

BAB II

LANDASAN TEORI

2.1 Landasan Teori

2.1.1 *Machine learning*

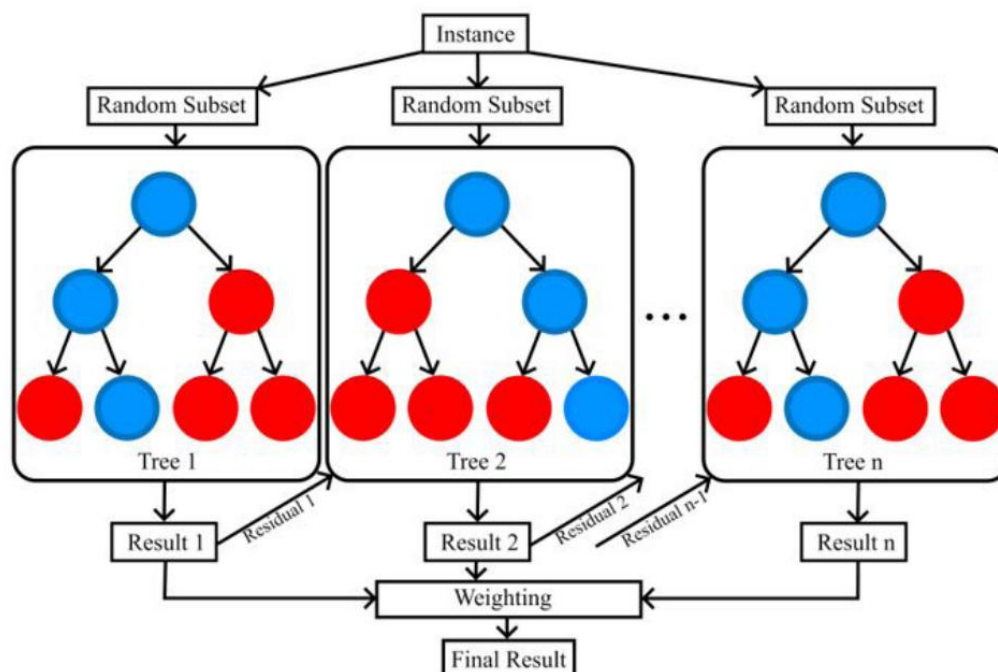
Machine learning (ML) hadir sebagai solusi efektif dalam menangani dan menganalisis kumpulan data berukuran besar secara efisien (Wahyono, 2019). Metode ini memungkinkan komputer untuk mengotomatisasi proses pengambilan keputusan, meningkatkan akurasi prediksi, serta mengoptimalkan berbagai tugas kompleks di berbagai bidang seperti kesehatan, keuangan, dan teknologi. *Machine learning* merupakan cabang khusus dari kecerdasan buatan yang berfokus pada mengajarkan mesin bagaimana belajar, sementara kecerdasan buatan secara umum bertujuan meniru kemampuan manusia.

Berbeda dengan pendekatan berbasis model dengan persamaan yang telah ditentukan sebelumnya, algoritma *Machine learning* memperoleh pemahaman langsung dari data melalui teknik komputasi (Cantt, 2023). Pendekatan ini memungkinkan mesin untuk mengidentifikasi pola, membuat prediksi, dan meningkatkan kinerjanya seiring bertambahnya data, sehingga dapat diterapkan dalam berbagai bidang seperti kesehatan, keuangan, dan teknologi. Prediksi harga tiket pesawat merupakan salah satu aplikasi *Machine learning* yang digunakan untuk memproses data historis dan mengidentifikasi pola harga. Berbagai algoritma telah digunakan untuk tugas ini, termasuk *Random Forest*, *XGBoost*, dan metode *Stacking Ensemble learning*. *Machine learning* dapat membantu dalam

mengoptimalkan strategi harga bagi maskapai penerbangan serta memberikan wawasan bagi pelanggan dalam menentukan waktu terbaik untuk membeli tiket.

2.1.1.1 Extreme Gradient Boosting

Extreme Gradient Boosting (XGBoost) membangun model secara bertahap dengan menambahkan pohon keputusan baru pada setiap tahap pelatihan, di mana pohon tersebut dirancang untuk mengoreksi kesalahan prediksi dari model sebelumnya (Oztornaci dkk., 2024). Proses ini terus berlangsung hingga jumlah pohon yang ditentukan tercapai. Setelah itu, seluruh *output* dari pohon-pohon yang telah dibentuk digabungkan, dan masing-masing pohon diberi bobot sesuai tingkat kontribusinya dalam meningkatkan akurasi prediksi akhir. *XGBoost* membangun model secara bertahap dengan menambahkan pohon keputusan baru pada setiap tahap pelatihan, di mana pohon tersebut dirancang untuk mengoreksi kesalahan prediksi dari model sebelumnya. Proses ini terus berlangsung hingga jumlah pohon yang ditentukan tercapai. Setelah itu, seluruh *output* dari pohon-pohon yang telah dibentuk digabungkan, dan masing-masing pohon diberi bobot sesuai tingkat kontribusinya dalam meningkatkan akurasi prediksi akhir. Gambar 2.1 merupakan Model *XGBoost* yang diperoleh dari penelitian (Oztornaci dkk., 2024).



Gambar 2. 1 Model *XGBoost* (Oztornaci dkk., 2024)

XGBoost adalah metode *Machine learning* yang membangun model secara bertahap dengan menambahkan pohon keputusan atau *decision tree* yang berfungsi memperbaiki kesalahan model sebelumnya (Li dkk., 2023). Struktur dasarnya adalah pohon *CART*, yang memisahkan data menjadi dua bagian secara rekursif untuk meminimalkan kesalahan.

Secara prosedural, alur pembentukan model *XGBoost* dengan parameter *default* dapat dijabarkan melalui notasi algoritma berikut. Algoritma ini menggambarkan bagaimana model membangun *ensemble* pohon keputusan secara bertahap untuk meminimalkan fungsi kerugian (*loss function*) berdasarkan parameter bawaan. Algoritma dari *XGBoost* dalam konteks regresi dapat dijelaskan melalui tahapan pada Tabel 2.1 yang diadaptasi oleh peneliti (Mienye & Sun, 2022) berikut.

Tabel 2. 1 Algoritma *XGBoost* Default

Algoritma <i>XGBoost</i> Default	
Input:	
$D = \{(x_i, y_i)\}$	: Data latih (Fitur dan Target Harga Tiket)
$L(y, F(x))$	: Fungsi Kerugian (Squared Error untuk Regresi)
Parameter Default (Inisialisasi):	
$n_estimators$ (K)	: 100 (Jumlah pohon iterasi)
$learning_rate$ (η)	: 0.3 (Langkah pembelajaran)
max_depth	: 6 (Kedalaman maksimum pohon)
reg_lambda (λ)	: 1 (Regularisasi L2)
$gamma$ (γ)	: 0 (Minimum loss reduction)
Output:	
$F(x)$	: Model Prediksi Akhir
Proses:	
1. Inisialisasi model awal dengan nilai konstan:	
$F_0(x) = \operatorname{argmin}_c \sum L(y_i, c)$	
2. Untuk $k = 1$ sampai K (iterasi pohon ke-1 hingga 100):	
a. Hitung statistik gradien orde pertama (g_i) dan kedua (h_i) untuk setiap data ke-i berdasarkan prediksi iterasi sebelumnya:	
$g_i = \partial L(y_i, F_{\{k-1\}}(x_i)) / \partial F_{\{k-1\}}(x_i)$	
$h_i = \partial^2 L(y_i, F_{\{k-1\}}(x_i)) / \partial F_{\{k-1\}}(x_i)^2$	
b. Bangun pohon keputusan $f_k(x)$ dengan batasan ' $max_depth = 6$ ':	
- Tentukan struktur pohon dengan mencari titik pemisahan (split) terbaik yang memaksimalkan Gain:	
$Gain = 1/2 * [(\sum g_L)^2 / (\sum h_L + \lambda) + (\sum g_R)^2 / (\sum h_R + \lambda) - (\sum g)^2 / (\sum h + \lambda)] - \gamma$	
- Tetapkan bobot optimal (w) pada setiap daun (j):	
$w_j = - (\sum_{i \in I_j} g_i) / (\sum_{i \in I_j} h_i + \lambda)$	
c. Perbarui model ensemble dengan menambahkan pohon baru yang telah diskalakan oleh learning rate ($\eta = 0.3$):	
$F_k(x) = F_{\{k-1\}}(x) + \eta * f_k(x)$	

3. Kembalikan model akhir $F_K(x)$

Pada Tabel 2.1 terdapat algoritma *XGBoost default* diawali dengan inisialisasi model prediksi dasar (F_0) yang memberikan nilai konstan untuk meminimalkan fungsi objektif awal. Proses pembelajaran kemudian berlangsung secara iteratif sebanyak 100 putaran ($n_estimators$), di mana pada setiap iterasinya, sistem menghitung statistik gradien orde pertama (g_i) dan orde kedua (h_i) dari fungsi kerugian (*loss function*) terhadap prediksi sebelumnya untuk mengidentifikasi arah kesalahan model. Berdasarkan nilai gradien tersebut, sebuah pohon keputusan baru dibangun untuk memprediksi *residual error*, dengan kompleksitas struktur yang dibatasi oleh kedalaman maksimal (*max_depth*) sebesar 6 level guna mencegah terjadinya *overfitting*. Struktur pohon ditentukan melalui pencarian titik pemisahan terbaik yang memaksimalkan nilai *Gain*, dan setiap daun pada pohon tersebut diberikan bobot optimal (w) yang merepresentasikan nilai koreksi harga. Akhirnya, model diperbarui secara aditif pada setiap langkah dengan menambahkan kontribusi pohon baru yang telah diskalakan oleh *learning rate* (η) sebesar 0.3, sehingga menghasilkan model *ensemble* akhir (F_K) yang merupakan akumulasi cerdas dari seluruh pohon keputusan yang terbentuk.

Berikut merupakan formulasi matematika dari *XGBoost* yang mana prediksi akhir $\hat{y}_i$ untuk data ke- i diperoleh dari penjumlahan seluruh fungsi prediksi dari κ pohon keputusan. Masing-masing fungsi $f_{k(x_i)}$ adalah model pohon ke- k yang termasuk dalam himpunan $\mathcal{F}$, yaitu kumpulan semua pohon keputusan yang dilatih. Penjumlahan seluruh fungsi prediksi dari κ pohon keputusan dihitung dengan Persamaan 2.1.

$$\hat{y}_i = \sum_{k=1}^K f_k(x_i), \quad f_k \in \mathcal{F} \quad (2.1)$$

Fungsi kerugian $l(y_i, \hat{y}_i)$ yang menghitung seberapa jauh prediksi terhadap nilai sebenarnya. Fungsi regularisasi $\Omega(f_t)$ yang mengontrol kompleksitas model agar tidak *overfitting*. Fungsi objektif yang terdiri dari dua bagian: *loss function* dan *regularization term* dihitung dengan Persamaan 2.2.

$$L^{(t)} = \sum_{i=1}^n l(y_i, \hat{y}_i^{(t-1)} + f_{t(x_i)}) + \Omega(f_t) \quad (2.2)$$

Fungsi regularisasi ini berfungsi untuk mengendalikan kompleksitas dari pohon keputusan f_k . Dimana γ dan λ adalah parameter regulasi untuk menghindari pohon menjadi terlalu rumit, T adalah jumlah daun pohon, dan ω_j adalah bobot pada daun ke- j . Fungsi regularisasinya sendiri dihitung dengan Persamaan 2.3.

$$\Omega(f_k) = \gamma T + \frac{1}{2} \lambda \sum_{j=1}^T \omega_j^2 \quad (2.3)$$

Untuk memudahkan optimasi, fungsi objektif ini dikembangkan dengan pendekatan Taylor, yang melibatkan turunan pertama dan kedua dari fungsi *loss*. Fungsi objektif dengan ekspansi Taylor orde dua dihitung dengan Persamaan 2.4.

$$L^{(t)} \approx \sum_{i=1}^n \left[g_i f_t(x_i) + \frac{1}{2} h_i f_t^2(x_i) \right] + \Omega(f_t) \quad (2.4)$$

Untuk menentukan seberapa besar kontribusi *output* dari daun ke- j terhadap prediksi akhir. Bobot ini akan digunakan saat pohon tersebut menyumbangkan nilainya dalam proses penjumlahan prediksi akhir ($\hat{y}_i$). Agar nilai *output* dari

setiap daun dalam pohon bisa meminimalkan kesalahan prediksi sambil tetap memperhatikan regularisasi dapat dihitung dengan menggunakan Persamaan 2.5.

$$\omega_j^* = - \frac{\sum_{i \in I_j} g_i}{\sum_{i \in I_j} h_i + \lambda} \quad (2.5)$$

Persamaan 2.6 digunakan untuk menghitung nilai minimum dari fungsi objektif setelah optimasi bobot daun. Semakin kecil nilai ini, semakin baik struktur pohon dalam meminimalkan kesalahan prediksi. Nilai minimal dari fungsi objektif setelah pemangkasan pohon dapat dihitung dengan Persamaan 2.6.

$$\tilde{L}^{(t)} = -\frac{1}{2} \sum_{j=1}^T \frac{(\sum_{i \in I_j} g_i)^2}{\sum_{i \in I_j} h_i + \lambda} + \gamma T \quad (2.6)$$

Nilai *Gain* mengukur seberapa besar peningkatan kualitas model jika data dibagi, semakin besar *Gain*, semakin baik pemisahan yang dipilih. Selanjutnya, untuk memilih pemisahan terbaik pada setiap *node*, *XGBoost* menghitung nilai *Gain* dengan menggunakan Persamaan 2.7.

$$Gain = \frac{1}{2} \left[\frac{(\sum_{i \in I_L} g_i)^2}{\sum_{i \in I_L} h_i + \lambda} + \frac{(\sum_{i \in I_R} g_i)^2}{\sum_{i \in I_R} h_i + \lambda} - \frac{(\sum_{i \in I} g_i)^2}{\sum_{i \in I} h_i + \lambda} \right] - \gamma \quad (2.7)$$

2.1.2 Hyperparameter Tuning

Hyperparameter Tuning merupakan aspek penting dalam *machine learning*, karena berperan dalam mengatur perilaku serta tingkat kompleksitas model, yang secara langsung memengaruhi kinerja dan akurasi prediksi (Ilemobayo dkk., 2024). *Hyperparameter Tuning* adalah proses untuk menemukan kombinasi parameter terbaik yang dapat meningkatkan performa model *machine learning*.

Hyperparameter Tuning sangat penting untuk meningkatkan kinerja dan kemampuan model *machine learning* dalam mengenali pola secara akurat.

Beberapa faktor utama yang memengaruhi performa model, seperti kualitas data, jenis algoritma yang digunakan, dan tingkat kompleksitas model, juga berperan besar dalam hasil akhir. Selain itu, pengaturan *Hyperparameter* seperti laju pembelajaran dan ukuran *batch* dapat berdampak langsung pada proses pelatihan (Ilemobayo dkk., 2024). *Hyperparameter* berbeda dari parameter model karena tidak diperoleh melalui proses pelatihan, melainkan harus ditentukan sebelum pelatihan dilakukan. Teknik *Tuning* yang umum digunakan meliputi *Grid Search*, *Random Search*, *Bayesian Optimization*, dan *Tree-structured Parzen Estimator* (TPE) digunakan untuk menemukan kombinasi terbaik guna meningkatkan akurasi dan efisiensi model. Penelitian tersebut juga menyajikan contoh hasil *Tuning* yang menunjukkan peningkatan akurasi model serta memberikan rekomendasi untuk penyetelan model yang lebih efisien dan efektif (Koch dkk., 2017). Dengan begitu penelitian ini akan menggunakan *Bayesian Optimization* sebagai metode *Hyperparameter Tuning* yang dipakai.

2.1.2.1 Bayesian Optimization

Bayesian Optimization adalah metode yang menggunakan model probabilistik untuk mencari kombinasi *Hyperparameter* terbaik secara efisien. Pengoptimalan *Bayesian* telah terbukti efektif dalam menyesuaikan pengaturan *Hyperparameter* secara otomatis, sehingga membantu berbagai model bekerja lebih optimal dan efisien (Yang dkk., 2024). Pendekatan ini menggunakan fungsi perkiraan untuk memprediksi kinerja model berdasarkan sampel yang telah diuji,

kemudian menentukan kombinasi *Hyperparameter* baru yang diharapkan dapat meningkatkan performa secara keseluruhan. Salah satu implementasi populer dari *Bayesian Optimization* adalah *BayesSearchCV*, yang digunakan dalam penelitian ini.

Penelitian ini membahas bagaimana teknik Optimasi *Bayesian* digunakan untuk menyelesaikan masalah optimasi dalam fisika akselerator. Algoritma ini dianggap unggul karena dapat memprediksi hasil dengan lebih efisien menggunakan model statistik sebagai perkiraan, tanpa harus menguji setiap kemungkinan secara langsung. Hal ini sangat bermanfaat dalam menangani tantangan yang kompleks, terutama ketika ada gangguan atau kebisingan selama operasi akselerator, serta dalam simulasi fisika yang membutuhkan banyak sumber daya (Roussel dkk., 2024). Dengan keberhasilan Optimasi *Bayesian* dalam menangani berbagai masalah kompleks, serta pengakuan luas terhadap keunggulannya dalam fisika akselerator, penelitian yang akan dilakukan memiliki potensi besar untuk mencapai hasil yang optimal dan berkontribusi secara signifikan dalam meningkatkan efisiensi serta akurasi model yang dikembangkan.

2.1.3 Evaluation Metrics

Metrik seperti *Mean Absolute Error* (MAE), *Mean Squared Error* (MSE), *Root Mean Squared Error* (RMSE), dan *R-squared* (R^2) digunakan untuk mengevaluasi model (Kumar, 2023).

2.1.3.1 Mean Absolute Error (MAE)

Mean Absolute Error (MAE) merupakan rata-rata dari selisih absolut antara nilai aktual dan nilai prediksi. MAE memberikan bobot yang sama pada setiap

perbedaan dan tidak terlalu sensitif terhadap *outlier*. Nilai MAE dapat dihitung dengan Persamaan 2.8.

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad (2.8)(\text{Kumar, 2023})$$

2.1.3.2 Mean Squared Error (MSE)

Mean Squared Error (MSE) adalah rata-rata dari kuadrat selisih antara nilai estimasi dan nilai aktual. MSE memberikan penalti yang lebih besar terhadap kesalahan yang lebih besar, sehingga model dengan nilai MSE yang lebih kecil dianggap lebih baik karena memiliki tingkat kesalahan yang lebih rendah. Nilai MSE dapat dihitung dengan Persamaan 2.9.

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (2.9)(\text{Kumar, 2023})$$

2.1.3.3 Root Mean Squared Error (RMSE)

Root Mean Squared Error (RMSE) dihitung dengan mengambil akar kuadrat dari MSE. RMSE memberikan penalti lebih besar untuk kesalahan yang besar, sehingga metrik ini berguna ketika kesalahan besar tidak diinginkan. Nilai RMSE dapat dihitung dengan Persamaan 2.10.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (2.10)(\text{Kumar, 2023})$$

2.1.3.4 Koefisien Determinasi (R^2)

Koefisien Determinasi (R^2) adalah ukuran statistik yang menunjukkan seberapa baik variabel independen dapat menjelaskan perubahan dalam variabel dependen. Semakin mendekati 1 nilai R^2 , semakin baik model dalam menjelaskan variasi dalam data. Nilai R^2 dapat dihitung dengan Persamaan 2.11.

$$R^2 = 1 - \frac{RSS}{TSS} \quad (2.11)(\text{Kumar, 2023})$$

Keterangan:

y_i = nilai aktual,

$\hat{y}_i$ = nilai prediksi,

n = jumlah total data,

RSS = *Residual Sum of Squares*,

TSS = *Total Sum of Squares*.

2.2 Penelitian Terkait

2.2.1 *State of The Art*

Pada tabel 2.2 dapat dilihat beberapa penelitian terdahulu yang telah mengkaji penerapan *machine learning* dalam prediksi harga tiket pesawat maupun optimasi model menggunakan *Bayesian Optimization*, yang menjadi landasan bagi penelitian ini.

Tabel 2. 2 *State of The Art*

No	Peneliti	Judul	Objek Penelitian	Algoritma	Hasil Penelitian
1	(Dalal dkk., 2022)	<i>Hybrid XGBoost model with Hyperparameter Tuning for prediction of liver disease with better accuracy</i>	<i>Prediction of liver disease</i>	<i>Hybrid XGBoost</i>	Algoritma pohon keputusan seperti CART dan CHAID tidak memberikan tingkat akurasi yang tinggi dalam memprediksi penyakit hati pada kelompok pasien di India. Namun, model <i>Hybrid XGBoost</i> yang dikembangkan dalam penelitian ini

					menunjukkan tingkat akurasi yang lebih baik, mencapai 93,65%.
2	(Porter & Tech, n.d.)	<i>Studying the Acquisition Function of Bayesian Optimization With Machine learning With DNA Reads.</i>	<i>DNA Reads</i>	<i>Bayesian Optimization With Machine learning</i>	Penelitian ini menganalisis pengaruh fungsi akuisisi dalam optimasi menggunakan random forests dan multilayer perceptrons pada prediksi performa pemetaan DNA pendek yang telah mengalami perlakuan bisulfit, metode sekuensing untuk mempelajari metilasi epigenetik. Hasilnya menunjukkan bahwa <i>Bayesian Optimization</i> memiliki performa serupa dengan grid search, namun kurang efisien dalam

					paralelisasi saat diimplementasikan menggunakan skopt.
3	(Tiasama & Budi, 2024)	Perbandingan <i>Random Search</i> dan Algoritma Genetika dalam Penyetelan <i>Hyperparameter XGBoost</i> pada <i>Retail Sales Forecasting</i>	<i>Retail Sales Forecasting</i>	Perbandingan <i>Random Search</i> dan Algoritma Genetika dalam Penyetelan <i>Hyperparameter XGBoost</i>	Penelitian ini memperoleh hasil <i>Random Search</i> lebih unggul dari algoritma genetika dengan nilai RMSE pada proses latih dan proses uji sebesar 0.123 dan 0.122. Sementara itu, nilai RMSE yang dihasilkan oleh algoritma genetika pada proses latih dan proses uji yaitu sebesar 0.333 dan 0.332.
4	(H. Darmawan dkk., 2023)	<i>GRU and XGBoost Performance with</i>	<i>Weather Prediction System</i>	<i>GRU and XGBoost</i>	Hasil pengujian dengan data dari BMKG Juanda dan IoT menunjukkan

		<i>Hyperparameter Tuning</i> <i>Using GridSearchCV and Bayesian Optimization on an IoT-Based Weather Prediction System</i>		Performance with Hyperparameter Tuning Using GridSearchCV and Bayesian Optimization	bahwa dalam tahap pertama, <i>XGBoost</i> regression lebih unggul dengan RMSE 1.2728 dibanding GRU regression (RMSE 1.5517). Namun, pada tahap kedua, GRU regression lebih baik (RMSE 2.23) dibanding <i>XGBoost</i> (RMSE 2.28). Dalam klasifikasi, GRU memiliki skor F1 0.88, lebih tinggi dibanding <i>XGBoost</i> (0.86), meskipun keduanya memiliki akurasi yang sama (0.75) dengan data IoT. Model GRU classification dianggap lebih unggul karena mempertimbangkan konteks
--	--	--	--	--	--

					prediksi, sehingga lebih akurat dalam memprediksi kemungkinan hujan.
5	(N dkk., 2024)	<i>XGBOOST HYPERPARAMETER OPTIMIZATION USING RANDOMIZEDSEARCHCV FOR ACCURATE FOREST FIRE DROUGHT CONDITION PREDICTION</i>	<i>ACCURATE FOREST FIRE DROUGHT CONDITION PREDICTION</i>	<i>XGBOOST HYPERPAR AMETER OPTIMIZA TION USING RANDOMI ZEDSEARC HCV</i>	Hasil penelitian menunjukkan bahwa model <i>XGBoost</i> yang telah dioptimasi menghasilkan penurunan <i>Mean Squared Error</i> (MSE) dan peningkatan nilai <i>R-squared</i> dibandingkan dengan model default. Model yang dioptimasi mencapai MSE sebesar 0.0210 dan R2 sebesar 0.9820 pada data uji, yang mengindikasikan peningkatan akurasi prediksi secara signifikan terhadap

					kondisi kekeringan bahan bakar hutan.
6	(Putatunda & Rama, 2019)	A Modified <i>Bayesian Optimization</i> based Hyper-Parameter <i>Tuning</i> Approach for <i>Extreme Gradient Boosting</i>	<i>Extreme Gradient Boosting</i>	<i>Bayesian Optimization</i> based Hyper-Parameter <i>Tuning</i>	Penelitian ini mengusulkan metode Randomized-Hyperopt untuk meningkatkan optimasi <i>Hyperparameter</i> pada <i>XGBoost</i> . Dibandingkan dengan Grid Search, Random Search, dan Hyperopt, uji coba pada 10 dataset menunjukkan bahwa Randomized-Hyperopt memberikan hasil terbaik dalam hal akurasi prediksi dan waktu eksekusi.

7	(Elina dkk., 2022)	Penerapan Metode <i>Extreme Gradient Boosting (XGBOOST)</i> pada Klasifikasi Nasabah Kartu Kredit	Klasifikasi Nasabah Kartu Kredit	<i>Extreme Gradient Boosting (XGBOOST)</i>	Hasil dari penelitian ini membuktikan bahwa penggunaan algoritma dengan <i>Hyperparameter Tuning</i> dapat meningkatkan kinerja algoritma <i>Extreme Gradient Boosting</i> dalam proses klasifikasi nasabah kartu kredit dengan akurasi sebesar 80,039%, presisi sebesar 81,338% dan nilai recall sebesar 96,854%.
8	(M, 2024)	<i>MODELING AND PREDICTING INDONESIA RICE PRICES USING HYPERPARAMETER OPTIMIZATION XGBOOST</i>	<i>PREDICTING INDONESIA RICE PRICES</i>	<i>HYPERPARAMETER OPTIMIZATION XGBOOST</i>	Studi ini menggunakan model <i>XGBoost</i> dengan <i>Hyperparameter Tuning</i> dan <i>cross-validation</i> untuk meramalkan harga beras di Indonesia hingga dua tahun setelah Juli 2024.

					Dengan data Januari 2010 – Juli 2024 dan <i>4-fold cross-validation</i> , model dievaluasi melalui 26.000+ fits. Hasilnya menunjukkan <i>XGBoost</i> mampu menangkap pola kompleks dalam data dan menghasilkan prediksi akurat.
9	(Dinanthi dkk., 2024)	<i>Diabetes Detection Using Extreme Gradient Boosting (XGBoost) with Hyperparameter Tuning</i>	<i>Diabetes Detection</i>	<i>Extreme Gradient Boosting (XGBoost) with Hyperpara</i>	Studi ini meneliti kinerja <i>XGBoost</i> dalam klasifikasi diabetes dengan optimasi GridSearchCV dan RandomSearchCV serta penyeimbangan data menggunakan SMOTE. Hasilnya, GridSearchCV

				<i>meter Tuning</i>	mencapai akurasi 85%, sementara RandomSearchCV mencapai 83%.
10	(Rhouas dkk., 2024)	<i>Enhancing currency prediction in international e-commerce: Bayesian-optimized random forest approach using the Klarna dataset</i>	<i>Prediction in international e-commerce</i>	<i>Bayesian-optimized random forest</i>	Penelitian ini mengoptimalkan algoritma Random Forest menggunakan <i>Bayesian Optimization</i> pada dataset Klarna <i>E-commerce</i> untuk meningkatkan prediksi nilai tukar mata uang dalam perdagangan global. Dengan 14 iterasi, 1 job, dan random state 684, model mencapai MSE -0.9891 dan akurasi 74.63%. Studi ini menyoroti peran <i>Bayesian Optimization</i> dalam meningkatkan kinerja model, memberikan strategi

					yang lebih baik bagi bisnis dalam manajemen risiko keuangan global.
--	--	--	--	--	---

Tabel 2.2 merangkum berbagai penelitian yang membahas optimasi *XGBoost* dan *Bayesian Optimization* di berbagai bidang, seperti prediksi penyakit, klasifikasi, peramalan harga, dan sistem berbasis IoT. Secara umum, penelitian-penelitian ini menunjukkan bahwa penyetelan *Hyperparameter*, terutama dengan *Bayesian Optimization*, *GridSearchCV*, dan *RandomizedSearchCV*, dapat meningkatkan performa model secara signifikan. Misalnya, dalam prediksi penyakit hati, penggunaan *Hybrid XGBoost* berhasil mencapai tingkat akurasi 93,65% (Dalal dkk., 2022). Selain itu, *Bayesian Optimization* juga diterapkan untuk meningkatkan efisiensi pemetaan DNA *Reads*, meskipun hasilnya masih menunjukkan bahwa metode ini perlu dioptimalkan lebih lanjut dalam hal paralelisasi (Porter & Tech, n.d.).

Dalam *Retail Sales Forecasting*, penelitian menemukan bahwa *Random Search* lebih unggul dibandingkan Algoritma Genetika, dengan nilai RMSE yang lebih rendah (Tiasama & Budi, 2024). Sementara itu, dalam prediksi cuaca berbasis IoT, kombinasi GRU dan *XGBoost* memberikan hasil yang bervariasi. *XGBoost* lebih unggul dalam regresi tahap awal, sementara GRU lebih baik dalam klasifikasi, ditunjukkan dengan skor F1 yang lebih tinggi (H. Darmawan dkk., 2023). Dalam bidang mitigasi bencana, optimasi *Hyperparameter XGBoost* menggunakan *RandomizedSearchCV* untuk prediksi kekeringan akibat kebakaran hutan berhasil meningkatkan akurasi model dengan R^2 sebesar 0.982, jauh lebih baik dibandingkan model default (N dkk., 2024)

Lebih lanjut, penelitian oleh Putatunda & Rama menemukan bahwa metode *Randomized-Hyperopt* lebih efektif dibandingkan *Grid Search* dan *Random Search*

dalam menyetel *Hyperparameter XGBoost*. Di bidang keuangan, model *XGBoost* diterapkan untuk klasifikasi nasabah kartu kredit, dan hasilnya menunjukkan bahwa penyetelan *Hyperparameter* dapat meningkatkan akurasi hingga 80,039% (Elina dkk., 2022). Studi lain juga membuktikan bahwa *XGBoost* dengan *Hyperparameter Tuning* dan *cross-validation* mampu mengidentifikasi pola kompleks dalam data historis harga beras di Indonesia, menghasilkan prediksi yang lebih akurat (M, 2024). Selain itu, dalam deteksi diabetes, penggunaan *GridSearchCV* menghasilkan akurasi tertinggi sebesar 85%, lebih tinggi dibandingkan dengan *RandomSearchCV* (Dinanthi dkk., 2024).

Terakhir, *Optimasi Bayesian* juga diterapkan dalam model *Random Forest* untuk prediksi nilai tukar mata uang dalam *e-commerce*, dengan hasil yang menunjukkan peningkatan akurasi hingga 74,63% (Rhous dkk., 2024). Berdasarkan berbagai penelitian sebelumnya, dapat disimpulkan bahwa optimasi *Hyperparameter*, khususnya dengan *Bayesian Optimization*, memiliki peran besar dalam meningkatkan akurasi model berbasis *XGBoost*. Oleh karena itu, penelitian ini bertujuan untuk memperkuat temuan tersebut dengan menerapkan metode serupa dalam prediksi harga tiket pesawat, serta mengevaluasi seberapa efektif *Bayesian Optimization* dibandingkan metode *Tuning* lainnya.