work-life balance with a correlation of -0.1407. In this study, PCC was not only utilized to reduce data dimensionality, but also to map the unique correlations of turnover-causing factors within each organizational unit, thereby providing more specific and targeted retention strategy recommendations for organizational management.

Keywords: Long Short-Term Memory (LSTM), Employee Turnover, Pearson Correlation Coefficient, Department Analysis, Feature Selection.

1. PENDAHULUAN

Dalam bisnis modern, pengelolaan SDM menjadi kunci keberhasilan organisasi. Pasalnya, tingkat *turnover* yang tinggi bisa menimbulkan kerugian finansial signifikan dari biaya rekrutmen dan pelatihan ulang. Dari data simulasi 1.470 karyawan,

angka *turnover* tercatat 40,6% (597 karyawan), dengan rincian: Sales 40,7%, R&D 40,4%, dan HR 41,8%.

Penelitian ini mengidentifikasi tiga tantangan utama prediksi *turnover*. Pertama, bersifat non-linear sehingga model statistik linear kurang optimal. Kedua, adanya ketergantungan jangka panjang yang butuh kapasitas

memori lebih besar dari RNN konvensional (Hochreiter & Schmidhuber, 1997). Ketiga, pemicu *turnover* berbeda antar departemen, butuh pendekatan spesifik.

Di sinilah LSTM berperan. Model ini mengatasi *vanishing gradient* pada RNN dengan kemampuannya menyimpan informasi penting dari riwayat kerja karyawan lewat mekanisme *gate*. Mahammad dkk. (2025) membuktikan LSTM-RNN-HMM efektif memprediksi *turnover*. Ganapathisamy & Narayan (2024) mengembangkan LSTM dengan pendekatan *Brownian Motion Butterfly Optimization* dan meraih akurasi unggul. Morteza pour Shiri dkk. (2025) menggunakan Bi-TCN, sementara Wardhani & Lhaksana (2022) menerapkan Logistic Regression dengan seleksi fitur dan mendapat hasil cukup baik.

PCC digunakan untuk mengukur hubungan linier antar variabel dan mengidentifikasi fitur paling relevan. Romadloni & Pardede (2019) menunjukkan bahwa seleksi fitur berbasis Pearson mampu meningkatkan kinerja klasifikasi pada Logistic Regression, Naïve Bayes, dan SVM.

Di dalam penelitian ini dilakukan integrasi analisis korelasi lintas departemen. Setiap unit bisnis memiliki pemicu *turnover* khas. Pendekatan tradisional dengan indikator statis, menurut Yakymiv & Yehorchenkova (2026), kurang mampu mencerminkan perilaku dinamis karyawan di lingkungan proyek.

Implementasi model LSTM dilakukan dengan Python di Google Colab. Alasannya sederhana: akses mudah, GPU gratis, dan hasil bisa langsung tersimpan ke Google Drive. Pada akhirnya, integrasi LSTM, PCC, dan analisis lintas departemen menawarkan kerangka prediksi *turnover* yang lebih adaptif dan akurat—sesuai kebutuhan manajemen proyek modern yang dinamis dan bertransformasi digital.

2. TINJAUAN PUSTAKA

2.1 Long Short-Term Memory (LSTM)

Long Short-Term Memory (LSTM) pertama kali diperkenalkan oleh Hochreiter & Schmidhuber (1997) untuk mengatasi masalah *vanishing gradient* pada RNN konvensional. Bedanya dengan RNN biasa, LSTM punya tiga mekanisme *gate* utama—yaitu *input gate*, *forget gate*, dan *output gate*—yang memungkinkan informasi bisa dipertahankan dalam jangka panjang. Lewat mekanisme ini, LSTM juga bisa mengatasi konflik antara bobot input dan output karena kontrolnya lebih peka terhadap konteks.

Yakymiv & Yehorchenkova (2026) menambahkan beberapa poin penting. LSTM dirancang khusus untuk menghindari gradien yang memudar pada sekuens

panjang, sehingga sangat efektif buat prediksi deret waktu. Kemampuan LSTM untuk "mengingat" peristiwa penting dalam riwayat data juga jadi nilai plus. Selain itu, LSTM bekerja dengan sekuens berpanjangan tetap melalui apa yang disebut *time window*, bukan dengan contoh tunggal seperti halnya regresi klasik.

Kalau ditarik ke konteks analisis sumber daya manusia, LSTM bisa dipakai untuk menangkap pola-pola temporal. Misalnya fluktuasi kepuasan kerja dari waktu ke waktu, tren lembur karyawan, atau perubahan beban kerja yang terjadi secara periodik. Beberapa penelitian sebelumnya sudah menunjukkan hasil yang menjanjikan. Mahammad dkk. (2025) membuktikan bahwa model LSTM-RNN-HMM efektif buat prediksi *turnover* dengan tingkat akurasi yang tinggi. Ganapathisamy & Narayan (2024) juga mengembangkan LSTM dengan pendekatan *Brownian Motion Butterfly Optimization* dan berhasil mencapai hasil yang tergolong superior. Sementara itu, Morteza pour Shiri dkk. (2025) memilih pendekatan lain, yaitu Bi-TCN (*Bidirectional Temporal Convolutional Network*), dan dilaporkan bahwa modelnya mampu mencapai akurasi yang kompetitif pada dataset IBM.

Ponmalar dkk. (2024) membuktikan bahwa LSTM—terutama ketika dipadukan dengan ensemble methods seperti CatBoost, AdaBoost, dan LightGBM—bisa mendeteksi sinyal-sinyal *turnover* yang tidak kasat mata oleh model konvensional, terutama dari data karyawan yang sifatnya longitudinal. Jadi, bukan sekadar teori, pendekatan ini benar-benar bisa dipakai tim HR untuk menyusun strategi retensi secara lebih proaktif.

2.2 Pearson Correlation Coefficient (PCC)

Pearson Correlation Coefficient (PCC) pertama kali dikembangkan oleh Karl Pearson pada tahun 1895. Secara sederhana, metode ini dipakai untuk mengukur seberapa kuat dan ke arah mana hubungan linier antara dua variabel. Hasilnya berupa koefisien korelasi (r) yang nilainya bergerak dari -1 hingga +1. Kalau nilainya mendekati +1, artinya hubungan positifnya kuat; kalau mendekati -1, hubungan negatifnya kuat; dan kalau mendekati nol, bisa dibilang hampir tidak ada hubungan linier sama sekali.

Romadloni & Pardede (2019) memberikan penjelasan yang cukup membantu. Analisis korelasi pada dasarnya digunakan untuk mengetahui kekuatan hubungan antara dua variabel, sementara variabel lainnya dikendalikan sebagai variabel kontrol. Kebetulan, data yang diteliti termasuk jenis data interval, jadi teknik Pearson ini memang pilihan yang tepat. Nilai korelasinya kemudian dipakai untuk seleksi fitur—khususnya fitur-fitur dengan

nilai korelasi Pearson tertinggi, karena itu dianggap paling berpengaruh.

Hasil penelitian Romadloni & Pardede (2019) sendiri cukup meyakinkan dan membuktikan bahwa PCC tidak hanya efektif untuk seleksi fitur, tapi juga mampu meningkatkan akurasi klasifikasi secara signifikan.

Sejalan dengan temuan dalam penelitian ini, Musanga & Chibaya (2023) juga mengkonfirmasi efektivitas Pearson Correlation Coefficient sebagai alat seleksi fitur. Dalam studi yang dilakukan, kombinasi antara PCC dan algoritma Random Forest terbukti mampu meningkatkan akurasi prediksi turnover secara signifikan. Hal ini memperkuat bukti bahwa metode PCC tidak terbatas manfaatnya hanya pada model deep learning seperti LSTM, tetapi juga sangat efektif jika diterapkan pada model machine learning konvensional.

2.3 Turnover Karyawan dan Faktor Pemicu

Secara sederhana, turnover adalah pergerakan keluar tenaga kerja dari suatu organisasi. Mahammad dkk. (2025) berpendapat bahwa sebelum menurunkan angka turnover, organisasi perlu menyelidiki penyebabnya terlebih dahulu—misalnya melalui pelatihan model, ekstraksi fitur, dan pra-pemrosesan data.

Beberapa penelitian sebelumnya telah mencoba pendekatan berbeda. Wardhani & Lhaksana (2022) berhasil memprediksi employee attrition dengan regresi logistik dan seleksi fitur. Morteza pour Shiri dkk. (2025) menegaskan bahwa prediksi turnover penting agar perusahaan bisa bertindak proaktif, karena penyebab attrition beragam: alasan pribadi, ketidakpuasan kerja, kompensasi tidak memadai, hingga kondisi kerja kurang menguntungkan.

3. METODE PENELITIAN

3.1 Rancangan Penelitian

Penelitian ini menggunakan pendekatan eksperimental kuantitatif dengan alur kerja sebagai berikut:

- (1) pengumpulan data simulasi IBM HR Analytics yang mencakup 1.470 karyawan dan 28 fitur;
- (2) preprocessing data yang meliputi label encoding untuk data kategorikal dan min-max scaling untuk normalisasi;
- (3) seleksi fitur berbasis Pearson Correlation Coefficient untuk menentukan 10 fitur dengan korelasi tertinggi terhadap target;
- (4) analisis korelasi per departemen untuk mengidentifikasi pemicu turnover yang unik di setiap unit bisnis;
- (5) pembangunan arsitektur LSTM dengan 2 layer (hidden size 64) dan dropout 0,2; serta

- (6) evaluasi model menggunakan confusion matrix untuk menghitung akurasi, presisi, recall, dan F1-score.

Untuk implementasinya, seluruh proses pelatihan dan evaluasi model LSTM dilakukan menggunakan Python. Beberapa library yang dipakai antara lain Pandas, NumPy, Scikit-learn, PyTorch, dan Matplotlib. Soal platform, Google Colab dengan akselerasi GPU T4 dipilih—tujuannya jelas, agar komputasi bisa berjalan lebih cepat. Data dan hasil pengolahan otomatis tersimpan ke Google Drive, yang berguna buat dokumentasi maupun kalau nanti ingin mereproduksi hasil penelitian.

Perlu dicatat, model LSTM itu cara kerjanya berbeda dengan regresi klasik. Seperti yang dijelaskan oleh Yakymiv & Yehorchenkova (2026), LSTM tidak bekerja dengan contoh tunggal. Sebaliknya, ia menggunakan pendekatan *time window*, di mana setiap sekuens yang diproses memiliki bentuk $T \times K$. Di sini, T adalah panjang jendela waktu, sementara K adalah jumlah fitur pada setiap langkah waktu.

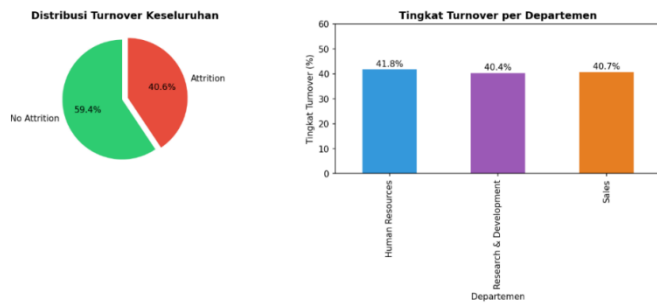
3.2 Data Penelitian

Data penelitian berupa data simulasi dengan karakteristik yang meniru dataset mengacu pada IBM HR Analytics. Tabel 1 menyajikan karakteristik lengkap data yang digunakan dalam penelitian ini.

Tabel 1. Karakteristik Data Penelitian

Parameter	Nilai
Jumlah sampel	1.470 karyawan
Jumlah Attrition (Yes)	597 karyawan (40,6%)
Jumlah Attrition (No)	873 karyawan (59,4%)
Jumlah fitur awal	28 fitur
Jumlah fitur setelah seleksi PCC	10 fitur
Pembagian data training	1.028 sampel (70%)
Pembagian data validation	147 sampel (10%)
Pembagian data test	295 sampel (20%)

Total data yang digunakan dalam penelitian ini berjumlah 1.470 karyawan. Dari jumlah tersebut, sebanyak 597 karyawan (40,6%) termasuk dalam kategori *attrition* (keluar), sementara sisanya—873 karyawan (59,4%)—tidak mengalami *attrition*. Komposisi kelas seperti ini tergolong relatif seimbang. Kondisi ini cukup baik dan mendukung proses pelatihan model prediksi karena tidak ada dominasi yang terlalu timpang antara satu kelas dengan kelas lainnya.



Gambar 1. Distribusi Data

Gambar 1 menyajikan dua visualisasi. Pertama, diagram pie distribusi *turnover*: 40,6% karyawan mengalami *attrition*, 59,4% tidak. Kedua, diagram batang *turnover* per departemen—Sales 40,7%, R&D 40,4%, HR 41,8%. Ketiganya tidak berbeda jauh.

Penelitian ini awalnya menggunakan 28 fitur (kondisi kerja, kepuasan kerja, pendapatan, lingkungan kerja). Setelah seleksi dengan PCC, tersisa 10 fitur utama yang paling signifikan pengaruhnya terhadap *attrition*. Tujuannya agar model lebih efisien tanpa kehilangan informasi penting.

Selanjutnya, dataset dibagi menjadi tiga: training 1.028 sampel (70%), validation 147 sampel (10%), dan test 295 sampel (20%). Pembagian ini supaya model bisa dilatih, divalidasi, dan diuji secara optimal—dengan harapan performa prediksi dan generalisasinya baik.

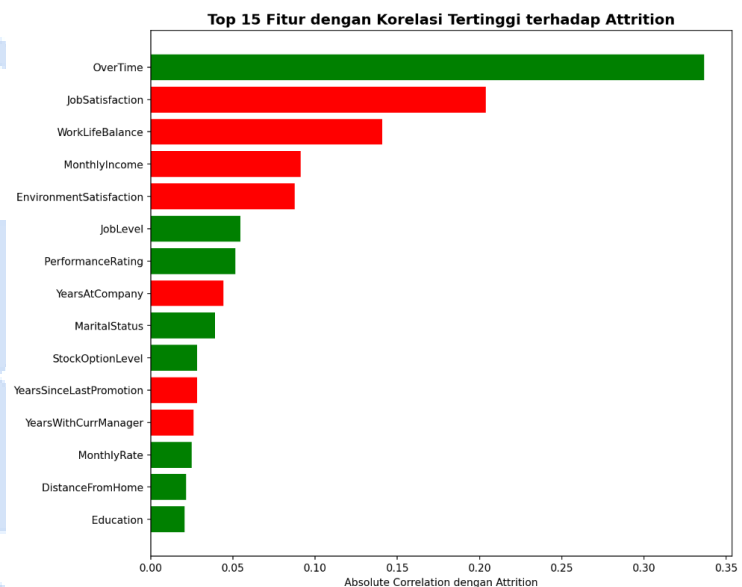
3.3 Rekayasa Fitur dan Seleksi dengan PCC

Sebanyak 28 fitur awal diolah melalui tiga tahapan. Pertama, *label encoding*—ini diterapkan pada 7 fitur kategorikal, yaitu Gender, MaritalStatus, EducationField, Department, JobRole, OverTime, dan BusinessTravel. Kedua, *min-max scaling* untuk menormalkan seluruh fitur numerik ke rentang [0,1]. Ketiga, seleksi fitur menggunakan Pearson Correlation Coefficient (PCC), di mana diambil fitur-fitur dengan nilai korelasi tertinggi terhadap target *Attrition*.

Tabel 2. Top 15 Fitur dengan Korelasi Tertinggi terhadap Attrition

Peringkat	Fitur	Korelasi	Arah
1	OverTime (Lembur)	0,3366	Positif
2	JobSatisfaction (Kepuasan Kerja)	-0,2038	Negatif
3	WorkLifeBalance (Keseimbangan Kerja-Kehidupan)	-0,1407	Negatif
4	MonthlyIncome (Pendapatan Bulanan)	-0,0912	Negatif
5	EnvironmentSatisfaction (Kepuasan Lingkungan Kerja)	-0,0877	Negatif
6	JobLevel	0,0544	Positif

Peringkat	Fitur	Korelasi	Arah
7	PerformanceRating	0,0516	Positif
8	YearsAtCompany	-0,0443	Negatif
9	MaritalStatus	0,0390	Positif
10	StockOptionLevel	0,0282	Positif
11	YearsSinceLastPromotion	-0,0282	Negatif
12	YearsWithCurrManager	-0,0261	Negatif
13	MonthlyRate	0,0249	Positif
14	DistanceFromHome	0,0214	Positif
15	Education	0,0208	Positif

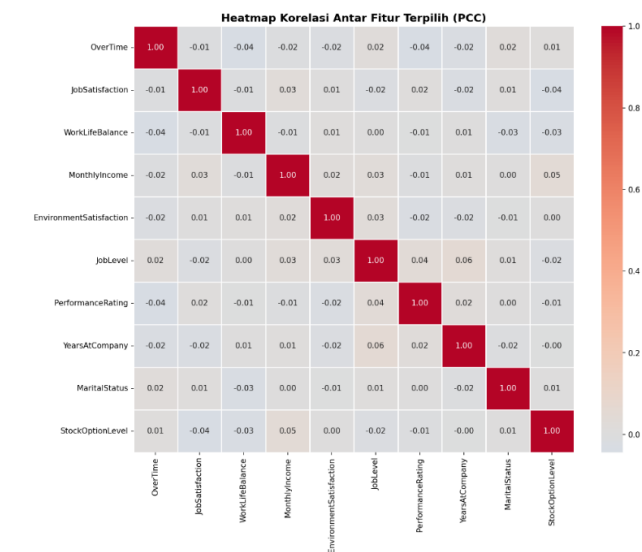


Gambar 2. 15 fitur dengan nilai korelasi tertinggi terhadap attrition.

Gambar 2 menampilkan 15 fitur dengan korelasi tertinggi terhadap *attrition*. Warna hijau = korelasi positif (risiko *turnover* naik), merah = korelasi negatif (risiko turun).

Dari Tabel 2, *OverTime* tertinggi (0,3366): makin sering lembur, makin tinggi risiko keluar. *JobSatisfaction* (-0,2038) dan *WorkLifeBalance* (-0,1407) berkorelasi negatif signifikan. *MonthlyIncome* dan *EnvironmentSatisfaction* juga negatif.

Fitur dengan korelasi positif rendah: *JobLevel*, *PerformanceRating*, *MaritalStatus*. Sementara *YearsAtCompany*, *YearsSinceLastPromotion*, *YearsWithCurrManager* berkorelasi negatif. Jasi, keseimbangan kerja, kepuasan kerja, dan beban lembur adalah tiga faktor utama pemicu *turnover*.



Gambar 3. Heatmap korelasi antarfitur terpilih

Gambar 3 menampilkan heatmap korelasi antar fitur terpilih, seperti OverTime, JobSatisfaction, WorkLifeBalance, MonthlyIncome, dan EnvironmentSatisfaction.

Beberapa pola menarik terlihat. OverTime berkorelasi positif rendah dengan JobLevel, namun berkorelasi negatif dengan WorkLifeBalance. Artinya, karyawan yang sering lembur cenderung memiliki keseimbangan kerja-kehidupan yang lebih buruk.

Heatmap ini menggunakan skala warna biru (korelasi negatif) hingga merah (korelasi positif), dengan angka korelasi di setiap sel untuk memudahkan interpretasi.

Sepuluh fitur dengan korelasi tertinggi terhadap Attrition adalah: OverTime, JobSatisfaction, WorkLifeBalance, MonthlyIncome, EnvironmentSatisfaction, JobLevel, PerformanceRating, YearsAtCompany, MaritalStatus, dan StockOptionLevel.

3.4 Analisis Korelasi per Departemen

Tabel 3 menyajikan hasil analisis korelasi per departemen untuk lima faktor teratas yang memengaruhi turnover di masing-masing unit bisnis.

Tabel 3. Analisis Korelasi per Departemen (Top 5 Faktor)

Departemen	Jumlah Sampel	Faktor 1	Faktor 2	Faktor 3	Faktor 4	Faktor 5
Sales	519	Over Time (0,3498)	JobSatisfaction (0,1897)	EnvironmentSatisfaction (0,1317)	WorkLifeBalance (0,1284)	JobRole (0,1099)
Research & Development	793	Over Time (0,3002)	JobSatisfaction (0,2156)	WorkLifeBalance (0,1544)	JobLevel (0,0869)	PerformanceRating (0,0700)

Departemen	Jumlah Sampel	Faktor 1	Faktor 2	Faktor 3	Faktor 4	Faktor 5
Human Resources	158	Over Time (0,4732)	JobSatisfaction (0,1898)	MonthlyIncome (0,1738)	StockOptionLevel (0,1678)	BusinessTravel (0,1612)

Dari Tabel 3, OverTime (lembur) menjadi faktor paling dominan penyebab turnover di semua departemen, dengan nilai korelasi tertinggi di HR (0,4732), disusul Sales (0,3498), lalu R&D (0,3002). Sementara itu, JobSatisfaction juga konsisten muncul sebagai faktor penting di seluruh departemen.

Tiap departemen punya karakteristik sendiri. Di Sales, selain lembur dan kepuasan kerja, EnvironmentSatisfaction, WorkLifeBalance, dan JobRole turut berpengaruh—wajar karena tekanan target penjualan cukup tinggi. Di R&D, WorkLifeBalance jadi faktor penting setelah lembur dan kepuasan kerja, mengingat karyawan di bidang ini butuh keseimbangan agar tidak cepat burnout. Lain lagi dengan HR: faktor finansial (MonthlyIncome dan StockOptionLevel) serta BusinessTravel berpengaruh signifikan.

Jadi, setiap departemen punya pemicu turnover berbeda, sehingga strategi retensi harus spesifik. Meski begitu, OverTime dan JobSatisfaction adalah dua faktor dominan di mana-mana. MonthlyIncome hanya masuk tiga besar di HR (0,1738), sementara WorkLifeBalance jadi faktor ketiga terpenting di R&D.

3.5 Penerapan Arsitektur LSTM

Dalam membangun arsitektur model LSTM, berpegang pada satu prinsip penting: LSTM termasuk model diferensial yang cukup sensitif terhadap skala variabel input. Karena itu, normalisasi data bukan sekadar langkah formalitas—ini benar-benar krusial sebelum pelatihan model dimulai.

Adapun arsitektur yang digunakan adalah sebagai berikut. *Input layer* terdiri dari 10 fitur yang sudah melalui proses seleksi PCC. Selanjutnya, memakai dua *layer LSTM* dengan masing-masing 64 *hidden units per layer*, ditambah *dropout* sebesar 0,2 untuk mencegah *overfitting*. Setelah itu, ada *fully connected layer* dengan 32 unit, yang kemudian terhubung ke *output layer* berupa satu unit dengan fungsi aktivasi *sigmoid*.

Tabel 4 menyajikan konfigurasi hyperparameter yang digunakan pada model Long Short-Term Memory (LSTM) untuk prediksi turnover karyawan.

Tabel 4. Hyperparameter Model

Parameter	Nilai	Keterangan
Sequence Length	1	Data <i>cross-sectional</i>
Hidden Size	64	Dimensi <i>hidden state</i>
Num LSTM Layers	2	Jumlah layer LSTM
Dropout	0,2	Tingkat regularisasi
Learning Rate	0,001	Optimizer Adam
Batch Size	32	Jumlah sampel per <i>batch</i>
Epochs	100	Maksimum iterasi (<i>early stopping</i> di epoch 20)
Loss Function	BCE	<i>Binary Cross Entropy</i>

Dalam proses menggunakan *sequence length* = 1. Karena data yang dipakai bersifat *cross-sectional*, jadi tiap sampel cukup diperlakukan sebagai satu urutan tunggal—tidak perlu sekuens panjang seperti pada data deret waktu biasa.

Untuk *hidden size*, dipilih 64. Angka ini sudah cukup optimal untuk menangkap pola antar fitur tanpa menambah kompleksitas yang tidak perlu. Selain itu, ditambah dua *layer* LSTM agar model bisa mempelajari pola yang lebih kompleks. Lalu, *dropout* 0,2 dipasang untuk mengurangi risiko *overfitting*.

Untuk *optimizer*, dipakai Adam dengan *learning rate* 0,001. Nilai ini dipilih karena terbukti mampu menjaga stabilitas sekaligus mencapai konvergensi yang optimal. Proses pelatihan sendiri dijalankan dengan *batch size* 32 dan maksimal 100 *epoch*. Tapi, juga diterapkan *early stopping*—pelatihan berhenti di *epoch* ke-20 karena setelah itu performa sudah tidak meningkat signifikan.

Sebagai fungsi *loss*, digunakan *Binary Cross Entropy* (BCE), menyesuaikan dengan masalah klasifikasi biner yang dihadapi (*Attrition* vs *No Attrition*). Secara keseluruhan, kombinasi *hyperparameter* ini dirancang agar model yang dihasilkan stabil, efektif, dan mampu generalisasi dengan baik pada data *turnover* karyawan.

Gambar 4. Grafik Training Loss dan Validation Loss selama proses training

Pada Gambar 4, disajikan grafik *Training Loss* dan *Validation Loss* dari *epoch* 0 hingga 17,5. Kedua kurva—yang biru dan oranye—terlihat menurun secara stabil dan mulai konvergen di sekitar *epoch* 15 hingga 20. Dari sini, diputuskan untuk menghentikan pelatihan di *epoch* ke-20 memang tepat, karena setelah itu performa model sudah tidak meningkat secara berarti.

Yang menarik, pola konvergensi ini juga menandakan bahwa model tidak mengalami *overfitting* dan kemampuan generalisasinya tergolong baik.

Untuk urusan teknis, implementasi LSTM dilakukan menggunakan PyTorch di Google Colab. Kode Python dikembangkan bersifat modular, mencakup beberapa tahap sekaligus: pembangkitan data, *preprocessing*, seleksi fitur dengan PCC, pelatihan, evaluasi, hingga visualisasi. Hasilnya, seluruh proses bisa berjalan otomatis dalam satu sesi tanpa perlu intervensi manual di tiap tahap.

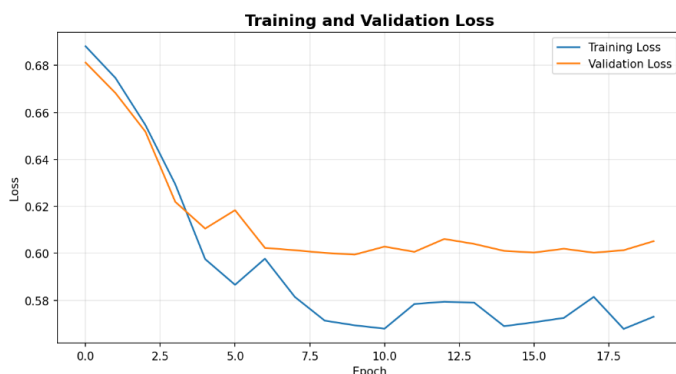
3.6 Metode Evaluasi

Untuk mengevaluasi kinerja model, digunakan empat metrik standar dalam klasifikasi. Pertama, akurasi—yaitu proporsi prediksi benar dari seluruh data. Kedua, presisi, yang mengukur proporsi prediksi positif yang memang benar-benar positif. Ketiga, recall (atau sensitivitas), yaitu proporsi aktual positif yang berhasil terdeteksi dengan benar oleh model. Keempat, F1-Score, yang merupakan rata-rata harmonik dari presisi dan recall—berguna untuk menyeimbangkan keduanya, terutama kalau data tidak seimbang.

Pilihan metrik ini sebenarnya tidak dibuat sembarangan, merujuk pada penelitian-penelitian sebelumnya. Morteza pour Shiri dkk. (2025), misalnya, juga menggunakan akurasi, presisi, recall, F1-score, plus AUC untuk mengevaluasi model Bi-TCN-nya. Demikian pula Wardhani & Lhaksana (2022) yang menerapkan metrik serupa pada model prediksi *attrition*—nyadengan regresi logistik.

4. HASIL DAN PEMBAHASAN

Semua hasil yang disajikan di bab ini berasal dari eksekusi program Python di Google Colab. Program tersebut dirancang untuk melakukan beberapa tugas sekaligus: membangkitkan data simulasi IBM HR Analytics, menjalankan *preprocessing*, melakukan seleksi fitur dengan PCC, melatih model LSTM, lalu menghasilkan berbagai visualisasi dan tabel hasil. Agar rapi dan mudah dilacak, seluruh data



dan *output* disimpan secara terstruktur ke dalam folder Google Drive.

4.1 Kinerja Model LSTM

Tabel 5 menyajikan hasil evaluasi model Long Short-Term Memory (LSTM) pada data pengujian.

Tabel 5. Hasil Evaluasi Model LSTM

Metrik	Nilai
Accuracy	71,86%
Precision	69,47%
Recall	55,00%
F1-Score	61,40%

Dari Tabel 5, terlihat bahwa model LSTM yang dibangun mencapai akurasi 71,86%. Angka ini berarti cukup baik—artinya, model cukup mampu mengklasifikasikan data turnover pada tahap pengujian. Lalu, nilai precision 69,47% menunjukkan bahwa prediksi positif yang dihasilkan model sebagian besar tepat. Dengan kata lain, ketika model memprediksi seorang karyawan akan turnover, biasanya prediksi itu benar. Namun, ada satu catatan penting. Recall-nya hanya 55,00%. Ini artinya model ini masih cukup sering kelewatan dalam mendeteksi kasus turnover yang sebenarnya terjadi. Akibatnya, tidak semua karyawan yang benar-benar keluar berhasil teridentifikasi.

Adapun F1-score 61,40%—yang merupakan rata-rata harmonik dari precision dan recall—memberi gambaran tentang keseimbangan antara keduanya. Jadi, model LSTM ini cukup efektif untuk memprediksi turnover, tapi masih perlu peningkatan, terutama di sisi recall, agar deteksi terhadap karyawan yang benar-benar turnover bisa lebih optimal ke depannya.

Tabel 6 menyajikan classification report model LSTM yang memberikan rincian performa untuk setiap kelas prediksi.

Tabel 6. Classification Report Model LSTM

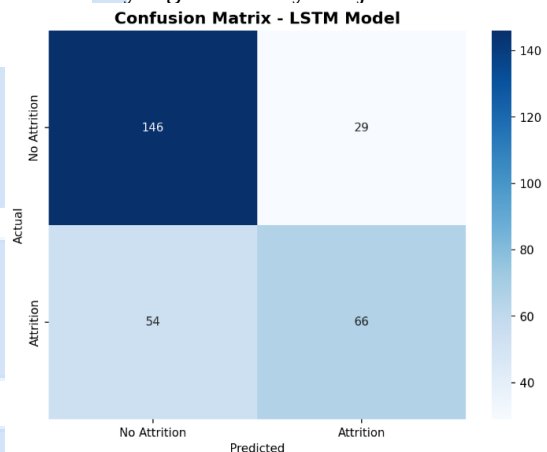
Kelas	Precision	Recall	F1-Score	Support
No Attrition	0,73	0,83	0,78	175
Attrition	0,69	0,55	0,61	120
Akurasi			0,72	295
Macro Avg	0,71	0,69	0,70	295
Weighted Avg	0,72	0,72	0,71	295

Sebagai ringkasan, model LSTM ini mencapai akurasi 71,86%, presisi 69,47%, recall 55,00%, dan F1-score 61,40%. Kalau dibandingkan dengan penelitian lain, angka ini tergolong sebanding—meskipun memang ada yang lebih tinggi. Misalnya, Morteza pour Shiri dkk.

(2025) berhasil meraih akurasi 89,65% pada dataset IBM yang sama menggunakan Bi-TCN. Demikian pula Wardhani & Lhaksana (2022) mendapatkan akurasi 86,5% dengan regresi logistik.

Dari *classification report* (Tabel 6), terlihat bahwa model ini ternyata lebih unggul dalam memprediksi karyawan yang *tidak* mengalami *attrition* (kelas *No Attrition*). Untuk kelas ini, *precision*-nya 0,73, *recall* 0,83, dan F1-score 0,78. Sementara untuk prediksi *attrition*, angkanya lebih rendah: *precision* 0,69, *recall* 0,55, F1-score 0,61.

Model ini cenderung konservatif. Maksudnya, lebih hati-hati agar tidak salah mengklasifikasikan karyawan yang sebenarnya bertahan sebagai karyawan yang berisiko keluar. Ini memang pilihan yang wajar, tapi konsekuensinya model jadi cukup banyak melewatkan kasus *turnover* yang sebenarnya terjadi.



Gambar 5. Confusion matrix hasil prediksi model LSTM

Pada Gambar 5, ditampilkan *confusion matrix* dari hasil prediksi model LSTM. Matriks ini memperlihatkan bagaimana distribusi prediksi—baik yang benar maupun yang salah—pada data uji yang berjumlah 295 sampel. Sel kiri atas berisi *True Negatives*, yaitu karyawan yang memang *tidak attrition* dan diprediksi benar oleh model. Sel kanan bawah adalah *True Positives*, karyawan yang memang *attrition* dan diprediksi dengan benar. Sementara itu, sel kanan atas menunjukkan *False Positives*—karyawan yang sebenarnya *tidak attrition*, tapi salah diprediksi sebagai *attrition*. Terakhir, sel kiri bawah adalah *False Negatives*, yaitu karyawan yang sebenarnya *attrition*, tetapi tidak terdeteksi oleh model.

Dengan visualisasi seperti ini, jadi lebih mudah melihat letak kesalahan klasifikasi yang dilakukan model.

4.2 Kinerja Model per Departemen

Tabel 7 menyajikan perbandingan kinerja model LSTM pada setiap departemen.

Tabel 7. Perbandingan Kinerja Model per Departemen

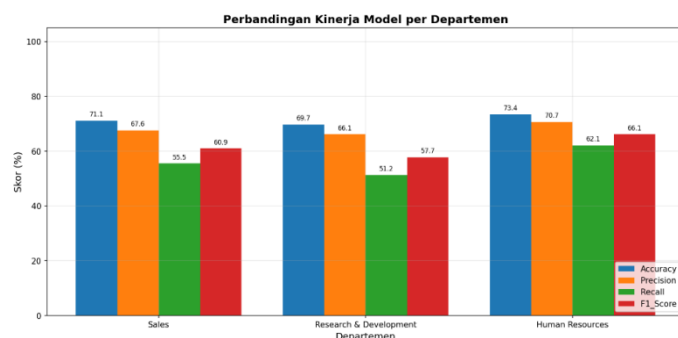
Departemen	Jumlah Sampel	Attrition Rate	Akurasi	Presisi	Recall	F1-Score
Human Resources	158	41,8%	73,42%	70,69%	62,12%	66,13%
Sales	519	40,7%	71,10%	67,63%	55,45%	60,94%
Research & Development	793	40,4%	69,74%	66,13%	51,25%	57,75%

Dari hasil pengujian di Tabel 7, terlihat bahwa departemen Human Resources (HR) mencatatkan performa terbaik. Akurasinya mencapai 73,42% dengan F1-Score 66,13%. Ini menandakan bahwa model ini lebih mampu memprediksi *turnover* di departemen HR dibandingkan dua departemen lainnya.

Di posisi berikutnya, Sales mencatat akurasi 71,10% (F1-Score 60,94%). Sementara itu, R&D justru menjadi yang terendah meskipun memiliki jumlah sampel paling besar—akurasinya hanya 69,74% dengan F1-Score 57,75%.

Menariknya, nilai *recall* tertinggi juga diraih oleh HR, yaitu 62,12%. Artinya, model paling peka dalam mendeteksi potensi *attrition* di departemen ini dibandingkan di Sales dan R&D.

Jadi, perbedaan performa ini menunjukkan bahwa karakteristik data dan pola *turnover* yang unik di tiap departemen memang ikut mempengaruhi kemampuan prediksi model. Jadi, pendekatan yang sama persis untuk semua unit/departemen mungkin kurang optimal.



Gambar 6. Perbandingan Kinerja Model

Pada Gambar 6, disajikan grouped bar chart yang membandingkan kinerja model antar departemen: HR, Sales, dan R&D. Sumbu vertikalnya dibuat dari 0 hingga

105% agar proporsinya terlihat jelas. Keempat metrik—Accuracy, Precision, Recall, dan F1-Score—dibedakan dengan warna berbeda, dan di atas setiap batang dicantumkan nilai numeriknya.

Secara visual, terlihat jelas bahwa HR unggul di semua metrik, disusul Sales, dan R&D di posisi terakhir. HR dengan 158 karyawan dan attrition rate 41,8% mencapai akurasi tertinggi 73,42% (F1-score 66,13%). Sales dengan 519 karyawan dan attrition rate 40,7% mencapai akurasi 71,10% (F1-score 60,94%). Sementara R&D dengan jumlah karyawan terbesar yaitu 793 dan attrition rate 40,4% justru paling rendah: akurasi 69,74% (F1-score 57,75%).

Lalu, HR bisa paling baik. Ada dua kemungkinan. Pertama, data HR mungkin lebih homogen meskipun jumlah sampelnya lebih kecil. Kedua, faktor pemicu turnover di HR lebih mudah diidentifikasi—terutama korelasi *OverTime* yang sangat tinggi (0,4732) seperti terlihat di Tabel 3.

Temuan ini sejalan dengan apa yang dikemukakan Romadloni & Pardede (2019), bahwa seleksi fitur berbasis PCC mampu meningkatkan kinerja klasifikasi secara signifikan. Jadi wajar kalau departemen dengan pola korelasi yang lebih jelas cenderung menghasilkan akurasi yang lebih tinggi.

4.3 Top Fitur Paling Berpengaruh

Tabel 8 menyajikan lima fitur teratas yang paling berpengaruh terhadap turnover karyawan berdasarkan hasil analisis Pearson Correlation Coefficient (PCC).

Tabel 8. Top 5 Fitur Paling Berpengaruh (PCC)

Peringkat	Fitur	Korelasi	Interpretasi
1	OverTime	+0,3366	Semakin sering lembur, semakin tinggi risiko turnover
2	JobSatisfaction	-0,2038	Semakin tinggi kepuasan kerja, semakin rendah risiko turnover
3	WorkLifeBalance	-0,1407	Semakin baik keseimbangan kerja, semakin rendah risiko turnover
4	MonthlyIncome	-0,0912	Semakin tinggi pendapatan, semakin rendah risiko turnover
5	EnvironmentSatisfaction	-0,0877	Semakin baik lingkungan kerja, semakin rendah risiko turnover

Dari Tabel 8, *OverTime* menempati posisi pertama dengan korelasi positif tertinggi terhadap turnover (+0,3366). Makin sering lembur, makin tinggi risiko keluar—karena kelelahan, stres, dan terganggunya keseimbangan kerja-kehidupan.

Sebaliknya, JobSatisfaction memiliki korelasi negatif terbesar (-0,2038): kepuasan kerja tinggi justru menurunkan kemungkinan keluar. WorkLifeBalance juga berkorelasi negatif, menegaskan pentingnya keseimbangan tersebut. MonthlyIncome dan EnvironmentSatisfaction turut berpengaruh negatif meski lebih rendah, artinya pendapatan memadai dan lingkungan nyaman bisa meningkatkan loyalitas.

Dengan demikian, OverTime adalah faktor paling dominan, diikuti JobSatisfaction dan WorkLifeBalance. Temuan ini sejalan dengan Morteza pour Shiri dkk. (2025) bahwa prediksi turnover memungkinkan perusahaan bertindak proaktif, karena penyebab attrition meliputi alasan pribadi, ketidakpuasan kerja, kompensasi tidak memadai, hingga kondisi kerja yang kurang menguntungkan.

4.4 Interpretasi dan Implikasi Manajerial

Berdasarkan hasil analisis yang telah dilakukan, rekomendasi strategi retensi yang berbeda dapat diberikan untuk setiap departemen. Tabel 9 menyajikan ringkasan rekomendasi strategi retensi karyawan per departemen.

Tabel 9. Rekomendasi Strategi Retensi Karyawan per Departemen

Departemen	Prioritas Utama	Tindakan yang Direkomendasikan
Human Resources	Mengurangi lembur & evaluasi kompensasi	• Kurangi frekuensi lembur • Tingkatkan kepuasan kerja • Evaluasi struktur kompensasi
Sales	Manajemen lembur & kepuasan kerja	• Manajemen lembur yang lebih baik • Program peningkatan kepuasan kerja • Perbaikan lingkungan kerja
Research & Development	Keseimbangan kerja-kehidupan	• Keseimbangan kerja-kehidupan • Manajemen lembur • Program kepuasan kerja

Tabel 9 menunjukkan bahwa setiap departemen membutuhkan strategi retensi yang berbeda. Di HR, prioritasnya mengurangi lembur, evaluasi kompensasi, dan meningkatkan kepuasan kerja. Di Sales, selain manajemen lembur dan kepuasan kerja, perlu juga perbaikan jadwal, motivasi, dan lingkungan kerja—wajar karena tekanan target cukup tinggi. Sementara di R&D, yang paling utama adalah menjaga *work-life*

balance, diikuti pengelolaan lembur yang efektif dan program peningkatan kepuasan kerja.

Hal ini sejalan dengan Mahammad dkk. (2025) bahwa organisasi harus menyelidiki penyebab *turnover* terlebih dahulu agar intervensi tepat sasaran. Yakymiv & Yehorchenkova (2026) juga menambahkan bahwa model prediksi *turnover* bernilai praktis karena bisa diintegrasikan ke sistem analitik SDM dan manajemen proyek untuk mendukung pengambilan keputusan manajerial.

5. PENUTUP

5.1 Kesimpulan

Berdasarkan hasil penerapan dan pembahasan yang telah dilakukan, dapat disimpulkan beberapa hal sebagai berikut:

a. Dari hasil proses pengolahan data, diperoleh:

- 1) Model LSTM dengan seleksi fitur PCC berhasil diterapkan untuk prediksi turnover karyawan dengan akurasi 71,86%, presisi 69,47%, recall 55,00%, dan F1-score 61,40%. Model LSTM yang dikombinasikan dengan seleksi fitur PCC berhasil diimplementasikan untuk memprediksi turnover karyawan. Dengan tingkat akurasi 71,86% (cukup baik), model ini menunjukkan kinerja yang lebih baik dalam memprediksi karyawan yang tidak keluar (No Attrition) dibandingkan memprediksi yang keluar (Attrition). Karakter ini dianggap lebih aman untuk praktik manajemen karena meminimalisir risiko melakukan intervensi terhadap karyawan yang sebenarnya loyal.
- 2) Seleksi fitur PCC menunjukkan bahwa OverTime (0,3366) sebagai faktor pemicu turnover paling dominan secara keseluruhan, diikuti JobSatisfaction (-0,2038) dan WorkLifeBalance (-0,1407). Analisis per departemen menunjukkan bahwa Human Resources memiliki korelasi OverTime tertinggi (0,4732), sementara Sales dan R&D memiliki pola korelasi yang berbeda.
- 3) Kinerja model bervariasi antar departemen: Human Resources (akurasi 73,42%), Sales (71,10%), dan R&D (69,74%). Hal ini menunjukkan bahwa pendekatan yang berbeda diperlukan untuk setiap departemen.
- 4) Direkomendasikan strategi retensi pada setiap departemen, yaitu Human Resources fokus pada pengurangan lembur dan evaluasi kompensasi; Sales fokus pada manajemen lembur dan peningkatan kepuasan kerja; sedangkan R&D fokus pada keseimbangan kerja-kehidupan.

- b. Google Colab dengan Python sebagai platform dapat digunakan untuk implementasi deep learning dalam analisis SDM.

5.2 Saran

Berdasarkan temuan-temuan dalam penelitian ini, beberapa saran untuk pengembangan selanjutnya adalah sebagai berikut:

- Pengembangan Model Time Series dengan cara menggunakan data longitudinal (time series) untuk menangkap pola temporal perubahan perilaku karyawan sebelum terjadinya attrition.
- Penambahan data eksternal dengan melakukan integrasi data eksternal seperti tingkat pengangguran, rata-rata gaji industri, atau data ekonomi makro lainnya untuk memperkaya fitur dan meningkatkan akurasi prediksi.
- Membandingkan metode LSTM dengan metode-metode lain seperti XGBoost, Random Forest, dan GRU.
- Mengembangkan dashboard secara real-time untuk memantau risiko turnover secara berkala.
- Pengembangan Model Spesifik Departemen dengancara melatih model secara terpisah untuk setiap departemen guna meningkatkan akurasi prediksi.
- Pengembangan Aplikasi Berbasis Web berbasis python dan script lainnya sehingga dapat diakses secara real-time oleh para praktisi HRD tanpa harus menjalankan Google Colab secara manual.

Proceedings of the NEMISA Digital Skills Conference 2023: Scaling Data Skills for Multidisciplinary Impact, EPiC Series in Education Science, 5, 17-30. EasyChair. <https://easychair.org/publications/paper/q5sh>

Pearson, K. (1895). Notes on regression and inheritance in the case of two parents. *Proceedings of the Royal Society of London*, 58, 240-242.

Ponmalar, S., Affrin Fowmiya, N., & Nandhini, C. (2024). AI-driven retention: A hybrid approach to employee turnover prediction. Dalam 2024 9th International Conference on Communication and Electronics Systems (ICCES) (hlm. 1554-1559). IEEE.

<https://doi.org/10.1109/ICCES63552.2024.10859722>

Romadloni, N. T., & Pardede, H. F. (2019). Seleksi Fitur Berbasis Pearson Correlation Untuk Optimasi Opinion Mining Review Pelanggan. *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 3(3), 505–510.

Wardhani, F. H., & Lhaksmana, K. M. (2022). Predicting Employee Attrition Using Logistic Regression with Feature Selection. *Sinkron: Jurnal dan Penelitian Teknik Informatika*, 6(4).

Yakymiv, A., & Yehorchenkova, N. (2026). Forecasting Employee Productivity Using LSTM Networks: Synthetic Data Approach and Baseline Comparison. *Economy and Society*, 84-67.

6. DAFTAR PUSTAKA

Ganapathisamy, S., & Narayan, V. (2024). A Long Short-Term Memory with Recurrent Neural Network and Brownian Motion Butterfly Optimization for Employee Attrition Prediction. *International Journal of Intelligent Engineering and Systems*, 17(1), 183-192.

Hochreiter, S., & Schmidhuber, J. (1997). Long short-term memory. *Neural Computation*, 9(8), 1735–1780.

Mahammad, D. R., Priya, S. L., Thulasimani, T., Mann, S., Mann, S., & Kalra, G. (2025). Analysis and prediction of employee turnover characteristics based on LSTM-RNN-HMM model. *Journal of Intelligent & Fuzzy Systems*.

Mortezapour Shiri, F., Yamaguchi, S., & Ahmadon, M. A. B. (2025). A Deep Learning Model Based on Bidirectional Temporal Convolutional Network (Bi-TCN) for Predicting Employee Attrition. *Applied Sciences*, 15(6), 2984.

Musanga, V., & Chibaya, C. (2023). A predictive model to forecast employee churn for HR analytics. Dalam