

BAB II

TINJAUAN PUSTAKA

2.1 Landasan Teori

2.1.1 *Turnover* Karyawan

Seorang karyawan meninggalkan organisasi secara sukarela atau tidak sukarela disebut *turnover* karyawan (Mihajlov & Mihajlov, 2020). *Turnover* menjadi salah satu indikator penting dalam manajemen sumber daya manusia karena dapat memengaruhi operasional, produktivitas, dan stabilitas suatu organisasi. Dalam praktiknya, *turnover* dapat terjadi karena berbagai alasan, termasuk kepuasan karyawan yang rendah di tempat kerja, gaji yang tidak kompetitif, atau kurangnya peluang untuk berkembang. *Turnover* sukarela biasanya terjadi karena keputusan karyawan sendiri, seperti pindah ke perusahaan lain, melanjutkan pendidikan, atau alasan organisasi lainnya. Sementara *turnover* tidak sukarela terjadi karena keputusan organisasi, seperti pemecatan karyawan karena kinerja yang buruk. Tingkat *turnover* yang tinggi, terutama jika melibatkan karyawan berkinerja tinggi, dapat menimbulkan kerugian besar bagi perusahaan karena harus mengalokasikan sumber daya tambahan untuk proses rekrutmen dan pelatihan karyawan baru. Oleh karena itu, memahami dan memprediksi *turnover* karyawan menjadi penting agar perusahaan dapat merumuskan strategi yang tepat untuk mempertahankan karyawan yang berpotensi.

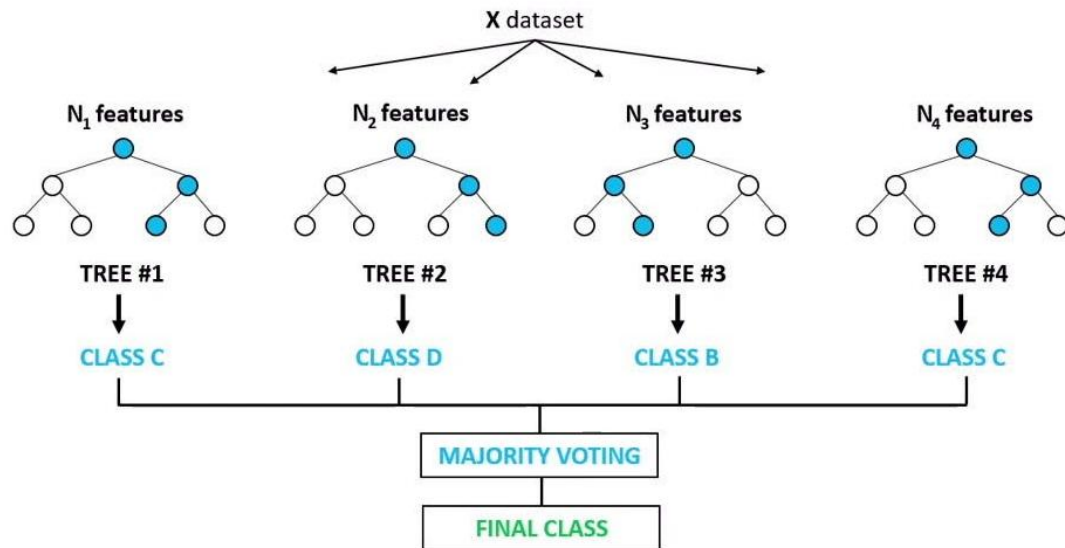
2.1.2 Machine Learning

Machine learning adalah cabang dari kecerdasan buatan yang memungkinkan komputer untuk belajar dari data dan menghasilkan prediksi atau ramalan tanpa pemrograman eksplisit (Vyawahare, 2022). Dalam konteks prediksi pergantian karyawan, *machine learning* digunakan untuk menganalisis data historis untuk mengidentifikasi faktor-faktor spesifik yang terkait dengan retensi karyawan atau pertumbuhan bisnis. Salah satu jenis *machine learning* yang biasanya digunakan untuk kasus seperti ini adalah *supervised learning*, dimana algoritma dilatih menggunakan data berlabel. Misalnya, informasi mengenai karyawan yang bekerja di perusahaan, termasuk karakteristik mereka, seperti gaji, tingkat kepuasan kerja, jam kerja, dan sebagainya. Dengan model yang dikembangkan, sistem dapat memprediksi kemungkinan seorang pekerja akan meninggalkan perusahaan di masa depan. Pendekatan ini tidak hanya membantu dalam meminimalkan tingkat *turnover*, tetapi juga memberikan wawasan strategis bagi perusahaan untuk meningkatkan retensi karyawan.

2.1.3 Random Forest

Random Forest adalah algoritma berbasis pohon keputusan yang membangun banyak model pada subset data berbeda dan menggabungkan hasilnya melalui voting atau rata-rata (Mahdi Abdulkareem & Mohsin Abdulazeez, 2021). Metode ini efektif mengurangi *overfitting*, cocok untuk data besar dan kompleks, serta mampu mengidentifikasi fitur-fitur penting. Dalam prediksi *turnover* karyawan, *Random Forest* memberikan akurasi tinggi dan membantu memahami

faktor-faktor yang memengaruhi keputusan karyawan. Berikut ilustrasi dari algoritma metode *Random Forest* pada Gambar 2. 1.

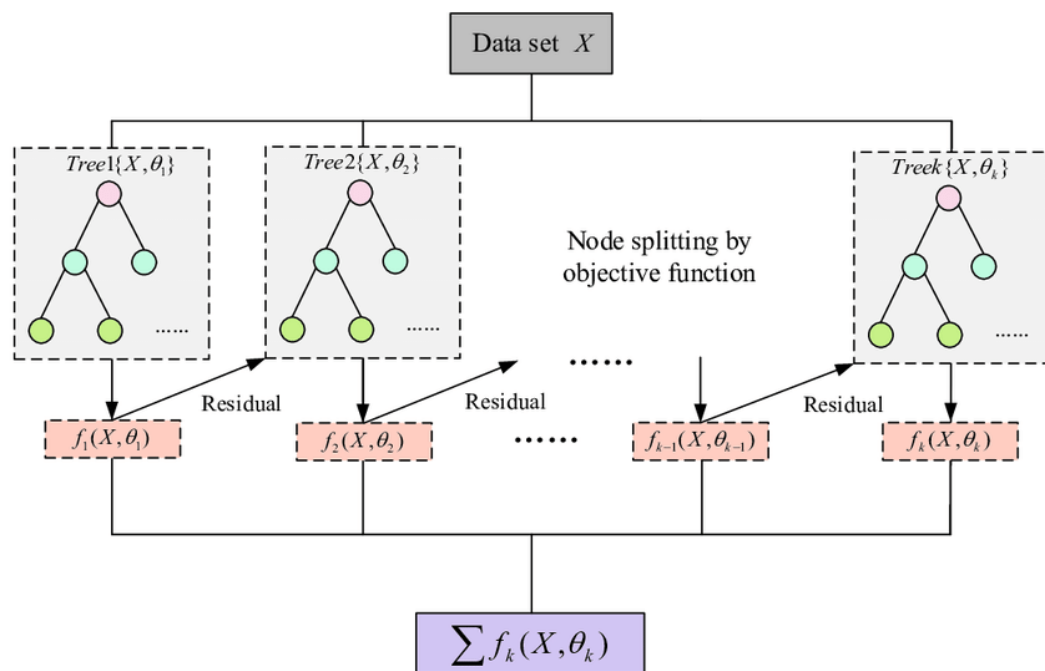


Gambar 2. 1 Ilustrasi *Random Forest*

2.1.4 *XGBoost*

Extreme Gradient Boosting atau *XGBoost*, adalah algoritma *machine learning* yang didasarkan pada *boosting* yang dirancang untuk meningkatkan kinerja prediksi dengan mengubah banyak model yang lemah menjadi model yang kuat (Hanif, 2020). Berbeda dengan *Random Forest* yang membangun model secara paralel, *XGBoost* membangun model dengan cara yang logis, di mana setiap model yang baru dikembangkan akan meningkatkan hasil dari model sebelumnya. Teknik ini sangat efektif untuk meningkatkan akurasi prediksi, terutama untuk data yang kompleks dan nonlinier. Selain itu, *XGBoost* didasarkan pada teknik regularisasi yang dapat mengurangi risiko *overfitting* dan eksekusi paralel untuk mempercepat proses pembelajaran. *XGBoost* telah banyak mengembangkan kemampuan *machine*

learning dan sangat cocok digunakan untuk memprediksi *turnover* karyawan karena kemampuannya dalam menganalisis data yang tidak konsisten dan memberikan hasil prediksi yang stabil. Berikut merupakan ilustrasi dari algoritma *XGBoost* pada Gambar 2. 2.

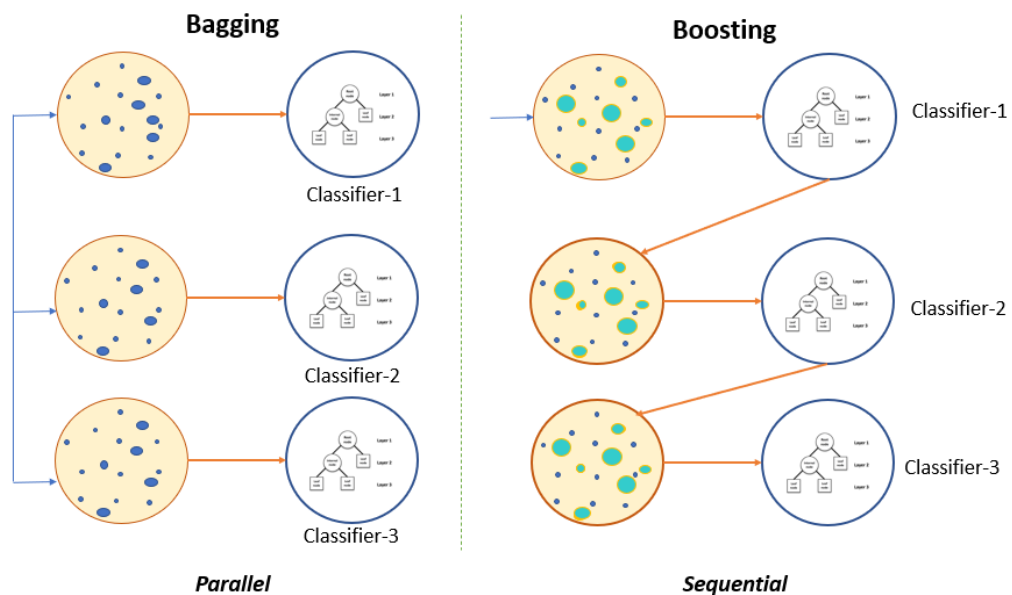


Gambar 2. 2 Ilustrasi *XGBoost*

2.1.5 Ensemble Learning

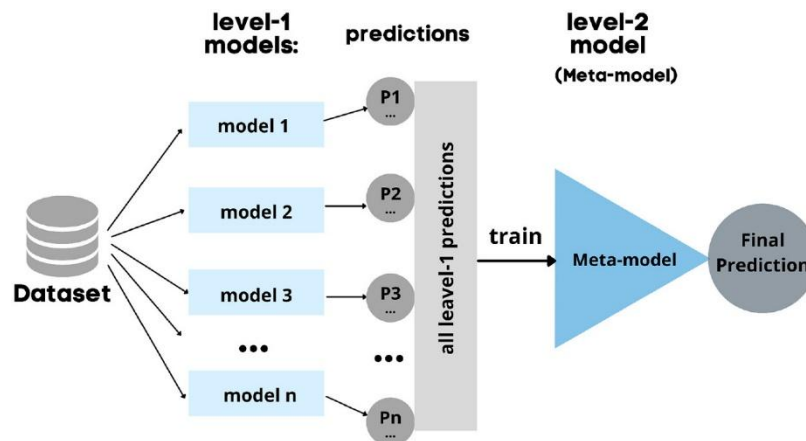
Ensemble learning adalah teknik *machine learning* yang menggabungkan beberapa model prediksi untuk menghasilkan hasil yang lebih akurat dan konsisten daripada menggunakan model tunggal (Dang dkk., 2022). Prinsip dasar pembelajaran *ensemble* adalah bahwa hubungan antara beberapa model, yang masing-masing memiliki kekuatan dan kelemahan yang unik, secara bertahap dapat meningkat dan menghasilkan prediksi yang lebih baik. Dalam praktiknya, beberapa

teknik pembelajaran *ensemble* umumnya digunakan, seperti *bagging*, *boosting*, dan *stacking*. *Bagging*, melatih beberapa model sejenis secara paralel pada data acak untuk mengurangi variansi, contohnya *Random Forest*. *Boosting*, seperti di *XGBoost*, membangun model secara langsung dengan meningkatkan kinerja model sebelumnya. Berikut merupakan alur kerja dari metode *Bagging* dan *Boosting* pada gambar 2. 3.



Gambar 2. 3 Ilustrasi *Bagging* dan *Boosting*

Di sisi lain, *stacking* adalah menggabungkan berbagai model berbeda dan memanfaatkan model meta untuk menyatukan prediksi, guna menangkap pola yang lebih kompleks.

Gambar 2. 4 Ilustrasi *Stacking*

2.1.6 Ensemble *Stacking*

Stacking adalah teknik *ensemble* yang menggabungkan dua jenis model: model dasar (*base learners*) dan model meta (*meta-learner*) (Wu dkk., 2024). Model dasar bertugas digunakan untuk membuat prediksi awal berdasarkan data pembelajaran, dan hasil prediksi digunakan sebagai input untuk model meta, yang kemudian menghasilkan prediksi akhir. Karena memungkinkan kombinasi berbagai algoritma dengan karakteristik yang berbeda, konsep *stacking* sangat fleksibel. Dalam konteks prediksi *turnover* karyawan, *stacking* menawarkan potensi peningkatan akurasi karena dapat menggabungkan kekuatan algoritma *Random Forest* dan *XGBoost*. *Random Forest*, yang sangat baik dalam menganalisis variabel yang kompleks, dan *XGBoost*, yang sangat baik dalam menganalisis masalah dengan cara yang sistematis, keduanya dapat menjadi sangat efektif. Hasilnya, model *stacking* yang menggabungkan keduanya berpotensi menghasilkan hasil prediksi yang lebih akurat. Model meta yang umum digunakan dalam *stacking*

adalah regresi logistik, tetapi dapat pula menggunakan algoritma lain tergantung pada kebutuhan dan kompleksitas data.

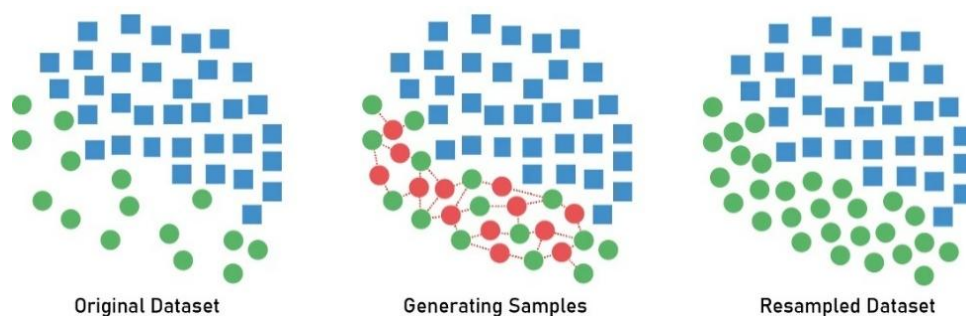
2.1.7 Multi Layer Perceptron (MLP)

Multi Layer Perceptron (MLP) merupakan jenis jaringan saraf tiruan berlapis yang terdiri atas input *layer*, satu atau lebih *hidden layer*, serta *output layer*, di mana setiap neuron saling terhubung secara penuh dan dipelajari menggunakan algoritma *backpropagation*. Arsitektur ini mampu memodelkan hubungan non-linear secara efektif melalui fungsi aktivasi seperti *relu* dan *tanh* (Intharasompong dkk., 2025). Berbagai penelitian menunjukkan bahwa MLP memiliki fleksibilitas tinggi dalam mempelajari pola kompleks, sehingga kinerjanya dapat melampaui algoritma tradisional pada masalah klasifikasi (Shafieian & Zulkernine, 2023). Selain itu, MLP juga terbukti efektif digunakan sebagai *meta-learner* dalam model *ensemble*, karena mampu menangkap interaksi non-linear dari keluaran *base learner* dan menghasilkan prediksi yang lebih stabil dan akurat dibandingkan *meta-learner* linear seperti *Logistic Regression* (Sumsion & Lee, 2025).

2.1.8 Imbalance Data dan SMOTE

Imbalance data adalah kondisi ketika proporsi antara kelas mayoritas dan minoritas tidak seimbang (Nemoto dkk., 2025), seperti pada kasus prediksi *turnover* karyawan di mana jumlah karyawan bertahan jauh lebih banyak daripada yang keluar. Ketimpangan ini dapat menyebabkan model bias terhadap kelas mayoritas dan gagal mengenali pola pada kelas minoritas. Untuk mengatasi hal ini, digunakan teknik SMOTE (*Synthetic Minority Oversampling Technique*) yang bekerja dengan

merekonstruksi sampel sintetis dari kelas minoritas berdasarkan *nearest neighbors* pada data pelatihan (Zhang & Li, 2022). SMOTE terbukti efektif dalam menyeimbangkan distribusi kelas, meningkatkan sensitivitas model terhadap kelas minoritas, serta mencegah *overfitting*, sehingga model seperti *Random Forest*, *XGBoost*, maupun *ensemble stacking* dapat menghasilkan prediksi yang lebih akurat dan seimbang (Syukron dkk., 2020).



Gambar 2. 5 Ilustrasi Cara Kerja SMOTE

2.1.9 Bayesian Optimization

Bayesian Optimization adalah metode optimasi berbasis probabilitas yang digunakan untuk mencari nilai optimal dari fungsi objektif yang kompleks dan sulit dievaluasi secara langsung (Daulton dkk., 2022). Tidak seperti *grid search* atau *random search* yang melakukan pencarian secara menyeluruh atau acak, *Bayesian Optimization* membangun model probabilitas (sering menggunakan *Gaussian Process*) dan memanfaatkan fungsi akuisisi untuk menentukan titik parameter yang paling potensial dievaluasi, sehingga proses pencarian menjadi lebih efisien.

Dalam *machine learning*, *Bayesian Optimization* banyak diterapkan untuk *hyperparameter tuning* karena mampu menangani ruang parameter yang besar dan multidimensi dengan jumlah percobaan yang relatif sedikit (Onorato, 2024).

Metode ini efektif meningkatkan performa algoritma seperti *Random Forest* dan *XGBoost*, serta dalam penelitian ini digunakan untuk mengoptimalkan model *ensemble stacking* agar menghasilkan prediksi turnover karyawan yang lebih akurat.

2.1.10 Evaluation Metrics

Dalam konteks klasifikasi, evaluasi kinerja model dilakukan dengan menggunakan metrik seperti *accuracy*, *precision*, *recall*, *F1-score* dan ROC-AUC yang masing-masing memiliki peran penting dalam menilai aspek spesifik dari performa model.

2.1.7.1 Accuracy

Metrik *accuracy* merupakan ukuran yang digunakan untuk menentukan sejauh mana model klasifikasi dapat memberikan prediksi yang tepat secara keseluruhan. Semakin tinggi nilai *accuracy*, semakin tinggi pula tingkat keberhasilan model dalam mengenali label kelas dengan benar.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

2.1.7.2 Precision

Precision adalah rasio yang mengukur ketepatan model dalam memprediksi kelas positif, yaitu seberapa besar proporsi prediksi positif yang benar-benar sesuai dengan kelas sebenarnya.

$$Precision = \frac{TP}{TP + FP}$$

2.1.7.3 Recall

Recall, atau dikenal pula sebagai sensitivitas, digunakan untuk mengukur kemampuan model dalam menangkap semua *instance* yang benar-benar termasuk ke dalam kelas positif.

$$Recall = \frac{TP}{TP + FN}$$

2.1.7.4 F1-Score

F1-Score merupakan metrik gabungan yang memperhitungkan keseimbangan antara *precision* dan *recall*. Nilai F1 tinggi menunjukkan bahwa model tidak hanya akurat dalam prediksi positif, tetapi juga mampu mendeteksi sebagian besar *instance* positif.

$$F1\ Score = \frac{Precision \times Recall}{Precision + Recall}$$

Keterangan:

TP = True Positive,

TN = True Negative,

FP = False Positive,

FN = False Negative.

2.1.7.5 ROC - AUC

Receiver Operating Characteristic (ROC) adalah kurva yang menggambarkan hubungan antara *True Positive Rate* (TPR atau *recall*) dan *False*

Positive Rate (FPR) pada berbagai ambang batas klasifikasi (Hillman & Hocking, 2023). Kurva ini menunjukkan kemampuan model dalam membedakan kelas positif dan negatif pada berbagai tingkat sensitivitas. Semakin dekat kurva ROC dengan sisi kiri atas, semakin baik performa model karena berarti memiliki TPR yang tinggi dan FPR yang rendah.

Area Under Curve (AUC) merupakan nilai numerik yang diperoleh dari luas area di bawah kurva ROC. Nilai AUC berada pada rentang 0 hingga 1, dengan nilai mendekati 1 menunjukkan kemampuan diskriminasi yang sangat baik, sedangkan nilai 0,5 menunjukkan performa model setara dengan tebakan acak (Li, 2024). Dalam konteks klasifikasi dengan data tidak seimbang, AUC menjadi metrik penting karena tidak hanya menilai akurasi keseluruhan, tetapi juga memperhatikan kemampuan model dalam membedakan antara karyawan yang bertahan dan yang meninggalkan perusahaan.

2.2 Penelitian Terkait

Saat ini, penerapan metode *Machine Learning* dalam bidang prediksi *turnover* karyawan semakin berkembang pesat, seiring dengan meningkatnya kebutuhan perusahaan dalam mempertahankan talenta terbaik. Berbagai algoritma telah digunakan untuk mengatasi permasalahan ini, seperti *Random Forest*, *XGBoost*, *Support Vector Machine (SVM)*, serta pendekatan *ensemble learning* seperti *bagging*, *boosting*, dan *stacking*. Dalam konteks prediksi *turnover*, *Random Forest* dan *XGBoost* menjadi dua algoritma yang sering digunakan karena

kemampuannya dalam menangani data yang kompleks dan menghasilkan performa prediksi yang tinggi.

Penelitian yang dilakukan oleh (Khasanah, 2024) dan (Jamroni dkk., 2025) menunjukkan bahwa *Random Forest* memiliki akurasi yang tinggi dalam memprediksi *turnover* karyawan, bahkan mampu mengungguli *SVM* dalam beberapa skenario. Selain itu, penelitian oleh (Maahiroh, 2024) membuktikan bahwa *XGBoost* juga memberikan hasil prediksi yang sangat baik, terlebih setelah diterapkannya teknik SMOTE untuk mengatasi ketidakseimbangan data. Sementara itu, studi yang dilakukan oleh (Biswas dkk., 2023) dan (Hossen dkk., 2021) menunjukkan bahwa penggunaan metode *ensemble*, khususnya *stacking* dan *feature-based ensemble* yang menggabungkan berbagai algoritma seperti *Gradient Boosting*, *Random Forest*, dan *K-Nearest Neighbour*, mampu meningkatkan akurasi serta generalisasi model dalam memprediksi perilaku *turnover* karyawan.

Selain itu, penelitian (Sudarto & Kusri, 2024) menunjukkan bahwa metode *Stacking Ensemble Learning* mampu menghasilkan performa prediksi *turnover* yang lebih baik dibandingkan *Random Forest* maupun model individual lainnya. Temuan ini sejalan dengan penelitian (Rayadin dkk., 2024) yang membandingkan teknik *bagging*, *boosting*, dan *stacking*, di mana pendekatan *stacking* memberikan hasil yang lebih unggul. Penelitian (Intharasompong dkk., 2025) juga memperkuat temuan tersebut dengan menunjukkan bahwa model *stacking* berbasis MLP mampu melampaui model dasar yang digunakan. Secara keseluruhan, hasil-hasil ini menegaskan bahwa *stacking* menawarkan keunggulan lebih konsisten dibandingkan metode *ensemble* lainnya dalam tugas klasifikasi.

Dalam upaya meningkatkan performa model, penelitian (Khurshid dkk., 2025) menunjukkan bahwa penerapan *Bayesian Optimization* sebagai teknik *hyperparameter tuning* mampu menghasilkan model *XGBoost* yang lebih optimal dibandingkan konfigurasi dengan metode pencarian parameter lainnya. Temuan ini menegaskan bahwa *Bayesian Optimization* efektif dalam menemukan kombinasi parameter terbaik secara efisien, sehingga dapat meningkatkan kemampuan model dalam melakukan prediksi pada berbagai konteks analisis data.

Berdasarkan penelitian-penelitian tersebut, terlihat bahwa pendekatan *Stacking Ensemble Learning* yang mengombinasikan *Random Forest* dan *XGBoost* telah menunjukkan kinerja yang kompetitif dalam berbagai studi prediksi *turnover* karyawan maupun permasalahan klasifikasi lainnya. Meskipun demikian, peluang peningkatan performa masih terbuka melalui penerapan teknik optimasi, khususnya *Bayesian Optimization*, yang terbukti mampu memperoleh konfigurasi parameter model yang lebih optimal. Oleh karena itu, penelitian ini berfokus pada pengembangan model *Stacking* menggunakan *base learner Random Forest* dan *XGBoost* dengan *meta-learner MLP* yang dioptimasi menggunakan *Bayesian Optimization*, dengan tujuan menghasilkan model yang lebih akurat dan stabil dalam memprediksi *turnover* karyawan. Tabel 2.1 berikut menyajikan rangkuman penelitian terkait yang menjadi landasan dalam pengembangan model ini.

Tabel 2. 1 Penelitian Terkait

Penelitian Terkait		
No.	Penelitian	Hasil
1.	Judul: Prediksi <i>Turnover</i> Karyawan Menggunakan Algoritma <i>Random Forest</i> Peneliti: (Khasanah, 2024) Metode: <i>Random Forest</i>	Hasil penelitian menunjukkan bahwa <i>Random Forest</i> mampu mencapai akurasi tertinggi sebesar 88% pada skenario penggunaan seluruh fitur dengan rasio data 80:20, menggunakan kombinasi <i>hyperparameter</i> terbaik ($n_estimators=100$, $max_depth=5$, $max_features=2$). Penelitian ini juga menunjukkan bahwa seleksi fitur dengan PCC tidak memberikan peningkatan signifikan jika seluruh fitur sudah digunakan, namun membantu dalam mengidentifikasi faktor-faktor penting yang memengaruhi <i>turnover</i>.
2.	Judul: Analisis Faktor dan Prediksi Atrisi untuk Optimalisasi Retensi Karyawan Menggunakan <i>Machine Learning</i> Peneliti: (Jamroni dkk., 2025)	Hasil penelitian menunjukkan bahwa algoritma <i>Random Forest</i> memberikan performa yang lebih baik dibandingkan algoritma <i>K-Nearest Neighbors</i> (KNN) dalam prediksi <i>turnover</i> karyawan. <i>Random Forest</i> mencapai tingkat

	Metode: <i>Random Forest, K-Nearest Neighbors (KNN)</i>	akurasi sebesar 93% dengan nilai AUC 0,98, serta menunjukkan nilai <i>precision</i>, <i>recall</i>, dan <i>F1-score</i> yang lebih seimbang dibandingkan KNN. Temuan ini mengindikasikan bahwa <i>Random Forest</i> memiliki kemampuan klasifikasi yang stabil dan efektif dalam mengidentifikasi karyawan yang berpotensi mengalami <i>attrition</i>.
3.	Judul: Implementasi <i>Ensemble Learning</i> Metode <i>XGBoost</i> dan <i>Random Forest</i> untuk Prediksi Waktu Penggantian Baterai Aki Peneliti: (Rayadin dkk., 2024) Metode: <i>XGBoost, Random Forest, Bagging, Stacking, Boosting</i>	Hasil penelitian menunjukkan bahwa model <i>ensemble learning</i> yang menggabungkan <i>XGBoost</i> dan <i>Random Forest</i> memberikan prediksi waktu penggantian baterai aki yang lebih akurat dibandingkan dengan model tunggal. Berdasarkan evaluasi, model <i>ensemble</i> dengan teknik <i>bagging</i> menunjukkan nilai <i>error</i> MSE terendah sebesar 145.448. Selain itu, evaluasi menggunakan metrik MAPE, MAE, dan RMSE menunjukkan bahwa model <i>boosting</i> memiliki nilai <i>error</i> terendah

		masing-masing sebesar 11,56%, 43,80, dan 38,760
4.	Judul: <i>Implementation of Machine Learning on Employee Attrition Based on Performance Parameters Using Particle Swarm Optimization and Ensemble Classifier Methods</i> Peneliti: (Fauziah dkk., 2024) Metode: <i>Particle Swarm Optimization (PSO), Support Vector Machine (SVM), Deep Learning (DL), Neural Network (NN), Bagging, Boosting</i>	Penelitian ini menunjukkan bahwa penggunaan <i>Particle Swarm Optimization</i> (PSO) untuk reduksi dimensionalitas dapat meningkatkan akurasi prediksi atrisi karyawan dengan metode <i>Support Vector Machine</i>, <i>Deep Learning</i>, dan <i>Neural Network</i>. Akurasi tertinggi dicapai oleh <i>Deep Learning</i> yang dikombinasikan dengan PSO dan dioptimasi menggunakan <i>Bagging</i>, mencapai 86,94%. Metode <i>Bagging</i> terbukti meningkatkan akurasi, sedangkan <i>Boosting</i> tidak memberikan peningkatan.
5.	Judul: <i>Unveiling diabetes onset: Optimized XGBoost with Bayesian optimization for enhanced prediction</i> Peneliti: (Khurshid dkk., 2025) Metode: <i>XGBoost, Bayesian Optimization</i>	Hasil utama dari penelitian ini menunjukkan bahwa model <i>XGBoost</i> yang dioptimalkan dengan <i>Bayesian optimization</i> mampu memberikan prediksi diabetes yang sedikit lebih akurat dibandingkan dengan <i>XGBoost</i> yang dioptimalkan dengan metode <i>grid</i>

		<i>search</i>. Secara spesifik, akurasi model yang dioptimalkan menggunakan <i>Bayesian optimization</i> mencapai 97,26%, sementara model dengan <i>grid search</i> mencapai 97,24%.
6.	Judul: Klasifikasi <i>Turnover</i> Karyawan Menggunakan Algoritma <i>XGBoost</i> (Studi kasus: Divisi <i>Engineering</i>, Perusahaan Jasa Pertambangan) Peneliti: (Maahiroh, 2024) Metode: <i>XGBoost</i>	Berdasarkan hasil penelitian, model <i>XGBoost</i> tanpa SMOTE memiliki akurasi sebesar 0,94 dengan <i>misclassification rate</i> 0,06 dan AUC 0,96. Setelah diterapkan SMOTE untuk menyeimbangkan proporsi kelas data, model <i>XGBoost</i> menunjukkan peningkatan performa dengan akurasi sebesar 0,97, <i>misclassification rate</i> menurun menjadi 0,03, dan AUC meningkat menjadi 0,98. Hal ini menunjukkan bahwa penggunaan SMOTE meningkatkan kemampuan model dalam memprediksi <i>turnover</i> karyawan, khususnya dalam menangani kelas minoritas (karyawan yang meninggalkan perusahaan), serta memberikan keseimbangan lebih baik

		antara <i>presisi</i> dan <i>recall</i> dengan <i>F1-Score</i> positif meningkat dari 0,67 menjadi 0,86.
7.	Judul: <i>Ensemble Method Based Architecture Using Random Forest Importance to Predict Employee's Turnover</i> Peneliti: (Hossen dkk., 2021) Metode: <i>Sequential Backward Selection (SBS), Chi-square, Support Vector Machine (SVM), Decision Tree (DT), Gaussian Naive Bayes (GNB), Random Forest (RF), Multi-Layer Perceptron (MLP), K-Nearest Neighbor (KNN)</i>	Penelitian ini berhasil mengidentifikasi faktor-faktor utama yang mempengaruhi pergantian karyawan seperti tingkat kepuasan kerja, evaluasi terakhir, jumlah proyek, rata-rata jam kerja bulanan, dan lama bekerja di perusahaan. Dengan menerapkan metode seleksi fitur menggunakan <i>Random Forest Importance</i> serta 10-fold <i>cross-validation</i> pada algoritma <i>Random Forest</i>, model prediksi yang dihasilkan mencapai akurasi tertinggi sebesar 0,99
8.	Judul: <i>Klasifikasi Tsunami Gempa Bumi dengan Teknik Stacking Ensemble Machine Learning</i> Peneliti: (Sudarto & Kusrini, 2024) Metode: <i>Decision Tree, Random Forest, Support Vector Machine (SVM), Naive Bayes, Stacking</i>	Hasil penelitian ini menunjukkan bahwa teknik <i>Stacking Ensemble</i> (dengan <i>Logistic Regression</i> sebagai <i>meta-learner</i>) mencapai tingkat akurasi tertinggi sebesar 94%, melampaui kinerja seluruh model individual yang diuji. Sebagai perbandingan, algoritma

		<i>Decision Tree</i> dan <i>Random Forest</i> mencatatkan akurasi terbaik di antara model dasar yaitu 90%, sedangkan <i>Support Vector Machine</i> (SVM) dan <i>Naive Bayes</i> masing-masing hanya mencapai 83% dan 78%. Temuan ini menegaskan bahwa metode <i>stacking</i> mampu mengintegrasikan keunggulan tiap algoritma untuk menghasilkan prediksi klasifikasi tsunami yang lebih presisi dan andal.
9.	Judul: <i>Customer Churn Prediction with Hybrid Resampling and Ensemble Learning</i> Peneliti: (Kimura, 2022) Metode: <i>XGBoost</i>, <i>LightGBM</i>, <i>CatBoost</i>, SMOTE-ENN, SMOTE-Tomek Links	Hasil penelitian menunjukkan bahwa algoritma <i>XGBoost</i> yang dipadukan dengan <i>hybrid resampling</i> seperti SMOTE-Tomek Links mencapai akurasi tertinggi, sekitar 83-85%. Selain itu, metode ini juga memperoleh <i>F1-score</i> tertinggi, sekitar 0.52, dan ROC-AUC mencapai sekitar 0.88, yang menunjukkan performa prediksi yang cukup baik dalam mengidentifikasi pelanggan yang akan churn. Penggunaan <i>ensemble learning</i> dan

		<i>hybrid resampling</i> secara bersamaan terbukti efektif dalam meningkatkan baik aspek akurasi maupun skor pengukuran lainnya dibandingkan dengan model tradisional maupun model tanpa <i>resampling</i>.
10.	Judul: <i>An Ensemble Learning Model for Predicting the Intention to Quit Among Employees Using Classification Algorithms</i> Peneliti: (Biswas dkk., 2023) Metode: <i>Gradient Boosting, Random Forest, K-Nearest Neighbour (KNN), Logistic Regression, Neural Network, Support Vector Machine (SVM), Naïve Bayes, Ensemble Learning</i>	Hasil penelitian ini menunjukkan bahwa model <i>ensemble</i> terbaik yang dikembangkan, yaitu kombinasi <i>Gradient Boosting, Logistic Regression</i>, dan <i>K-Nearest Neighbour</i> (KNN), mencapai tingkat akurasi sebesar 94,7% (0,947) dalam memprediksi niat keluar karyawan. Model ini juga memiliki nilai <i>F1, precision</i>, dan <i>recall</i> yang sama, yaitu 0,947, yang menunjukkan performa yang sangat baik dan konsisten dalam mengidentifikasi karyawan yang berniat keluar maupun yang tidak.
11.	Judul: <i>Stacked Ensemble Models for SME Credit Risk Assessment:</i>	Penelitian ini mengembangkan model prediksi risiko kredit UMKM menggunakan teknik <i>Stacking</i> yang

	<i>Integrating Data Balancing and Feature Selection Techniques</i> Peneliti: (Intharasompong dkk., 2025) Metode: <i>Decision Tree, SVM, Gradient Boosting, KNN, Naïve Bayes, Logistic Regression (Meta-Learning), Gradient Boosting (Meta-Learning), XGBoost (Meta-Learning), dan MLP (Meta-Learning)</i>	dipadukan dengan metode penyeimbangan data. Setelah penerapan SMOTEENN dan Stepwise Feature Selection, seluruh model <i>Stacking</i> menunjukkan kinerja lebih baik dibandingkan model dasar. META-MLP menjadi model terbaik dengan akurasi 0,942, <i>F1-score</i> 0,953, <i>recall</i> 0,990, dan AUC 0,990, menandakan kemampuannya yang unggul dalam mendeteksi nasabah berisiko tinggi pada kondisi data yang tidak seimbang. Hasil ini menunjukkan bahwa kombinasi <i>resampling</i> dan <i>ensemble learning</i> mampu meningkatkan performa prediksi secara signifikan.
--	--	---