

Penerapan Algoritma XGBoost untuk Prediksi Tingkat Pemahaman Siswa pada Praktikum Laboratorium Komputer

Kecitaan Harefa^{1*}, Rinna Rachmatika², Oktafpianus Gea³, Arojasa Harefa⁴

^{1,2,3,4}Ilmu Komputer, Teknik Informatika, Universitas Pamulang, Tangerang Selatan, Indonesia

Email: ^{1*}dosen00842@unpam.ac.id, ²rinnarachmatika@unpam.ac.id, ³oxthagea@gmail.com,

⁴arojasaHarefa03@gmail.com

(*Email Corresponding Author: dosen00842@unpam.ac.id)

Received: 20 Agustus 2026 | Revision: 28 Agustus 2026 | Accepted: 3 September 2026

Abstrak

Evaluasi tingkat pemahaman siswa pada praktikum laboratorium komputer masih sering dilakukan secara manual berdasarkan nilai dan observasi guru, sehingga membutuhkan waktu dan berpotensi menimbulkan subjektivitas. Penelitian ini bertujuan menerapkan algoritma Extreme Gradient Boosting (XGBoost) sebagai solusi untuk memprediksi tingkat pemahaman siswa secara objektif, cepat, dan berbasis data. Penelitian menggunakan pendekatan kuantitatif dengan 120 data siswa yang terbagi seimbang ke dalam kategori rendah, sedang, dan tinggi. Variabel prediktor meliputi nilai kuis, nilai praktikum, kehadiran, waktu pengerjaan, keaktifan, dan observasi guru. Tahapan penelitian mencakup pengumpulan dan prapemrosesan data, pengodean label, serta pembagian data secara terstratifikasi dengan komposisi 80% data latih dan 20% data uji. Model kemudian dilatih menggunakan XGBoost dan dievaluasi melalui matriks konfusi, akurasi, presisi, recall, dan F1-score. Hasil pengujian menunjukkan bahwa model berhasil mengklasifikasikan 23 dari 24 data uji dengan benar, dengan akurasi sebesar 95,83%, presisi 95,62%, recall 95,83%, dan F1-score 95,69%. Analisis kepentingan fitur menunjukkan bahwa nilai kuis merupakan variabel paling berpengaruh. Dengan demikian, XGBoost efektif digunakan sebagai pendukung evaluasi pembelajaran untuk membantu guru mengidentifikasi tingkat pemahaman siswa secara lebih akurat dan efisien.

Kata Kunci: Xgboost, Pembelajaran Mesin, Prediksi, Tingkat Pemahaman Siswa, Praktikum Laboratorium Komputer.

Abstract

Evaluation of students' understanding levels in computer laboratory practicum is still often done manually based on teacher grades and observations, so it takes time and has the potential to cause subjectivity. This study aims to apply the Extreme Gradient Boosting (XGBoost) algorithm as a solution to predict the level of student understanding objectively, quickly, and based on data. The research uses a quantitative approach with 120 student data which is divided into low, medium, and high categories. Predictive variables include quiz scores, practicum scores, attendance, processing time, activeness, and teacher observation. The research stages include data collection and preprocessing, label coding, and stratified data distribution with a composition of 80% training data and 20% test data. The model is then trained using XGBoost and evaluated through a matrix of confusion, accuracy, precision, recall, and F1-score. The test results showed that the model managed to correctly classify 23 out of 24 test data, with an accuracy of 95.83%, precision of 95.62%, recall of 95.83%, and an F1-score of 95.69%. Analysis of the importance of features showed that quiz scores were the most influential variables. Thus, XGBoost is effectively used as a support for learning evaluations to help teachers identify students' levels of understanding more accurately and efficiently.

Keywords: Xgboost, Machine Learning, Prediction, Student Comprehension Level, Computer Lab Practicum.

1. PENDAHULUAN

Praktikum laboratorium komputer merupakan bagian penting dalam pembelajaran Teknologi Informasi dan Komunikasi karena memberikan kesempatan kepada siswa untuk menerapkan konsep yang telah dipelajari melalui kegiatan secara langsung [1]. Keberhasilan kegiatan praktikum tidak hanya ditentukan oleh terselesaikannya tugas, tetapi juga oleh tingkat pemahaman siswa terhadap materi, prosedur, dan penggunaan perangkat komputer [2]. Setiap siswa memiliki tingkat pemahaman yang berbeda karena dipengaruhi oleh kemampuan akademik, intensitas kehadiran, keaktifan selama pembelajaran, ketepatan menyelesaikan tugas, serta hasil pengamatan guru. Perbedaan tersebut mengharuskan guru melakukan evaluasi secara tepat agar siswa yang belum memahami materi dapat segera memperoleh pendampingan [3]. Namun, penilaian tingkat pemahaman siswa masih sering dilakukan secara manual dengan berfokus pada nilai akhir atau berdasarkan pengamatan guru. Cara tersebut memerlukan waktu, terutama ketika jumlah siswa cukup banyak, serta berpotensi menghasilkan penilaian yang kurang konsisten karena belum memanfaatkan seluruh data pembelajaran secara terintegrasi.

Nilai kuis dan nilai praktikum pada dasarnya belum sepenuhnya menggambarkan tingkat pemahaman siswa. Seorang siswa dapat memperoleh nilai yang baik, tetapi membutuhkan waktu pengerjaan yang lebih lama atau kurang aktif selama kegiatan praktikum. Sebaliknya, siswa yang memperoleh nilai lebih rendah mungkin menunjukkan keaktifan, kehadiran, dan kemampuan praktik yang baik [4]. Oleh karena itu, penilaian tingkat pemahaman siswa perlu mempertimbangkan beberapa indikator secara bersamaan, seperti nilai kuis, nilai praktikum, kehadiran, waktu pengerjaan, keaktifan, dan observasi guru. Pemanfaatan data tersebut diharapkan dapat menghasilkan evaluasi yang lebih

menyeluruh dan objektif. Permasalahan lainnya adalah belum tersedianya suatu model prediksi yang dapat mengolah berbagai indikator tersebut untuk mengelompokkan tingkat pemahaman siswa ke dalam kategori rendah, sedang, dan tinggi, khususnya dalam pelaksanaan praktikum laboratorium komputer pada tingkat Sekolah Menengah Pertama [5].

Salah satu solusi yang dapat diterapkan adalah pemanfaatan machine learning melalui algoritma Extreme Gradient Boosting atau XGBoost. Algoritma XGBoost merupakan pengembangan dari metode gradient boosting yang membangun sejumlah pohon keputusan secara bertahap. Setiap pohon berikutnya diarahkan untuk memperbaiki kesalahan prediksi dari pohon sebelumnya [6]. XGBoost dilengkapi dengan teknik regularisasi yang dapat mengendalikan kompleksitas model, mengurangi risiko overfitting, dan meningkatkan kemampuan generalisasi. Algoritma ini juga mampu mengolah hubungan yang kompleks dan tidak linier antarfaktor serta menyediakan informasi mengenai tingkat kepentingan setiap fitur. Dengan karakteristik tersebut, XGBoost dapat digunakan untuk mempelajari pola dari data pembelajaran dan memprediksi tingkat pemahaman siswa. Hasil prediksi diharapkan dapat menjadi informasi pendukung bagi guru dalam melakukan evaluasi, memberikan pendampingan, dan menentukan strategi pembelajaran yang sesuai dengan kebutuhan siswa [7].

Pemanfaatan machine learning dalam memprediksi hasil belajar siswa telah dikaji dalam beberapa penelitian selama lima tahun terakhir. Yang dkk. [8] mengembangkan kerangka machine learning yang mengintegrasikan XGBoost dan beberapa algoritma lainnya untuk memprediksi kinerja akademik siswa. Model tersebut menghasilkan akurasi sebesar 99,6% dan 97,5%. Namun, penelitian ini masih berfokus pada identifikasi siswa berisiko secara umum dan belum secara khusus memprediksi tingkat pemahaman siswa pada praktikum laboratorium komputer berdasarkan indikator akademik dan perilaku belajar.

Cheng dkk. [9] menerapkan XGBoost yang dioptimalkan dengan Enhanced Artificial Ecosystem-Based Optimization (EAEBO) untuk memprediksi prestasi akademik siswa. Model tersebut memperoleh akurasi 94,17% dan F1-score 94,13%. Namun, penelitian ini masih berfokus pada prestasi akademik secara umum dan belum secara khusus memprediksi tingkat pemahaman siswa dalam praktikum laboratorium komputer berdasarkan hasil penilaian, keaktifan, dan observasi guru.

Guo dkk. [10] menggunakan XGBoost dan SHAP untuk memprediksi tingkat depresi dan stres mahasiswa berdasarkan data perilaku, psikometrik, dan fisiologis. Model menunjukkan kemampuan prediksi yang baik, dengan nilai AUC lebih dari 0,90 untuk kategori depresi berat. Namun, penelitian tersebut berfokus pada kesehatan mental mahasiswa dan belum memprediksi tingkat pemahaman siswa dalam kegiatan praktikum laboratorium komputer berdasarkan indikator akademik dan perilaku belajar.

Natras dkk. [11] membandingkan XGBoost, Random Forest, AdaBoost, dan Decision Tree untuk memprediksi Vertical Total Electron Content (VTEC). Hasilnya menunjukkan bahwa metode ensemble, termasuk XGBoost, mampu memberikan akurasi lebih baik dibandingkan Decision Tree. Namun, penelitian tersebut diterapkan pada bidang cuaca antariksa dan belum digunakan untuk memprediksi tingkat pemahaman siswa berdasarkan indikator akademik dan perilaku selama praktikum.

Hakkal dan Lahcen [12] menerapkan XGBoost untuk meningkatkan prediksi kinerja siswa pada Intelligent Tutoring System menggunakan delapan dataset. XGBoost meningkatkan kinerja model PFA pada tujuh dataset dengan nilai AUC hingga 0,88. Namun, penelitian tersebut berfokus pada riwayat jawaban siswa dalam pembelajaran daring dan belum memprediksi tingkat pemahaman pada praktikum laboratorium komputer berdasarkan indikator akademik, keaktifan, dan observasi guru.

Selanjutnya, Villar dan Andrade [13] membandingkan XGBoost dengan beberapa algoritma untuk memprediksi keberhasilan akademik dan risiko putus kuliah. Hasilnya menunjukkan bahwa LightGBM dan CatBoost memberikan kinerja lebih baik daripada XGBoost dan metode klasifikasi tradisional. Namun, penelitian tersebut berfokus pada status kelulusan dan putus kuliah mahasiswa, bukan tingkat pemahaman siswa dalam praktikum laboratorium komputer berdasarkan indikator akademik dan perilaku belajar.

Berdasarkan penelitian terdahulu tersebut, terdapat kesenjangan penelitian (research gap) pada aspek objek, variabel, konteks penerapan, dan keluaran prediksi. Penelitian dalam bidang pendidikan umumnya berfokus pada prediksi prestasi akademik, risiko putus kuliah, dan kemampuan siswa menjawab soal dengan memanfaatkan data akademik, demografis, serta rekaman interaksi pada sistem pembelajaran daring [14]. Sementara itu, penelitian yang secara khusus mengklasifikasikan tingkat pemahaman siswa SMP dalam kegiatan praktikum laboratorium komputer masih terbatas. Belum banyak penelitian yang mengintegrasikan enam indikator, yaitu Nilai Kuis, Nilai Praktikum, Kehadiran, Waktu Pengerjaan, Keaktifan, dan Observasi Guru, ke dalam satu model XGBoost. Penelitian ini mengevaluasi kinerja model menggunakan akurasi, presisi, recall, dan F1-score, serta menganalisis tingkat kepentingan fitur untuk mengetahui indikator yang paling berpengaruh terhadap tingkat pemahaman siswa.

Dengan demikian, penelitian ini bertujuan menerapkan algoritma XGBoost untuk mengklasifikasikan tingkat pemahaman siswa pada praktikum laboratorium komputer ke dalam kategori rendah, sedang, dan tinggi. Penelitian ini juga bertujuan mengevaluasi kinerja model serta mengidentifikasi variabel yang paling berpengaruh terhadap hasil prediksi. Model yang dihasilkan diharapkan dapat membantu guru melakukan evaluasi secara lebih cepat, konsisten, dan berbasis data. Hasil prediksi dapat menjadi bahan pertimbangan dalam menentukan tindak lanjut pembelajaran, seperti pemberian pengayaan kepada siswa dengan tingkat pemahaman tinggi, penguatan materi kepada siswa dengan tingkat pemahaman sedang, serta pendampingan atau remedial kepada siswa dengan tingkat pemahaman rendah.

2. METODOLOGI PENELITIAN

Penelitian ini menggunakan metode kuantitatif dengan pendekatan machine learning melalui algoritma Extreme Gradient Boosting (XGBoost). Algoritma tersebut digunakan untuk membangun model klasifikasi tingkat pemahaman siswa berdasarkan data hasil praktikum laboratorium komputer. Tingkat pemahaman siswa diklasifikasikan ke dalam tiga kategori, yaitu rendah, sedang, dan tinggi. Pengolahan data dilakukan menggunakan bahasa pemrograman Python dengan dukungan pustaka Pandas, Scikit-learn, dan XGBoost.

2.1 Tahapan Penelitian

Tahapan penelitian disusun secara sistematis mulai dari identifikasi masalah hingga analisis hasil prediksi. Tahapan tersebut meliputi identifikasi masalah, pengumpulan data, preprocessing data, pembagian dataset, penerapan algoritma XGBoost, evaluasi model, dan analisis hasil [15]. Alur tahapan penelitian ditunjukkan pada Gambar 1.



Gambar 1. Tahapan Penelitian

Berdasarkan Gambar 1, penelitian diawali dengan identifikasi permasalahan pada kegiatan praktikum laboratorium komputer. Hasil observasi menunjukkan adanya perbedaan tingkat pemahaman siswa terhadap materi praktikum. Selain itu, proses evaluasi masih dilakukan secara manual sehingga guru mengalami kesulitan dalam mengidentifikasi siswa yang memiliki tingkat pemahaman rendah secara cepat dan objektif.

Tahap berikutnya adalah pengumpulan data akademik dan aktivitas siswa selama mengikuti praktikum laboratorium komputer. Data yang dikumpulkan meliputi nilai kuis, nilai praktikum, kehadiran, waktu pengerjaan, keaktifan, dan observasi guru. Dataset penelitian terdiri atas 120 data siswa SMP yang digunakan untuk proses pelatihan dan pengujian model.

Data yang telah dikumpulkan kemudian melalui tahap preprocessing. Tahap ini meliputi pemeriksaan missing value, pemeriksaan data duplikat, pengodean variabel target menggunakan label encoding, dan penentuan variabel masukan serta variabel target [16]. Hasil pemeriksaan menunjukkan bahwa tidak terdapat missing value maupun data duplikat sehingga seluruh data dapat digunakan pada tahap berikutnya.

Setelah preprocessing, dataset dibagi menjadi data training dan testing menggunakan Stratified Train-Test Split. Model XGBoost kemudian dilatih, diuji, dan dievaluasi untuk mengukur kinerja serta mengetahui variabel yang paling berpengaruh.

2.2 Dataset dan Preprocessing Data

Dataset penelitian terdiri atas 120 data siswa SMP yang mengikuti praktikum laboratorium komputer. Data disusun menyerupai kondisi riil pembelajaran dan memiliki variasi nilai pada setiap kategori tingkat pemahaman. Dataset terdiri atas enam variabel independen sebagai masukan model dan satu variabel dependen sebagai target prediksi. Rincian variabel penelitian disajikan pada Tabel 1.

Tabel 1. Variabel Penelitian

No.	Variabel	Jenis Variabel	Keterangan
1	Nilai Kuis	Independen	Nilai teori komputer
2	Nilai Praktikum	Independen	Nilai tugas praktikum
3	Kehadiran	Independen	Persentase kehadiran siswa
4	Waktu Pengerjaan	Independen	Lama penyelesaian tugas
5	Keaktifan	Independen	Tingkat keaktifan siswa
6	Observasi Guru	Independen	Penilaian guru
7	Tingkat Pemahaman	Dependen	Target prediksi

Berdasarkan Tabel 1, nilai kuis, nilai praktikum, kehadiran, waktu pengerjaan, keaktifan, dan observasi guru digunakan sebagai variabel masukan atau fitur (X). Sementara itu, tingkat pemahaman digunakan sebagai variabel target (Y) yang diklasifikasikan ke dalam kategori rendah, sedang, dan tinggi.

Dataset memiliki distribusi kelas yang seimbang. Masing-masing kategori tingkat pemahaman terdiri atas 40 data siswa. Distribusi yang seimbang memberikan kesempatan yang sama kepada model untuk mempelajari karakteristik setiap kelas dan mengurangi potensi bias terhadap salah satu kategori.

Pada tahap preprocessing, pemeriksaan missing value menunjukkan bahwa seluruh variabel memiliki nilai missing value sebesar 0. Pemeriksaan data duplikat juga menunjukkan bahwa tidak terdapat data yang berulang. Dengan demikian, seluruh 120 data dapat digunakan tanpa proses penghapusan maupun pengisian data.

Variabel target yang masih berbentuk kategorikal selanjutnya dikonversi menjadi data numerik menggunakan teknik label encoding. Hasil pengodean kategori tingkat pemahaman disajikan pada Tabel 2.

Tabel 2. Label Encoding Tingkat Pemahaman Siswa

Tingkat Pemahaman	Label
Rendah	0
Sedang	1
Tinggi	2

Berdasarkan Tabel 2, kategori rendah dikonversi menjadi label 0, kategori sedang menjadi label 1, dan kategori tinggi menjadi label 2. Proses tersebut hanya mengubah representasi data dari bentuk teks menjadi numerik agar dapat diproses oleh algoritma XGBoost tanpa mengubah makna setiap kategori.

Dataset selanjutnya dibagi menggunakan metode Stratified Train-Test Split dengan rasio 80:20. Metode ini digunakan untuk mempertahankan proporsi setiap kelas pada data training dan data testing. Nilai random_state=42 digunakan agar pembagian dataset menghasilkan komposisi yang sama apabila penelitian dijalankan kembali. Hasil pembagian dataset disajikan pada Tabel 3.

Tabel 3. Pembagian Dataset

Jenis Dataset	Jumlah Data	Persentase
Data Training	96	80%
Data Testing	24	20%
Total	120	100%

Sebanyak 96 data digunakan untuk melatih model XGBoost, sedangkan 24 data digunakan untuk menguji kemampuan model dalam mengklasifikasikan tingkat pemahaman siswa pada data yang belum pernah dipelajari.

2.3 Penerapan Algoritma XGBoost

Penerapan XGBoost dilakukan melalui tahapan input data, pelatihan model, pembentukan decision tree, proses boosting, prediksi, dan evaluasi. Alur penerapan algoritma XGBoost ditunjukkan pada Gambar 2.



Gambar 2. Tahapan Penerapan Algoritma XGBoost

Berdasarkan Gambar 2, sebanyak 96 data training dimasukkan ke dalam model. Data tersebut terdiri atas enam variabel masukan dan satu variabel target yang telah dikonversi menjadi numerik. Model kemudian mempelajari

hubungan antara nilai kuis, nilai praktikum, kehadiran, waktu pengerjaan, keaktifan, dan observasi guru terhadap kategori tingkat pemahaman siswa.

XGBoost bekerja dengan membentuk sejumlah decision tree secara bertahap. Pohon pertama menghasilkan prediksi awal, kemudian kesalahan prediksi digunakan sebagai dasar dalam pembentukan pohon berikutnya. Proses tersebut disebut boosting. Setiap pohon baru berfungsi memperbaiki kesalahan dari pohon sebelumnya sehingga kombinasi seluruh pohon dapat menghasilkan prediksi yang lebih akurat. Secara umum, hasil prediksi XGBoost dirumuskan dalam Persamaan (1).

$$\hat{y} = \sum_{k=1}^K f_k(x_i) \quad (1)$$

Keterangan:

$\hat{y}$ merupakan hasil prediksi untuk data ke- i ;

K merupakan jumlah decision tree;

f_k merupakan fungsi pohon keputusan ke- k ;

x_i merupakan data masukan ke- i .

Model XGBoost dalam penelitian ini menggunakan parameter yang disesuaikan dengan karakteristik dataset. Konfigurasi parameter model disajikan pada Tabel 4.

Tabel 4. Parameter Algoritma XGBoost

Parameter	Nilai	Keterangan
learning_rate	0,1	Mengatur kecepatan pembelajaran model
max_depth	4	Kedalaman maksimum setiap <i>decision tree</i>
n_estimators	100	Jumlah <i>decision tree</i> yang dibangun
subsample	0,8	Proporsi data <i>training</i> pada setiap iterasi
colsample_bytree	0,8	Proporsi fitur pada pembentukan setiap pohon
objective	multi:softmax	Fungsi objektif klasifikasi multikelas
eval_metric	mlogloss	Metrik kesalahan klasifikasi multikelas
random_state	42	Menjaga konsistensi hasil pelatihan

Parameter learning_rate sebesar 0,1 digunakan agar proses pembelajaran berlangsung secara bertahap. Kedalaman maksimum pohon dibatasi menggunakan max_depth=4, sedangkan jumlah pohon ditetapkan sebanyak 100 melalui parameter n_estimators. Parameter subsample dan colsample_bytree masing-masing sebesar 0,8 digunakan untuk meningkatkan variasi model dan mengurangi risiko overfitting. Parameter multi:softmax digunakan karena penelitian memiliki tiga kelas target, sedangkan mlogloss digunakan untuk mengukur kesalahan prediksi selama pelatihan.

Setelah proses pelatihan selesai, model digunakan untuk memprediksi 24 data testing. Hasil prediksi berupa kategori rendah, sedang, atau tinggi, yang kemudian dibandingkan dengan kelas aktual untuk mengetahui kinerja model.

2.4 Evaluasi Model

Evaluasi model dilakukan menggunakan confusion matrix, accuracy, precision, recall, dan F1-score. Selain itu, analisis feature importance digunakan untuk mengetahui kontribusi setiap variabel terhadap hasil prediksi.

Accuracy menunjukkan persentase keseluruhan data yang berhasil diklasifikasikan dengan benar dan dihitung menggunakan Persamaan (2).

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \times 100\% \quad (2)$$

Precision mengukur ketepatan model ketika memprediksi suatu kelas dan dihitung menggunakan Persamaan (3).

$$\text{Precision} = \frac{TP}{TP+FP} \quad (3)$$

Recall mengukur kemampuan model dalam mengenali seluruh data yang sebenarnya termasuk dalam suatu kelas dan dihitung menggunakan Persamaan (4).

$$\text{Recall} = \frac{TP}{TP+FN} \quad (4)$$

F1-score merupakan rata-rata harmonis antara precision dan recall yang dihitung menggunakan Persamaan (5).

$$\text{F1 - Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (5)$$

Keterangan:

TP (True Positive) merupakan data positif yang diprediksi dengan benar;

TN (True Negative) merupakan data negatif yang diprediksi dengan benar;

FP (False Positive) merupakan data negatif yang diprediksi sebagai positif;
FN (False Negative) merupakan data positif yang diprediksi sebagai negatif.

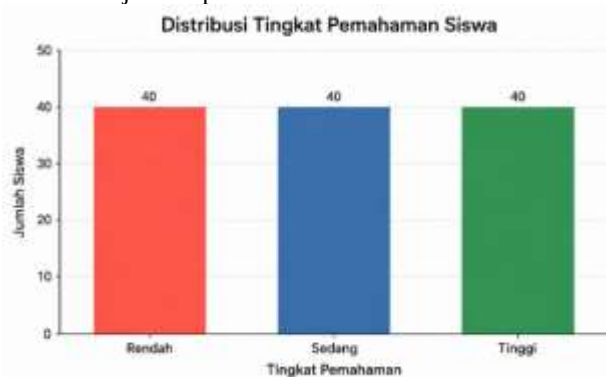
Pada klasifikasi multikelas, perhitungan precision, recall, dan F1-score dilakukan pada setiap kategori tingkat pemahaman, kemudian dirangkum berdasarkan keseluruhan hasil klasifikasi. Confusion matrix digunakan untuk menunjukkan jumlah prediksi benar dan salah pada kelas rendah, sedang, dan tinggi. Sementara itu, feature importance digunakan untuk menentukan besarnya kontribusi nilai kuis, nilai praktikum, kehadiran, waktu pengerjaan, keaktifan, dan observasi guru terhadap hasil prediksi model.

3. HASIL DAN PEMBAHASAN

Bagian iBagian ini menyajikan hasil pengolahan data, implementasi algoritma XGBoost, pengujian model, dan pembahasan hasil prediksi tingkat pemahaman siswa pada praktikum laboratorium komputer. Analisis dilakukan terhadap 120 data siswa dengan enam variabel masukan, yaitu nilai kuis, nilai praktikum, kehadiran, waktu pengerjaan, keaktifan, dan observasi guru. Variabel target berupa tingkat pemahaman siswa yang terdiri atas kategori rendah, sedang, dan tinggi.

3.1 Hasil Pengolahan Data

Dataset penelitian terdiri atas 120 data siswa yang terbagi secara seimbang ke dalam tiga kategori tingkat pemahaman. Setiap kategori rendah, sedang, dan tinggi terdiri atas 40 data siswa atau masing-masing sebesar 33,33%. Distribusi tingkat pemahaman siswa ditunjukkan pada Gambar 3.



Gambar 3. Distribusi Tingkat Pemahaman Siswa

Berdasarkan Gambar 3, dataset memiliki distribusi kelas yang seimbang atau balanced dataset. Kondisi ini memberikan kesempatan yang sama kepada algoritma XGBoost untuk mempelajari karakteristik setiap kategori serta mengurangi kecenderungan model untuk lebih banyak memprediksi salah satu kelas. Selain itu, keseimbangan kelas membuat hasil pengukuran akurasi lebih mudah diinterpretasikan karena setiap kelas memberikan kontribusi yang sama terhadap proses pelatihan dan pengujian.

Analisis statistik deskriptif selanjutnya dilakukan untuk mengetahui karakteristik setiap variabel sebelum digunakan dalam proses pelatihan. Analisis meliputi nilai minimum, maksimum, rata-rata, dan standar deviasi. Hasilnya disajikan pada Tabel 5.

Tabel 5. Statistik Deskriptif Variabel Penelitian

Variabel	Minimum	Maksimum	Rata-rata	Standar Deviasi
Nilai Kuis	45	98	72,83	13,21
Nilai Praktikum	48	100	75,17	12,48
Kehadiran (%)	70	100	88,75	8,34
Waktu Pengerjaan (menit)	18	60	36,42	9,15
Keaktifan	2	10	6,87	2,15
Observasi Guru	2	10	7,12	2,08

Berdasarkan Tabel 5, nilai kuis memiliki rentang antara 45 dan 98 dengan rata-rata 72,83 serta standar deviasi 13,21. Rentang tersebut menunjukkan adanya variasi kemampuan siswa dalam menguasai materi teori komputer. Variasi nilai ini memberikan informasi yang diperlukan model untuk membedakan siswa berdasarkan tingkat pemahamannya.

Nilai praktikum memiliki nilai minimum 48 dan maksimum 100 dengan rata-rata 75,17 serta standar deviasi 12,48. Hasil tersebut menunjukkan bahwa secara umum siswa mampu menyelesaikan kegiatan praktikum dengan cukup baik. Adanya variasi nilai praktikum juga memperlihatkan perbedaan kemampuan siswa dalam menerapkan materi yang telah dipelajari.

Variabel kehadiran memiliki rata-rata 88,75%, sedangkan waktu pengerjaan memiliki rata-rata 36,42 menit dengan rentang 18–60 menit. Kehadiran menunjukkan keterlibatan siswa dalam mengikuti pembelajaran, sedangkan waktu pengerjaan memberikan informasi mengenai kecepatan siswa dalam menyelesaikan tugas. Namun, waktu pengerjaan tetap perlu dianalisis bersama nilai praktikum karena penyelesaian yang lebih cepat belum tentu menunjukkan pemahaman tinggi apabila hasil pekerjaannya kurang tepat.

Keaktifan memiliki rata-rata 6,87, sedangkan observasi guru memiliki rata-rata 7,12. Kedua variabel tersebut mencerminkan perilaku dan keterlibatan siswa selama kegiatan praktikum. Secara keseluruhan, seluruh variabel memiliki variasi yang dapat dimanfaatkan oleh XGBoost untuk mempelajari hubungan antara karakteristik siswa dan tingkat pemahamannya.

Sebelum dilakukan pelatihan, dataset diperiksa untuk memastikan tidak terdapat data kosong. Hasil pemeriksaan missing value ditunjukkan pada Gambar 4.

Nilai_Kuis	0
Nilai_Praktikum	0
Kehadiran	0
Waktu_Pengerjaan	0
Keaktifan	0
Observasi_Guru	0
Tingkat_Pemahaman	0

Gambar 4. Hasil Pemeriksaan Missing Value

Berdasarkan Gambar 4, nilai kuis, nilai praktikum, kehadiran, waktu pengerjaan, keaktifan, observasi guru, dan tingkat pemahaman memiliki nilai missing value sebesar 0. Dengan demikian, seluruh data dinyatakan lengkap sehingga tidak diperlukan penghapusan data maupun pengisian nilai menggunakan teknik imputasi.

Pemeriksaan data duplikat juga menunjukkan bahwa tidak terdapat data yang memiliki nilai identik pada seluruh atribut. Seluruh 120 data bersifat unik dan dapat digunakan pada tahap berikutnya. Variabel tingkat pemahaman yang semula berbentuk kategorikal kemudian ditransformasikan menggunakan label encoding. Kategori rendah dikonversi menjadi label 0, sedang menjadi label 1, dan tinggi menjadi label 2.

Dataset selanjutnya dibagi menggunakan metode Stratified Train-Test Split dengan rasio 80:20 dan random_state=42. Hasil pembagian dataset ditunjukkan pada Gambar 5.



Gambar 5. Pembagian Dataset Training dan Testing

Sebanyak 96 data atau 80% digunakan sebagai data training, sedangkan 24 data atau 20% digunakan sebagai data testing. Penggunaan metode stratified mempertahankan proporsi setiap kelas sehingga data testing terdiri atas delapan data kategori rendah, delapan data kategori sedang, dan delapan data kategori tinggi. Nilai random_state=42 digunakan agar pembagian data dapat direproduksi ketika penelitian dijalankan kembali menggunakan dataset yang sama.

3.2 Implementasi Algoritma XGBoost

Implementasi XGBoost dilakukan menggunakan 96 data training. Model mempelajari pola hubungan antara enam variabel masukan dan tiga kategori tingkat pemahaman siswa. Proses pembentukan model ditunjukkan pada Gambar 6.

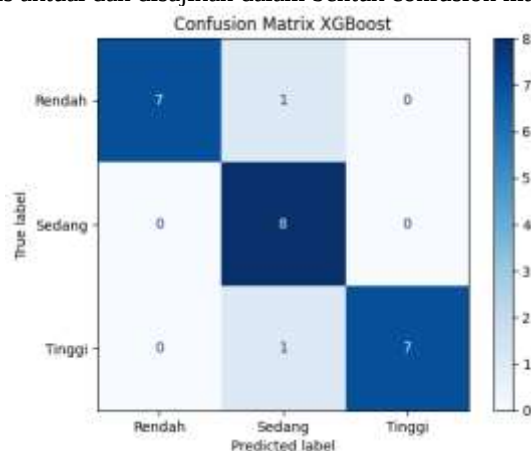


Gambar 6. Proses Pembentukan Model XGBoost

Berdasarkan Gambar 6, algoritma membentuk pohon keputusan pertama menggunakan data training, kemudian menghitung kesalahan prediksi yang dihasilkan. Kesalahan tersebut digunakan sebagai dasar untuk membangun pohon berikutnya. Setiap pohon baru berfungsi memperbaiki kesalahan dari kumpulan pohon sebelumnya. Proses tersebut dilakukan secara berulang hingga jumlah pohon yang ditetapkan terpenuhi, kemudian seluruh pohon digabungkan menjadi satu model klasifikasi.

Model dibangun menggunakan 100 decision tree dengan `learning_rate=0,1` dan `max_depth=4`. Nilai `learning_rate` tersebut membuat proses pembelajaran berlangsung secara bertahap, sedangkan kedalaman maksimum empat digunakan untuk mengendalikan kompleksitas pohon. Parameter `subsample=0,8` menyebabkan setiap pohon menggunakan 80% data training, sedangkan `colsample_bytree=0,8` menggunakan 80% fitur dalam pembentukan pohon. Kedua parameter tersebut digunakan untuk meningkatkan variasi model dan mengurangi risiko overfitting. Fungsi objektif yang digunakan adalah `multi:softmax` karena target terdiri atas tiga kelas, sedangkan `mlogloss` digunakan sebagai metrik kesalahan selama pelatihan.

Model yang telah dilatih kemudian digunakan untuk memprediksi tingkat pemahaman pada 24 data testing. Hasil prediksi dibandingkan dengan kelas aktual dan disajikan dalam bentuk confusion matrix pada Gambar 7.



Gambar 7. Confusion Matrix Algoritma XGBoost

Berdasarkan Gambar 7, model berhasil mengklasifikasikan 23 dari 24 data testing dengan benar. Pada kategori rendah, seluruh delapan data berhasil diprediksi sebagai kategori rendah. Tidak ditemukan siswa berkategori rendah yang

diprediksi sebagai sedang maupun tinggi. Hal tersebut menunjukkan bahwa karakteristik siswa dengan tingkat pemahaman rendah dapat dikenali secara konsisten oleh model.

Pada kategori sedang, sebanyak tujuh dari delapan data berhasil diprediksi dengan benar. Satu data yang sebenarnya termasuk kategori sedang diprediksi sebagai kategori tinggi. Kesalahan tersebut menunjukkan bahwa karakteristik siswa pada kategori sedang dan tinggi memiliki batas yang berdekatan. Siswa yang salah diklasifikasikan kemungkinan memiliki kombinasi nilai kuis, praktikum, kehadiran, atau observasi guru yang menyerupai pola kategori tinggi.

Pada kategori tinggi, seluruh delapan data berhasil diprediksi dengan benar. Tidak terdapat siswa kategori tinggi yang diprediksi sebagai kategori rendah maupun sedang. Secara keseluruhan, model tidak menghasilkan kesalahan klasifikasi yang jauh, seperti kategori rendah menjadi tinggi. Satu-satunya kesalahan terjadi pada dua kategori yang berdekatan, yaitu sedang dan tinggi.

Kinerja model selanjutnya dihitung menggunakan accuracy, precision, recall, dan F1-score. Hasil evaluasinya disajikan pada Tabel 6.

Tabel 6. Hasil Evaluasi Model XGBoost

Metrik	Nilai
<i>Accuracy</i>	95,83%
<i>Precision</i>	95,62%
<i>Recall</i>	95,83%
<i>F1-score</i>	95,69%

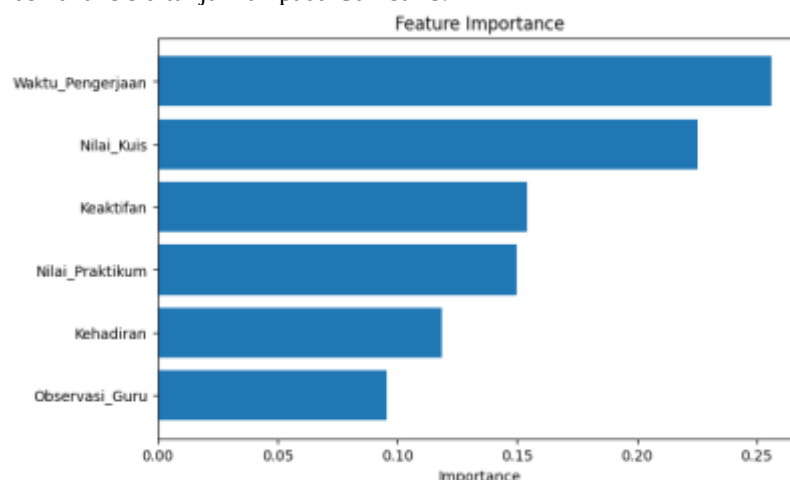
Berdasarkan Tabel 6, model memperoleh nilai accuracy sebesar 95,83%. Hasil tersebut sesuai dengan confusion matrix, yaitu 23 prediksi benar dari 24 data pengujian. Nilai tersebut menunjukkan bahwa hubungan antara enam variabel dan kategori tingkat pemahaman dapat dipelajari dengan baik oleh XGBoost.

Nilai precision sebesar 95,62% menunjukkan bahwa sebagian besar siswa yang diprediksi ke dalam suatu kategori memang berasal dari kategori tersebut. Dalam konteks evaluasi pembelajaran, nilai precision yang tinggi dapat mengurangi kemungkinan siswa ditempatkan pada kategori yang tidak sesuai dengan kondisi aktualnya.

Nilai recall sebesar 95,83% menunjukkan kemampuan model dalam mengenali hampir seluruh anggota setiap kategori. Kemampuan tersebut penting, terutama dalam mengenali siswa berkategori rendah agar siswa yang membutuhkan bantuan tidak terlewatkan. Hasil confusion matrix menunjukkan bahwa seluruh siswa kategori rendah pada data pengujian berhasil dikenali.

Nilai F1-score sebesar 95,69% menunjukkan adanya keseimbangan yang baik antara precision dan recall. Kedekatan nilai seluruh metrik memperlihatkan bahwa kinerja model tidak hanya tinggi berdasarkan jumlah prediksi benar, tetapi juga cukup seimbang dalam mengenali setiap kelas.

Analisis berikutnya dilakukan terhadap feature importance untuk mengetahui kontribusi setiap variabel dalam pembentukan model. Hasil analisis ditunjukkan pada Gambar 8.



Gambar 8. Feature Importance Algoritma XGBoost

Berdasarkan Gambar 8, nilai kuis memiliki skor kepentingan tertinggi sebesar 137, diikuti nilai praktikum 51, kehadiran 28, waktu pengerjaan 16, observasi guru 5, dan keaktifan 4. Hasil ini menunjukkan bahwa penguasaan teori dan kemampuan praktik menjadi faktor utama dalam memprediksi tingkat pemahaman siswa, sedangkan variabel lainnya berperan sebagai faktor pendukung. Skor tersebut menunjukkan kontribusi relatif fitur dalam model, bukan persentase pengaruh atau hubungan sebab-akibat.

3.3 Pembahasan Hasil Penelitian

XGBoost mampu memprediksi tingkat pemahaman siswa dengan akurasi 95,83% dan berhasil mengklasifikasikan 23 dari 24 data pengujian dengan benar. Hasil ini memperkuat penelitian terdahulu bahwa XGBoost memiliki kemampuan yang baik dalam melakukan klasifikasi dan memprediksi kinerja siswa. Berbeda dari penelitian sebelumnya, penelitian ini secara khusus mengklasifikasikan tingkat pemahaman siswa pada praktikum laboratorium komputer berdasarkan enam indikator akademik dan perilaku belajar. Nilai kuis menjadi variabel paling dominan, diikuti nilai praktikum. Meskipun menghasilkan kinerja yang tinggi, penelitian ini terbatas pada 120 data dari satu sekolah sehingga diperlukan pengujian dengan dataset yang lebih besar dan perbandingan dengan algoritma lain.

4. KESIMPULAN

Berdasarkan hasil penelitian, tingkat pemahaman siswa dalam mengikuti praktikum laboratorium komputer dapat diklasifikasikan ke dalam tiga kategori, yaitu rendah, sedang, dan tinggi, berdasarkan enam variabel yang meliputi nilai kuis, nilai praktikum, kehadiran, waktu pengerjaan, keaktifan, dan observasi guru. Penerapan algoritma XGBoost dilakukan melalui tahapan identifikasi masalah, pengumpulan data, preprocessing, pengodean label, pembagian dataset menggunakan Stratified Train-Test Split, pelatihan model, pengujian, dan evaluasi. Dari keseluruhan 120 data siswa, sebanyak 96 data digunakan sebagai data training dan 24 data sebagai data testing. Hasil pengujian menunjukkan bahwa model mampu mengklasifikasikan 23 dari 24 data pengujian dengan benar. Model memperoleh nilai accuracy sebesar 95,83%, precision sebesar 95,62%, recall sebesar 95,83%, dan F1-score sebesar 95,69%. Hasil tersebut menunjukkan bahwa XGBoost memiliki tingkat ketepatan yang sangat baik dalam memprediksi tingkat pemahaman siswa berdasarkan data praktikum laboratorium komputer. Analisis feature importance menunjukkan bahwa nilai kuis menjadi variabel paling berpengaruh, diikuti oleh nilai praktikum, kehadiran, waktu pengerjaan, observasi guru, dan keaktifan. Temuan ini mengindikasikan bahwa penguasaan materi secara konseptual dan kemampuan menerapkannya dalam kegiatan praktikum memiliki peran utama dalam menentukan tingkat pemahaman siswa. Dengan demikian, algoritma XGBoost dapat dimanfaatkan sebagai alat bantu bagi guru dalam mengidentifikasi tingkat pemahaman siswa secara lebih cepat, objektif, dan berbasis data sehingga pemberian bimbingan, remedial, maupun pengayaan dapat dilakukan sesuai kebutuhan siswa. Meskipun demikian, penerapannya tetap perlu disertai pertimbangan dan evaluasi langsung dari guru.

REFERENCES

- [1] M. O. Capao, V. S. Rosales, J. V. H. Leuterio, Q. Faith, and R. B. Requino, "Mobile Simulation-Based Learning (MSBL): An Integrated Basic Electronics," *J. Tech. Educ. Train.*, vol. 17, no. 3, pp. 64–78, 2025, doi: <https://doi.org/10.30880/jtet.2025.17.03.005>.
- [2] Ç. Atmaca, "Online teaching practicum procedures during the COVID-19 pandemic in Turkish EFL context," vol. 20, no. 4, pp. 352–369, 2023, doi: <https://doi.org/10.1177/20427530231156677>.
- [3] H. Subandiyah, R. Nasrullah, R. Ramadhan, H. Supratno, R. P. Raharjo, and F. Lukman, "The impact of differentiated instruction on student engagement and achievement in Indonesian language learning," *Cogent Educ.*, vol. 12, no. 1, p., 2025, doi: <https://doi.org/10.1080/2331186X.2025.2516378>.
- [4] A. Aydınlar *et al.*, "Awareness and level of digital literacy among students receiving health-based education," *BMC Med. Educ.*, vol. 24, no. 38, pp. 1–13, 2024, doi: [10.1186/s12909-024-05025-w](https://doi.org/10.1186/s12909-024-05025-w).
- [5] C. Lu and M. Cutumisu, "Online engagement and performance on formative assessments mediate the relationship between attendance and course performance," *Int. J. Educ. Technol. High. Educ.*, vol. 19, no. 2, 2022, doi: <https://doi.org/10.1186/s41239-021-00307-5>.
- [6] W. Liu, Z. Chen, and Y. Hu, "XGBoost algorithm-based prediction of safety assessment for pipelines," *Int. J. Press. Vessel. Pip.*, vol. 197, p. 104655, 2022, doi: <https://doi.org/10.1016/j.ijpvp.2022.104655>.
- [7] H. Liu, X. Chen, and X. Liu, "Factors influencing secondary school students' reading literacy: An analysis based on XGBoost and SHAP methods," *Front. Psychol.*, vol. 13, no. 948612, 2022, doi: [10.3389/fpsyg.2022.948612](https://doi.org/10.3389/fpsyg.2022.948612).
- [8] M. Yang, Z. Li, and S. Liu, "A Machine Learning Based Framework for Predictive School Management Using Student and Faculty Analytics," *Sci. Rep.*, vol. 16, no. 16224, pp. 1–31, 2026, doi: <https://doi.org/10.1038/s41598-026-47278-z>.
- [9] B. Cheng, Y. Liu, and Y. Jia, "Evaluation of students' performance during the academic period using the XG-Boost Classifier-Enhanced AEO hybrid model," *Expert Syst. Appl.*, vol. 238, no. PD, p. 122136, 2024, doi: <https://doi.org/10.1016/j.eswa.2023.122136>.
- [10] H. Guo, "An Interpretable Cross-Regional Prediction Study of University Students' Mental Health Based on XGBoost and SHAP," *IEEE Access*, vol. 14, pp. 48025–48039, 2026, doi: <https://doi.org/10.1109/ACCESS.2026.3676575>.
- [11] R. Natras and B. Soja, "Ensemble Machine Learning of Random Forest, AdaBoost and XGBoost for Vertical Total Electron Content Forecasting," *Remote Sens.*, vol. 14, pp. 1–34, 2022, doi: <https://doi.org/10.3390/rs14153547>.
- [12] S. Hakkal and A. A. Lahcen, "Computers and Education: Artificial Intelligence XGBoost To Enhance Learner Performance Prediction," *Comput. Educ. Artif. Intell.*, vol. 7, p. 100254, 2024, doi: <https://doi.org/10.1016/j.artint.2024.100254>.



- <https://doi.org/10.1016/j.caeai.2024.100254>.
- [13] A. Villar, C. Robledo, and V. De Andrade, "Supervised machine learning algorithms for predicting student dropout and academic success : a comparative study," *Discov. Artif. Intell.*, vol. 4, no. 2, 2024, doi: <https://doi.org/10.1007/s44163-023-00079-z>.
- [14] L. R. Pelima, Y. Sukmana, and Y. Rosmansyah, "Predicting University Student Graduation Using Academic Performance and Machine Learning : A Systematic Literature Review," *IEEE Access*, vol. 12, no. January, pp. 23451–23465, 2024, doi: <https://doi.org/10.1109/ACCESS.2024.3361479>.
- [15] M. Chaudhry *et al.*, "A Systematic Literature Review on Identifying Patterns Using Unsupervised Clustering Algorithms : A Data Mining Perspective," *Symmetry (Basel)*, vol. 15, no. 1679, pp. 1–44, 2023, doi: <https://doi.org/10.3390/sym15091679>.
- [16] M. Bilal, G. Ali, M. W. Iqbal, M. Anwar, M. Sheraz, and A. Malik, "Auto-Prep : Efficient and Automated Data Preprocessing Pipeline," *IEEE Access*, vol. 10, no. August, pp. 107764–107784, 2022, doi: <https://doi.org/10.1109/ACCESS.2022.3198662>.