

**Optimasi Model Prediksi Hasil Belajar Siswa dalam
Pembelajaran IPA Berbasis Media Sosial Menggunakan XGBoost
dan Bayesian Optimization**

SKRIPSI

DISUSUN OLEH

M Chairil Fawwaz

NPM. 2109020060



**PROGRAM STUDI TEKNOLOGI INFORMASI
FAKULTAS ILMU KOMPUTER DAN TEKNOLOGI INFORMASI
UNIVERSITAS MUHAMMADIYAH SUMATERA UTARA
MEDAN
2025**

**Optimasi Model Prediksi Hasil Belajar Siswa dalam
Pembelajaran IPA Berbasis Media Sosial Menggunakan XGBoost
dan Bayesian Optimization**

SKRIPSI

**Diajukan sebagai salah satu syarat untuk memperoleh gelar Sarjana Komputer
(S.Kom) dalam Program Studi Teknologi Informasi pada Fakultas Ilmu
Komputer dan Teknologi Informasi, Universitas Muhammadiyah Sumatera Utara**

M Chairil Fawwaz

NPM. 2109020060

**PROGRAM STUDI TEKNOLOGI INFORMASI
FAKULTAS ILMU KOMPUTER DAN TEKNOLOGI INFORMASI
UNIVERSITAS MUHAMMADIYAH SUMATERA UTARA
MEDAN
2025**

LEMBAR PENGESAHAN

Judul Skripsi : OPTIMASI MODEL PREDIKSI HASIL BELAJAR
SISWA DALAM PEMBELAJARAN IPA BERBASIS
MEDIA SOSIAL MENGGUNAKAN XGBOOST DAN
BAYESIAN OPTIMIZATION

Nama Mahasiswa : M CHAIRIL FAWWAZ

NPM : 2109020060

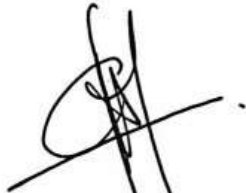
Program Studi : TEKNOLOGI INFORMASI

Menyetujui
Komisi Pembimbing



(Zuli Agustina Gultom, M.Si)
NIDN. 0130089003

Ketua Program Studi



(Fatmasari Hutagalung, S.Kom. M.Kom.)
NIDN. 0117088902

Dekan



(Dr. Al-Khowarizmi, S.Kom., M.Kom.)
NIDN. 0127099201

PERNYATAAN ORISINALITAS

Optimasi Model Prediksi Hasil Belajar Siswa dalam Pembelajaran IPA Berbasis Media Sosial Menggunakan XGBoost dan Bayesian Optimization

SKRIPSI

Saya menyatakan bahwa karya tulis ini adalah hasil karya sendiri, kecuali beberapa kutipan dan ringkasan yang masing-masing disebutkan sumbernya.



**PERNYATAAN PERSETUJUAN PUBLIKASI
KARYA ILMIAH UNTUK KEPENTINGAN
AKADEMIS**

Sebagai sivitas akademika Universitas Muhammadiyah Sumatera Utara, saya bertanda tangan dibawah ini:

Nama	: M Chairil Fawwaz
NPM	2109020060
Program Studi	: Teknologi Informasi
Karya Ilmiah	: Skripsi

Demi pengembangan ilmu pengetahuan, menyetujui untuk memberikan kepada Universitas Muhammadiyah Sumatera Utara Hak Bedas Royalti Non-Eksekutif (*Non-Exclusive Royalty free Right*) atas penelitian skripsi saya yang berjudul:

**Optimasi Model Prediksi Hasil Belajar Siswa dalam
Pembelajaran IPA Berbasis Media Sosial Menggunakan XGBoost
dan Bayesian Optimization**

Beserta perangkat yang ada (jika diperlukan). Dengan Hak Bebas Royalti Non-Eksekutif ini, Universitas Muhammadiyah Sumatera Utara berhak menyimpan, mengalih media, memformat, mengelola dalam bentuk database, merawat dan mempublikasikan Skripsi saya ini tanpa meminta izin dari saya selama tetap mencantumkan nama saya sebagai penulis dan sebagai pemegang dan atau sebagai pemilik hak cipta.

Demikian pernyataan ini dibuat dengan sebenarnya.

Medan, Agustus 2025

Yang membuat pernyataan



M Chairil Fawwaz

NPM. 2109020060

RIWAYAT HIDUP

DATA PRIBADI

Nama Lengkap	: M Chairil Fawwaz
Tempat dan Tanggal Lahir	: Medan, 24 Januari 2004
Alamat Rumah	: Jl.M.Yakub No 117
Telepon/Faks/HP	081396602383
E-mail	: chairilfawas@gmail.com
Instansi Tempat Kerja	: -
Alamat Kantor	: -

DATA PENDIDIKAN

SD	: SD Swasta Taman Harapan	TAMAT: 2015
SMP	: MTSN 2 MEDAN	TAMAT: 2018
SMA	: MAN 1 MEDAN	TAMAT: 2021

KATA PENGANTAR



Dengan memanjatkan puji syukur kehadiran Tuhan Yang Maha Esa, penulis akhirnya dapat menyelesaikan penyusunan skripsi yang berjudul “Optimasi Model Prediksi Hasil Belajar Siswa dalam Pembelajaran IPA Berbasis Media Sosial Menggunakan XGBoost dan Bayesian Optimization” sebagai salah satu bentuk tanggung jawab akademik dan syarat dalam menyelesaikan studi di Program Studi Teknologi Informasi, Fakultas Ilmu computer dan Teknologi Informasi, Universitas Muhammadiyah Sumatera Utara. Penulis tentunya berterima kasih kepada berbagai pihak dalam dukungan serta doa dalam penyelesaian skripsi. Penulis juga mengucapkan terima kasih kepada:

1. Bapak Prof. Dr. Agussani, M.AP., Rektor Universitas Muhammadiyah Sumatera Utara (UMSU)
2. Bapak Dr. Al-Khowarizmi, S.Kom., M.Kom. Dekan Fakultas Ilmu Komputer dan Teknologi Informasi (FIKTI) UMSU.
3. Ibu Fatma Sari Hutagalung S.Kom., M.Kom. Ketua Program Studi Teknologi Informasi
4. Bapak Mhd. Basri, S.Si, M.Kom Sekretaris Program Studi Teknologi Informasi
5. Ibu Zuli Agustina Gultom, M.Si. Pembimbing Skripsi
6. Ayah tercinta (Alm.) Bapak Chairunnas, meskipun engkau telah tiada, semangat dan doamu selalu hidup dalam setiap langkahku. Terima kasih atas cinta, perjuangan, dan nilai-nilai hidup yang kau wariskan. Semoga Allah SWT menempatkanmu di tempat terbaik di sisi-Nya.
7. Ibunda tercinta saya, Ibu Siti Sundari yang selalu mendoakan saya setiap hari, memberi semangat, dan menjadi sumber kekuatan dalam setiap langkah penulis.

8. abang kandung penulis yang selalu memberikan semangat dan dukungan, baik secara langsung maupun tidak langsung. Terima kasih atas kebersamaan, tawa, serta dorongan yang membuat penulis tidak merasa sendiri dalam proses ini.
9. Semua pihak yang terlibat langsung ataupun tidak langsung yang tidak dapat penulis ucapkan satu-persatu yang telah membantu penyelesaian skripsi ini

Optimasi Model Prediksi Hasil Belajar Siswa dalam Pembelajaran IPA Berbasis Media Sosial Menggunakan XGBoost dan Bayesian Optimization

ABSTRAK

Penelitian ini memiliki tujuan untuk menciptakan model yang dapat memprediksi hasil belajar siswa dengan memanfaatkan teknik Extreme Gradient Boosting (XGBoost) yang dioptimalkan menggunakan Bayesian Optimization. Penelitian ini terfokus pada pembelajaran Ilmu Pengetahuan Alam (IPA) melalui media sosial, dan efektivitasnya diukur dengan membandingkan hasil pre-test, kuis, dan post-test. Data dalam penelitian ini diperoleh dari 30 siswa SMP yang berperan sebagai responden, dengan variabel input berupa nilai pre-test dan kuis, sementara variabel outputnya adalah nilai post-test.

Tahapan dalam penelitian ini meliputi pra-pemrosesan data, yang mencakup penanganan nilai yang hilang dengan metode imputasi rata-rata, pembagian data menjadi set pelatihan dan set pengujian dengan rasio 80:20, pelatihan model XGBoost, optimasi hyperparameter melalui Bayesian Optimization, dan evaluasi model menggunakan metrik MAE (Mean Absolute Error), RMSE (Root Mean Square Error), dan R^2 (Coefficient of Determination).

Kata Kunci: Prediksi hasil belajar, XGBoost, Bayesian Optimization, pembelajaran Berbasis media sosial.

Optimizing a Prediction Model for Student Learning Outcomes in Social Media-Based Science Learning Using XGBoost and Bayesian Optimization

ABSTRACT

This study aims to develop a model capable of predicting students' learning outcomes by utilizing the Extreme Gradient Boosting (XGBoost) technique optimized with Bayesian Optimization. The research focuses on Science (IPA) learning through social media, with its effectiveness measured by comparing pre-test, quiz, and post-test results. The data were obtained from 30 junior high school students who served as respondents, with the input variables consisting of pre-test and quiz scores, while the output variable was the post-test score.

The stages of this research include data preprocessing, which involves handling missing values using the mean imputation method, splitting the data into training and testing sets with an 80:20 ratio, training the XGBoost model, optimizing hyperparameters using Bayesian Optimization, and evaluating the model using MAE (Mean Absolute Error), RMSE (Root Mean Square Error), and R^2 (Coefficient of Determination) metrics.

Keywords: Learning outcome prediction, XGBoost, Bayesian Optimization, social media-based learning.

DAFTAR ISI

LEMBAR PENGSAHAN	i
PENYATAAN ORISINALITAS	ii
PENYATAAN PERSETUJUAN PUBLIKASI	iii
RIWAYAT HIDUP	iv
KATA PENGANTAR	v
ABSTRAK.....	vi
ABSTRACT	vii
DAFTAR ISI	viii
DAFTAR TABEL	ix
DAFTAR GAMBAR	x
BAB I PENDAHULUAN	1
1.1 Latar Belakang Masalah	1
1.2 Rumusan Masalah.....	2
1.3 Batasan Masalah	2
1.4 Tujuan Penelitian	3
1.5 Manfaat Penelitian.....	3
BAB II LANDASAN TEORI.....	5
2.1 Hasil Belajar Siswa.....	5
2.2 Media Sosial	6
2.2.1 Dampak Pendidikan Media Sosial.....	7
2.3 Faktor Yang Mempengaruhi Keberhasilan Pembelajaran Online	8
2.4 IPA (ILMU PENGETAHUAN ALAM)	8
2.4.1 visualisasi Konsep IPA dalam Pembelajaran Berbasis Digital....	11

2.4.2 Cuplikan Konten Pembelajaran IPA di Media Sosial.....	11
2.5 Machine Learning.....	12
2.5.1 Tahapan dalam Machine Learning	13
2.5.1.1 Ide / Masalah Bisnis.....	13
2.5.1.2 Pengumpulan Data.....	13
2.5.1.3 Pemilihan model	14
2.5.1.4 Pra-pemrosesan Data.....	14
2.5.1.5 Melatih Model.....	15
2.5.1.6 Evaluasi Model	15
2.6 Jenis jenis machine learning	16
2.7 XGboost	17
2.8 Bayesian optimization.....	19
2.9 Python.....	20
2.10 Metode Evaluasi	20
2.10.1 Root Mean Squared Error (RMSE)	20
2.10.2 Mean Absolute Error (MAE).....	21
2.10.3 Coefficient of Determination (R ²).....	21
2.11 Penelitian terdahulu	22
BAB III METODE PENELITIAN	26
3.1 Jenis Penelitian	26
3.2 Metode Penelitian	26
3.3 Teknik Pengumpulan Data	26
3.4 Alat Bantu Penelitian.....	27
3.5 Alur penelitian	28

BAB IV HASIL DAN PEMBAHASAN.....	32
4.1 Pengumpulan data.....	32
4.2 Pra-pemrosesan data	38
4.2.1 Import Library	32
4.2.2 Import Dataset	33
4.2.3 Pengecekan dan Penanganan Nilai Kosong	33
4.2.4 Pemisahan Fitur dan Target.....	34
4.2.5 Split Data.....	34
4.2.6 Deskripsi Dataset.....	34
4.3 Pelatihan model XGBoost	36
4.4 Optimasi dengan Bayesian Optimization.....	37
4.4.1 Fungsi Evaluasi dan Ruang Parameter	37
4.5 Uji Validasi dan Reliabilitas.....	38
4.6 Evaluasi Model	40
4.7 Visualisasi Hasil	43
4.7.1 Grafik Hasil Prediksi	43
4.7.2 Riwayat Bayesian Optimization	44
4.8 Analisis Hasil.....	45
BAB V KESIMPULAN.....	47
5.1 Kesimpulan.....	47
5.2 Saran	48
DAFTAR PUSTAKA	49

DAFTAR TABEL

Tabel 2. 1 Alur Penelitian	21
Tabel 4. 1 Dataset Awal.....	35
Tabel 4. 2 Dataset Hasil Pra-Pemrosesan	35
Tabel 4. 3 Tabel uji validasi	38
Tabel 4. 4 Tabel Reliabilitas.....	39
Tabel 4. 5 Data Aktual dan Prediksi	40

DAFTAR GAMBAR

Gambar 2. 1 Organ pernapasan manusia	9
Gambar 2. 2 Bumi dan Tata Surya	10
Gambar 2. 3 Ekosistem	10
Gambar 2. 4 Perbedaan antara pembelajaran konvensional dan pembelajaran berbasis media sosial	11
Gambar 2. 5 Video Pembelajaran sistem pernapasan manusia	12
Gambar 2. 6 Visualisasi XGBoost	19
Gambar 3. 1 Alur Peneletian	29
Gambar 4. 1 Import Library	32
Gambar 4. 2 Import Dataset.....	33
Gambar 4. 3 Mengisi Nilai Kosong.....	33
Gambar 4. 4 Memisahkan fitur prediksi dan target	34
Gambar 4. 5 Pembagian Data Training Dan Testing.....	34
Gambar 4. 6 Pelatihan Model Xgboost	36
Gambar 4. 7 Optimasi Bayesian Optimization	37
Gambar 4. 8 Fungsi Evaluasi dan Ruang Parameter.....	37
Gambar 4. 9 Visualisasi hasil prediksi dan optimasi.....	38
Gambar 4. 10 Hasil Aktual Dengan Prediksi.....	43
Gambar 4. 11 Hasil riwayat Bayesian Optimization	44

BAB I

PENDAHULUAN

1.1 Latar Belakang Masalah

Media sosial telah menjadi bagian penting dari kehidupan sehari-hari di era komputer dan internet saat ini, terutama bagi siswa. Situs seperti Facebook, Instagram, dan YouTube tidak hanya berfungsi sebagai alat komunikasi tetapi juga membantu siswa belajar. Penggunaan media sosial dalam pendidikan dapat meningkatkan motivasi dan keterlibatan siswa.

Meskipun ada potensi positif, masalah utamanya adalah bagaimana memprediksi hasil belajar siswa dalam konteks pembelajaran berbasis media sosial. Dengan jumlah data yang meningkat yang berasal dari interaksi siswa di media sosial, peneliti memerlukan alat yang dapat menganalisis dan menginterpretasikan data tersebut dengan cara yang akan meningkatkan hasil belajar siswa. Model prediksi hasil belajar siswa menjadi penting untuk mengidentifikasi siswa yang berisiko mengalami kesulitan akademik dan merancang intervensi yang tepat.

Fokus penelitian ini adalah IPA (Ilmu Pengetahuan Alam) karena konten pembelajaran IPA yang beredar di media sosial seperti di platform YouTube, Instagram, Tiktok cenderung lebih menarik secara visual dan interaktif. Banyak video IPA menampilkan animasi, simulasi, dan ilustrasi yang menarik, yang membantu siswa memahami konsep abstrak. Hal ini menunjukkan bahwa IPA adalah mata pelajaran yang dapat dikembangkan dengan menggunakan pendekatan pembelajaran berbasis media sosial yang dapat meningkatkan pemahaman dan hasil belajar siswa.

Algoritma *Machine Learning XGBoost (Extreme Gradient Boosting)* dipilih dalam penelitian ini karena telah terbukti efektif dalam berbagai Aktivitas prediksi, termasuk di bidang pendidikan. Ciri khas utama *XGBoost* adalah kemampuannya dalam menangani data berukuran besar dan kompleks dengan tingkat akurasi yang tinggi. Meskipun

penggunaan algoritma ini dalam konteks prediksi hasil belajar siswa masih belum sebanyak algoritma klasik seperti *Naive Bayes* atau *Decision Tree*, penelitian ini akan mengintegrasikannya dengan *Bayesian Optimization* sebagai metode pencarian *hyperparameter* terbaik secara otomatis, guna meningkatkan performa model secara optimal.

Potensi untuk meningkatkan kualitas pendidikan bergantung pada penyelesaian masalah ini. Pendidik dapat meningkatkan hasil belajar secara keseluruhan dengan memberikan intervensi cepat kepada siswa yang membutuhkan dengan menggunakan teknologi dan data yang ada saat ini. Diharapkan penelitian ini akan membantu pendidik memahami cara terbaik untuk menggunakan media sosial untuk membantu siswa.

1.2 Rumusan Masalah

Berdasarkan latar belakang yang telah dikemukakan di atas, terdapat empat pokok masalah pada penelitian ini yaitu:

1. Bagaimana karakteristik hasil belajar siswa dalam pembelajaran IPA berbasis media sosial berdasarkan data yang tersedia?
2. Seberapa akurat model prediksi hasil belajar siswa yang dikembangkan menggunakan algoritma *XGBoost* dalam konteks pembelajaran berbasis media sosial?
3. Bagaimana teknik *Bayesian Optimization* dapat meningkatkan kinerja model prediksi hasil belajar siswa?
4. Apa saja faktor-faktor yang mempengaruhi hasil belajar siswa dalam pembelajaran IPA berbasis media sosial?

1.3 Batasan Masalah

Dalam penelitian ini memiliki beberapa batasan untuk di teliti, berikut batasan masalah penelitian yang dilakukan yaitu :

1. Penelitian ini akan difokuskan pada siswa di tingkat sekolah menengah yang mengikuti pembelajaran IPA.
2. Penelitian ini akan membatasi analisis pada konten pembelajaran yang tersedia di platform media sosial seperti YouTube dan TikTok.
3. Model prediksi yang akan dikembangkan hanya menggunakan algoritma *XGBoost* dan teknik *Bayesian Optimization* untuk optimasi hyperparameter.
4. Data yang akan dianalisis mencakup nilai kuis, interaksi siswa dengan konten media sosial, dan hasil ujian akhir (post-test) siswa.

1.4 Tujuan Penelitian

Berdasarkan rumusan dan batasan masalah di atas, maka tujuan dalam penelitian ini adalah untuk:

1. Mengidentifikasi pengaruh penggunaan media sosial terhadap hasil belajar siswa dalam pembelajaran IPA.
2. Mengembangkan model prediksi hasil belajar siswa menggunakan algoritma *XGBoost* untuk meningkatkan akurasi prediksi.
3. Menerapkan teknik *Bayesian Optimization* untuk mengoptimalkan hyperparameter model prediksi yang dikembangkan.
4. Menganalisis faktor-faktor yang mempengaruhi hasil belajar siswa dalam konteks pembelajaran berbasis media sosial.

1.5 Manfaat Penelitian

Berdasarkan tujuan di atas, maka manfaat yang diharapkan dalam penelitian ini adalah sebagai berikut:

1. Bagi Pendidik: Penelitian ini diharapkan dapat memberikan wawasan bagi pendidik tentang bagaimana media sosial dapat dimanfaatkan secara efektif dalam proses pembelajaran, serta memberikan alat untuk memprediksi kinerja siswa.

2. Bagi Siswa: Dengan adanya model prediksi yang akurat, siswa yang berisiko mengalami kesulitan akademik dapat diidentifikasi lebih awal, sehingga intervensi yang tepat dapat diberikan untuk meningkatkan hasil belajar mereka.
3. Bagi Peneliti: Penelitian ini dapat menjadi referensi bagi penelitian selanjutnya dalam bidang pendidikan dan teknologi, khususnya dalam penggunaan algoritma machine learning untuk analisis data pendidikan.
4. Bagi Pengembangan Kebijakan Pendidikan: Hasil penelitian ini dapat memberikan masukan bagi pengembangan kebijakan pendidikan yang lebih responsif terhadap kebutuhan siswa di era digital.

BAB II

LANDASAN TEORI

2.1 Hasi Belajar Siswa

Hasil belajar siswa adalah sebuah interaksi antara proses pembelajaran dan tindakan pengajaran. Oleh sebab itu, untuk mendapatkan hasil belajar yang terbaik, terdapat banyak faktor yang berperan di dalamnya. Beberapa di antara faktor tersebut adalah kurikulum, tenaga pendidik, orang tua, lingkungan pendidikan, serta aktivitas siswa itu sendiri. Aktivitas belajar mempunyai peranan penting dalam mencapai kesuksesan pendidikan. Di sekolah, terdapat berbagai aktivitas yang berkaitan dengan proses belajar, seperti menulis, mencatat, mengamati, membaca, mengingat, berpikir, berlatih, dan lain-lain. Siswa yang sedang belajar pasti melaksanakan berbagai kegiatan untuk membantu mencapai tujuan pembelajaran yang diinginkan. Aktivitas yang dimaksud menekankan pada siswa, karena melalui aktivitas siswa dalam proses pembelajaran, terciptalah lingkungan belajar yang aktif. Siswa adalah individu yang turut serta dalam proses pembelajaran secara langsung. (Aisa, 2021).

Dalam alur pembelajaran terdapat tiga domain hal yang perlu diperhatikan adalah aspek kognitif, afektif, dan psikomotorik. Kognitif adalah kemampuan yang berkaitan dengan pengetahuan, pemikiran, atau logika. Afektif adalah kemampuan yang lebih mengedepankan perasaan, emosi, dan berbagai reaksi yang berbeda dibandingkan dengan penalaran. Hal ini mencakup kategori penerimaan, partisipasi, penilaian atau penentuan sikap, organisasi, serta pembentukan pola hidup. Psikomotorik adalah kemampuan yang menekankan keterampilan fisik, yang meliputi kategori seperti persepsi, kesiapan, gerakan terarah, gerakan yang sudah terbiasa, gerakan yang kompleks, penyesuaian pola gerakan, serta aspek kreativitas. (Siregar et al., 2023).

Faktor yang mempengaruhi motivasi belajar seorang siswa dapat dilihat dari faktor dorongan internal hasrat pribadi untuk belajar dan keinginan akan hal keberhasilan dan

dorongan kebutuhan belajar, pencapaian cita cita. Sedangkan faktor pengaruh eksternal adalah adanya penghargaan, lingkungan belajar supportif dan kegiatan belajar yang memotivasi. Pada dasarnya motivasi hasil belajar siswa adalah dorongan internal dan eksternal kepada siswa untuk mengubah tingkah laku mereka, biasanya dengan beberapa indikator atau elemen yang mendukung.

Penilaian diartikan sebagai sebuah cara untuk mengidentifikasi sejumlah pengetahuan yang disimpan dalam pikiran peserta didik, penilaian ini juga untuk menyediakan informasi atau umpan balik (feedback) baik kepada guru maupun peserta didik untuk memandu cara mengajar dan belajar untuk mencapai tujuan bersama. Penilaian skor hasil *pre-test* dan *post-test* dianalisis secara deskriptif untuk melihat skor rata rata dan deviasi standar.

Dalam kegiatan pembelajaran terdiri dari dua bagian, yaitu *pre-test* dan *post-test*. *Pre-test* merupakan tes yang dilaksanakan sebelum proses pengajaran dimulai, dengan tujuan untuk mengukur tingkat pemahaman siswa terhadap materi ajar yang akan disampaikan. Dalam situasi ini, fungsi *pre-test* adalah untuk menilai sejauh mana efektivitas proses pengajaran. *Post-test* adalah evaluasi yang dilaksanakan di akhir setiap program pembelajaran. Tujuan dari *post-test* adalah untuk mengevaluasi sejauh mana pencapaian siswa terhadap materi ajar (baik pengetahuan maupun keterampilan) setelah mengikuti suatu proses pembelajaran. (Purwanto dalam Siregar et al., 2023).

2.2 MEDIA SOSIAL

Media sosial adalah jaringan sistem komunikasi yang difasilitasi komputer yang dijalankan dengan teknologi web 2.0 untuk memfasilitasi pembuatan dan peningkatan situs jejaring sosial di lingkungan digital (Ravichandran dalam Helal Mridha, 2023). Dalam konteks pembelajaran media sosial tidak hanya membantu orang berkomunikasi, tetapi juga menjadi wadah interaktif untuk siswa. Mereka dapat menonton video pembelajaran, berkomentar, berbagi, dan mengakses informasi kapan saja dan di mana saja. Siswa dapat memperoleh pelajaran dalam bentuk yang lebih visual, sederhana, dan mudah dipahami

melalui media sosial seperti YouTube, Tiktok, dll. Pembelajaran siswa tidak lagi terbatas oleh ruang dan waktu seperti pembelajaran konvensional, sehingga lebih fleksibel. Siswa, misalnya, dapat berpartisipasi dalam diskusi di kolom komentar, menonton ulang video pelajaran, atau menjawab kuis yang termasuk dalam konten. Selain itu, fitur media sosial yang interaktif, berkolaborasi, dan berbasis komunitas memungkinkan siswa untuk berpartisipasi secara aktif dalam proses pendidikan, bukan hanya sebagai penerima informasi tetapi juga sebagai peserta yang dapat menanggapi dan berkontribusi.

Dengan perkembangan Teknologi informasi dan Komunikasi, dunia bergerak maju menuju penggunaan fasilitas teknologi modern secara massif di setiap bidang kehidupan. Media sosial juga merupakan salah satu fasilitas modern yang banyak digunakan dalam berbagai aspek. Dunia Pendidikan, dalam hal ini, tidak ketinggalan untuk memanfaatkan media sosial dalam kegiatan akademik (Redecker dalam Helal Mridha, 2023). Istilah “media sosial” sekarang digambarkan sebagai jaringan untuk mentebarkan pengetahuan antara komunitas dan pelajar. (Al-Rahmi & Zeki dalam Helal Mridha, 2023).

2.2.1 Dampak Pendidikan Media Sosial

Media sosial membawa beberapa risiko. Dua jenis dampak yang dimaksud adalah dampak positif dan negative:

- 1) Dampak positif Media sosial menjadi semakin penting dalam dunia pendidikan saat ini. Kemudahan untuk mendapatkan informasi dan materi pembelajaran adalah manfaatnya. Sekarang, siswa dan mahasiswa dapat mencari referensi, mengikuti kursus online, dan berbicara dengan teman atau pendidik melalui platform seperti YouTube, Instagram, dan TikTok, yang sudah mulai digunakan sebagai media pembelajaran. Ini mendorong siswa untuk menjadi lebih aktif, kreatif, dan mandiri dalam pembelajaran mereka.
- 2) Sebaliknya, penggunaan media sosial juga dapat berdampak negatif, terutama jika digunakan dengan tidak bijak. Terlalu banyak distraksi dari konten yang tidak

relevan dapat menyebabkan penurunan fokus belajar, dan penggunaan media sosial yang berlebihan juga dapat menyebabkan kecanduan, yang mengganggu waktu belajar dan aktivitas produktif lainnya.

2.3 Faktor Yang Mempengaruhi Keberhasilan Pembelajaran Online

- 1) Motivasi belajar yang tinggi: Jika kita memiliki motivasi untuk belajar, kita akan sangat termotivasi untuk belajar. Jika kita tidak memilikinya, kita kemungkinan besar akan terjebak dalam kemalasan yang tidak berujung, yang menyebabkan rasa semangat untuk belajar berkurang.
- 2) Konsistensi dalam pembelajaran: Menjadi konsisten dalam belajar secara teratur adalah penting untuk keberhasilan dan kesuksesan dalam pembelajaran. Ini membantu Anda menjadi lebih memahami apa yang Anda pelajari dan meningkatkan kualitas belajar Anda.
- 3) Aspek penting lainnya yang berkontribusi pada keberhasilan pembelajaran adalah partisipasi dalam proses belajar. Murid tidak hanya diharapkan untuk mendengarkan atau membaca, tetapi juga harus berperan aktif dalam pembelajaran, seperti mengajukan pertanyaan dan berinteraksi saat berdiskusi dengan teman dan guru. Dengan memperhatikan berbagai faktor yang telah dijelaskan, tujuannya adalah untuk meraih keberhasilan dalam belajar. Penting untuk diingat bahwa belajar bukan hanya sebatas menghafal informasi, tetapi juga memerlukan pemahaman dan penerapan dalam kehidupan sehari-hari. Oleh karena itu, kita perlu memahami dan menguasai materi pembelajaran agar dapat membantu dalam menghadapi tantangan di masa mendatang. (Sudiantini, 2023).

2.4 IPA (ILMU PENGETAHUAN ALAM)

Ilmu Pengetahuan Alam (IPA) adalah salah satu pelajaran yang diajarkan di sekolah. Ilmu Pengetahuan Alam bertujuan mengembangkan kemampuan siswa dalam memahami gejala-gejala alam, baik yang timbul dengan sendirinya maupun timbul akibat

campur tangan manusia itu sendiri, memahami konsep dan teori serta berlatih dan mengatasi isu-isu IPA yang muncul di dalam lingkungan masyarakat. Cakupan pelajaran IPA mencakup permasalahan-permasalahan alam yang terjadi di sekitar siswa hingga lingkungan yang paling jauh.(Suendarti & Hasbullah, 2020).

Untuk memahami konsep IPA diperlukan kemampuan berpikir logis, analitis, dan kemampuan untuk mengaitkan teori dengan kehidupan nyata. Siswa harus memahami sebab-akibat, melakukan pengamatan, dan menarik kesimpulan dari gejala alam. Konsep-konsep dalam IPA seringkali abstrak, seperti gaya, energi, atau sistem organ tubuh, sehingga tidak cukup dijelaskan secara lisan atau melalui teks. Media yang mendukung secara visual dan interaktif akan membantu siswa memahami materi IPA. Misalnya, mempelajari sistem pernapasan dengan melihat animasi alur udara yang masuk ke paru-paru bisa jauh lebih efektif daripada hanya membaca deskripsi di buku. Sebagai contoh, gambar Organ pernapasan manusia dapat digunakan untuk menjelaskan alur udara dari hidung hingga paru-paru, dan bagaimana proses pertukaran oksigen terjadi. Gambar ini membantu siswa memahami materi yang sulit dibayangkan hanya melalui teks.



Gambar 2. 1 Organ pernapasan manusia

Sistem pernapasan manusia terdiri dari berbagai organ, termasuk rongga hidung, faring, laring, trakea, bronkus, paru-paru, dan diafragma. Semua organ ini berperan dalam mengalirkan oksigen dari udara masuk ke dalam darah serta mengeluarkan karbon dioksida dari dalam tubuh melalui proses pernapasan.



Gambar 2. 2 Bumi dan Tata Surya

Tata Surya adalah suatu sistem yang terdiri dari berbagai objek langit, dengan matahari sebagai pusatnya. Di sekeliling matahari, terdapat planet-planet seperti merkurius, venus, bumi, mars, jupiter, saturnus, uranus, dan neptunus yang melintas di orbit berbentuk elips. Tata Surya ini terbentuk sekitar 4,6 miliar tahun yang lalu dari kumpulan gas dan debu di ruang angkasa.



Gambar 2. 3 Ekosistem

Ekosistem merupakan suatu sistem yang terbentuk melalui interaksi timbal balik antara makhluk hidup (biotik), seperti produsen (tumbuhan), konsumen (hewan), dan pengurai (dekomposer), serta faktor abiotik seperti cahaya matahari, air, dan udara. Keanekaragaman hayati dalam ekosistem ini meliputi berbagai jenis organisme yang saling berinteraksi demi menjaga keseimbangan lingkungan.

2.4.1 Visualisasi Konsep IPA dalam Pembelajaran Berbasis Digital

Dunia pendidikan telah mengalami transformasi karena kemajuan teknologi, termasuk cara pelajaran disampaikan. Pembelajaran konvensional, yang biasanya dilakukan di ruang kelas melalui metode ceramah, cenderung kurang fleksibel dan bersifat satu arah. Metode ini memungkinkan siswa untuk mendapatkan informasi secara pasif tanpa terlibat secara aktif. Pembelajaran melalui media sosial, di sisi lain, lebih fleksibel, menarik, dan dua arah. Siswa dapat menjangkau konten kapan saja dan di mana saja, serta memberikan umpan balik melalui percakapan atau komentar di dunia maya. Gambar berikut menunjukkan perbedaan antara pembelajaran berbasis media sosial dan pembelajaran konvensional:



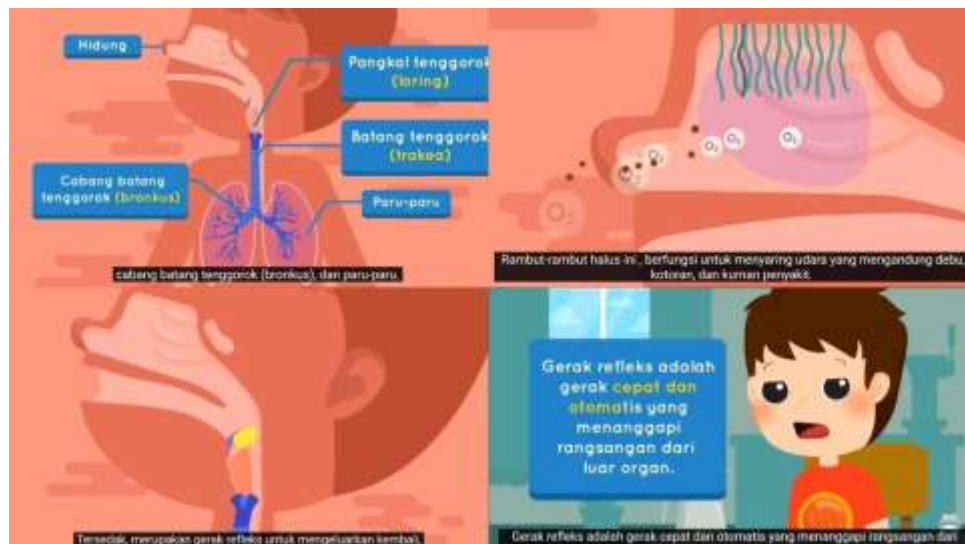
Gambar 2. 4 Perbedaan antara pembelajaran konvensional dan pembelajaran berbasis media sosial

Dengan ilustrasi seperti di atas kita tahu bahwa pembelajaran media sosial menciptakan peluang untuk berkomunikasi yang lebih banyak, dan mendorong siswa untuk lebih aktif serta mandiri dalam proses belajarnya.

2.4.2 Cuplikan Konten Pembelajaran IPA di Media Sosial

Saat ini, banyak platform media sosial menawarkan konten pembelajaran IPA dengan animasi, eksperimen sederhana, atau penjelasan berbasis cerita. Siswa dapat mengakses video pembelajaran ini melalui akun pendidikan mereka, yang tersedia secara gratis untuk semua orang. Buku yang berisi teks tidak hanya memberikan teori tetapi juga

memberikan gambaran konsep yang sulit dipahami. Gambar berikut menunjukkan cuplikan konten pembelajaran IPA yang dikemas dalam video pendek yang dapat diakses di platform media sosial:



Gambar 2. 5 Video Pembelajaran sistem pernapasan manusia

Dengan melihat tampilan seperti di atas siswa akan lebih mudah memahami proses ilmiah yang tidak bisa diamati secara langsung dengan melihat cuplikan seperti ini. Selain itu, presentasi yang sederhana dan visual dapat meningkatkan minat siswa dan keterlibatan mereka selama proses pembelajaran.

2.5 Machine Learning

Machine learning merupakan salah satu kategori dari kecerdasan buatan (AI) yang memungkinkan sistem untuk memperoleh pengetahuan dari data tanpa harus diprogram secara langsung. *machine learning* memberi kesempatan bagi komputer untuk "mempelajari" dari pengalaman yang telah ada dan data sebelumnya dengan menggunakan algoritma untuk menganalisis data, mengenali pola, dan membuat prediksi. sebuah teknologi yang membuat komputer tanpa perlu diprogram secara langsung untuk melakukan setiap tugasnya dapat belajar dari data. Sistem menggunakan pola-pola dalam data untuk memprediksi atau membuat keputusan di masa depan.

2.5.1 Tahapan dalam Machine Learning

Dalam *machine learning* Ada beberapa tahapan penting yang harus dilakukan agar model bisa bekerja dengan baik, antara lain ide/masalah bisnis, pengumpulan data, pemilihan model. Pra-pemrosesan data, melatih model, dan evaluasi model.

2.5.1.1 Ide / Masalah Bisnis

Hal Ini adalah aspek terpenting dalam *machine learning* yang berfungsi sebagai penggerak utama untuk diselesaikan dengan algoritma yang sesuai dalam pembelajaran mesin. (Suma, 2020) Tahapan ini adalah awal dari semua proses *machine learning*. Di sinilah merumuskan:

1. Apa yang ingin diselesaikan?
2. Apakah ini masalah klasifikasi, regresi, atau clustering?
3. Data seperti apa yang dibutuhkan?

2.5.1.2 Pengumpulan Data

Pengumpulan data merupakan langkah fundamental dalam proses *machine learning* yang melibatkan identifikasi dan pengambilan data mentah yang akan digunakan sebagai bahan pelatihan model. Proses pengumpulan data melibatkan identifikasi dan pengumpulan sumber data, baik internal maupun eksternal, yang relevan untuk menyelesaikan suatu masalah. Namun, sebagian besar model *machine learning* tidak dapat bekerja efektif dengan data mentah karena penggunaan atribut seperti nama dan alamat secara langsung dapat menyebabkan *poor generalization* (kesalahan generalisasi). Hal ini terjadi karena atribut kategorikal dengan domain sangat luas (seperti nama atau alamat) sulit diproses oleh algoritma *machine learning* dan tidak memberikan pola yang dapat digeneralisasikan. Oleh karena itu, data mentah perlu diproses terlebih dahulu menjadi representasi yang lebih bermakna melalui *feature engineering* (rekayasa fitur) atau *feature extraction*. Proses ini mencakup berbagai transformasi data tergantung pada jenis data dan seberapa informatif suatu atribut terhadap target prediksi. (Suma, 2020).

2.5.1.3 Pemilihan Model

Pemilihan model dalam *machine learning* bertujuan untuk menemukan algoritma yang dapat memprediksi data dengan tingkat akurasi terbaik. Model dengan bias tinggi cenderung kurang akurat (*underfitting*), sedangkan model dengan varians tinggi cenderung terlalu rumit dan sensitif terhadap data pelatihan. Keseimbangan antara varians dan bias sangat penting untuk keberhasilannya. Kita ingin model yang cukup fleksibel untuk menangkap pola data tetapi juga stabil terhadap perubahan data. Evaluasi berbagai algoritma dan penyesuaian parameter adalah bagian dari proses ini untuk mendapatkan kinerja terbaik. (Suma, 2020)

Dalam dunia nyata, pemilihan algoritma Machine learning tidak hanya bergantung pada tingkat akurasi prediksi semata, Kecepatan proses (*runtime*), biaya komputasi yang diperlukan, ketersediaan alat dan infrastruktur pelatihan, kemudahan implementasi, dan tingkat interpretabilitas model yang diharapkan adalah beberapa faktor penting yang harus dipertimbangkan. Komponen interpretabilitas ini sangat penting ketika model harus memberikan penjelasan yang dapat dipahami tentang bagaimana keputusan atau prediksi dibuat berdasarkan input yang diberikan.

Faktor-faktor tertentu, seperti:

1. Sifat permasalahan
2. Jumlah data yang dapat diakses untuk latihan
3. Dukungan keputusan atau kasus kognitif ?
4. Membutuhkan kemampuan untuk menjelaskan hubungan antara input dan keputusan. (Suma, 2020).

2.5.1.4 Pra-Pemrosesan Data

Tahapan preprocessing diperlukan untuk mempersiapkan data mentah menjadi bentuk yang siap diolah. Tahap pertama adalah pembersihan data (*data cleaning*) yang bertujuan menghilangkan noise dan ketidakkonsistenan dalam dataset, seperti menangani

missing values, menghapus data duplikat, serta mengatasi *outlier* yang dapat mengganggu analisis. Selanjutnya dilakukan transformasi data (*data transformation*) untuk mengkonversi data ke format yang lebih sesuai, meliputi proses normalisasi/standarisasi nilai numerik dan encoding data kategorikal menjadi representasi numerik. Tahap terakhir adalah seleksi fitur (*feature selection*) dimana hanya fitur-fitur paling relevan yang dipilih berdasarkan kajian literatur dan hasil eksplorasi data awal, sehingga dapat meningkatkan efisiensi dan akurasi model. (Sitio & Sianturi, 2024)

2.5.1.5 Melatih Model

Proses Melatih model melibatkan Tiga jenis dataset berbeda digunakan dalam proses pembuatan model pembelajaran mesin. Setiap dataset memiliki fungsi unik. Dataset pelatihan, atau dataset pelatihan, berfungsi sebagai dasar utama di mana model belajar mengenali pola dan hubungan dalam data. Selama proses pelatihan, dataset validasi berfungsi sebagai penguji sementara yang memungkinkan kita mengevaluasi kinerja model dan melakukan penyempurnaan parameter. Dataset validasi ini membantu menentukan apakah model sudah mampu digeneralisasi atau masih membutuhkan penyesuaian. Setelah model menunjukkan bahwa dia memerlukan penyesuaian, dataset validasi menunjukkan bahwa

2.5.1.6 Evaluasi Model

Evaluasi model merupakan tahap kritis Untuk mengukur kemampuan algoritma untuk memprediksi data, evaluasi model adalah langkah penting. Berbagai metrik, seperti akurasi, presisi, *recall*, atau skor F1 untuk klasifikasi, serta *MSE*, *RMSE*, atau *R-squared* untuk regresi, digunakan dalam proses ini. Untuk memastikan bahwa model tidak mengalami *overfitting* (terlalu kompleks) atau *underfitting* (terlalu sederhana), evaluasi dilakukan pada data pelatihan serta data validasi dan pengujian. Hasil pemeriksaan ini berfungsi sebagai dasar untuk menentukan apakah model sudah siap untuk digunakan atau memerlukan penyempurnaan.

2.6 Jenis jenis machine learning

Secara garis besar, Algoritma machine learning dibagi menjadi tiga jenis, yaitu *supervised learning*, *unsupervised learning*, dan *reinforcement learning*

1) *Supervised Learning*

Supervised learning mengadopsi konsep pendekatan fungsi, dimana pada dasarnya algoritma dilatih agar dapat memilih fungsi-fungsi yang paling menggambarkan input dimana X tertentu membuat estimasi terbaik dari y. Namun, pada kenyataannya tidak sedikit orang yang kesulitan menemukan fungsi yang paling cocok. Hal ini karena sebenarnya algoritma bergantung pada asumsi yang digunakan. Jika ada asumsi yang tidak terpenuhi, tidak jarang hasil pengolahan data akan menimbulkan bias. Oleh karena itu, algoritma ini membutuhkan data latih yang benar sehingga sistem dapat mempelajari polanya dan regresi, klasifikasi, KNN, *Naive Bayes*, *Decision Trees*, *Regresi linier*, *Support Vector Machine*, dan *neural network*. (Simanullang, 2021)

2) *Unsupervised Learning*

Algoritma *unsupervised learning* adalah algoritma yang tidak membutuhkan data berlabel. Pada *unsupervised learning*, algorithm tidak membutuhkan data training. Algoritma ini digunakan dalam mendeteksi pola dan pemodelan deskriptif yang tidak membutuhkan kategori atau output berlabel yang menjadi dasar algoritma untuk mencari model yang tepat. Algoritma ini digunakan untuk clustering dan association rule. Keunggulan dari *unsupervised learning* adalah karena tidak membutuhkan label, algoritma lebih leluasa untuk mencari pola yang mungkin sebelumnya belum diketahui. Sedangkan kekurangan dari algoritma ini adalah sulitnya menemukan informasi dalam data karena tidak ada label dan lebih sulit untuk membandingkan output dengan inputnya.(Simanullang, 2021)

3) *Reinforcment Learning*

Reinforcement learning adalah untuk menggunakan observasi yang dikumpulkan dari interaksi bersama lingkungan guna mengambil tindakan yang akan memaksimalkan output dan meminimalkan resiko. Algoritma ini akan terus belajar secara berulang-ulang. Dalam algoritma ini ada agen yang akan belajar dari interaksi dengan *environment*-nya. Untuk menghasilkan model, algoritma *reinforcement learning* melalui beberapa tahap antara lain agen mengamati data input, setelah itu agen melakukan suatu tindakan untuk mengambil keputusan. Setelah keputusan diambil, agen akan memperoleh "reward" atau penguatan dari lingkungan. Lalu kembali mengamati input dan proses pengambilan keputusan kembali dilakukan namun dengan tambahan penguatan dari lingkungan sehingga hasil keputusan yang diambil lebih akurat. (Kushariyadi et al., 2024).

2.7 XGBoost

XGBoost (Extreme Gradient Boosting) adalah salah satu algoritma ML berbasis ensemble yang menggunakan teknik boosting. Boosting adalah teknik yang secara iteratif menggabungkan beberapa model sederhana (biasanya pohon keputusan atau pohon keputusan) untuk membuat model yang lebih kuat dan akurat. Pendekatan gradient boosting yang digunakan *XGBoost* mengurangi kesalahan secara iteratif untuk meningkatkan kinerja prediksi. Model baru dibuat untuk mengurangi kesalahan yang dibuat oleh model sebelumnya. *XGBoost* bekerja dengan prinsip *boosting aditif*, di mana prediksi akhir dituliskan sebagai :

$$\hat{y}_i = \sum_{t=1}^T f_t(x_i), \quad f_t \in \mathcal{F} \quad (2.1)$$

Proses diawali dengan memberikan prediksi awal (*base score*) yang sama untuk semua data, biasanya berupa rata-rata target. Pada tahap berikutnya dihitung residual atau selisih antara nilai aktual dengan prediksi sebelumnya:

$$r_i^{(t)} = y_i - \hat{y}_i^{(t-1)} \quad (2.2)$$

Residual ini kemudian dipelajari oleh pohon baru. Pohon dibangun dengan menentukan split terbaik menggunakan perhitungan gain yang menunjukkan seberapa besar peningkatan kualitas jika pemisahan dilakukan:

$$\text{Gain} = \frac{1}{2} \left(\frac{(\sum g_L)^2}{\sum h_L + \lambda} + \frac{(\sum g_R)^2}{\sum h_R + \lambda} - \frac{(\sum g)^2}{\sum h + \lambda} \right) - \gamma \quad (2.3)$$

Setiap daun dari pohon menghasilkan bobot (*leaf weight*) yang dihitung dengan rumus:

$$w^* = - \frac{\sum g}{\sum h + \lambda} \quad (2.4)$$

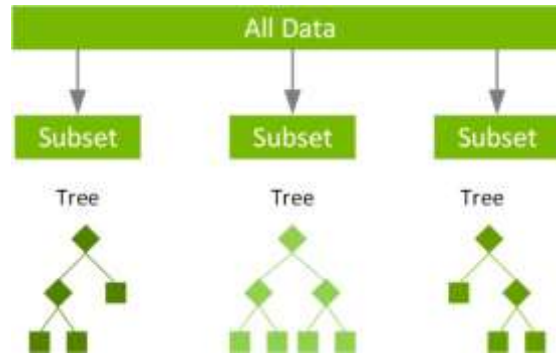
Bobot ini digunakan untuk memperbarui prediksi. Proses update dilakukan dengan persamaan:

$$\hat{\mathbf{y}}^{(t)} = \hat{\mathbf{y}}^{(t-1)} + \eta \mathbf{f}_t(\mathbf{x}) \quad (2.5)$$

dengan η adalah learning rate yang mengatur seberapa besar kontribusi tiap pohon. Langkah ini diulang hingga jumlah pohon sesuai parameter *n_estimators*. Semua mekanisme tersebut berjalan otomatis ketika perintah *fit()* pada XGBoost dijalankan, seperti berikut:

```
model = XGBRegressor(objective='reg:squarederror')
model.fit(X_train, y_train).
```

XGboost menghasilkan model baru berdasarkan kesalahan dari model sebelumnya, sehingga model yang dihasilkan menjadi lebih tepat. XGBoost menerapkan metode boosting yang membangun serangkaian model lemah secara bertahap pada subset data dengan tujuan mengurangi nilai mean square error ($\hat{y} - y$) dari model F , di mana $\hat{y} = f(x)$, dan memberikan bobot pada setiap prediksi lemah sesuai dengan kinerjanya. Prediksi akhir dihasilkan dengan menjumlahkan prediksi pohon keputusan yang sudah diberi bobot. Algoritma inti XGBoost yang berfokus pada minimisasi fungsi kerugian dijelaskan dalam Persamaan.(L. Mauro and G. Carmeci dalam Khairunnisa, 2023).



Gambar 2. 6 Visualisasi XGBoost

2.8 Bayesian optimization

Bayesian Optimization (BO) adalah sebuah teknik pengoptimalan yang berlandaskan pada model probabilitas dan sangat berguna untuk fungsi objektif yang mahal serta tidak dapat dianalisis dengan cara biasa. BO menggunakan model probabilistik, seperti Proses Gaussian, untuk merepresentasikan fungsi objektif dan memilih titik evaluasi berikutnya dengan memperhatikan keseimbangan antara eksplorasi dan pemanfaatan. (Dwi Utami et al., 2025).

Bayesian Optimization adalah teknik untuk mengoptimalkan hyperparameter yang berfungsi dengan menciptakan model perantara untuk memperkirakan bentuk dari fungsi tujuan. Fungsi tujuan yang diterapkan dalam studi ini adalah $-MSE$, sehingga persamaannya dapat dituliskan sebagai berikut:

$$f(\theta) = -MSE_{val}(\theta), \quad \theta^* = \arg \max_{\theta} f(\theta) \quad (2.6)$$

Proses pencarian akan berlangsung hingga tercapai jumlah iterasi yang diinginkan. Pada tahap awal, pencarian dilakukan dengan memilih beberapa titik secara acak dari pengaturan hyperparameter. Model representatif akan dibentuk berdasarkan penilaian awal ini dan akan diperbaharui terus-menerus dengan menambahkan data hasil penilaian di setiap putaran. (Simamora et al., 2025).

2.9 Python

Python memiliki kode sumber yang mudah, sehingga tidak sulit untuk dibuat, diingat, dan digunakan kembali. Ini sangat mempermudah proses pembuatan aplikasi, mulai dari tahap penulisan kode, pengujian, sampai perbaikan jika terdapat kesalahan atau cacat. Dengan menggunakan *Python*, aplikasi dapat dikembangkan dengan efisien berkat keberadaan pustaka yang lengkap, sehingga membuat kode yang ditulis menjadi lebih mudah. *Python* juga memiliki banyak fungsi, memberikan kesempatan bagi pengembang untuk membuat berbagai tipe aplikasi, mulai dari situs web, aplikasi jaringan, aplikasi dalam bidang robotika, hingga AI. Terdapat banyak modul siap guna yang dapat digunakan untuk memfasilitasi pengembangan aplikasi sesuai dengan kebutuhan. (Retnoningsih & Pramudita, 2020)

Salah satu keunggulan *Python* adalah interoperabilitas yang tinggi, yang memungkinkan interaksi dengan berbagai bahasa pemrograman yang berbeda. Aplikasi yang dibuat dengan *Python* bisa dijalankan di hampir semua platform sistem operasi.

2.10 Metode Evaluasi

2.10.1 Root Mean Squared Error (RMSE)

RMSE berfungsi untuk membedakan antara nilai yang diprediksi dan nilai yang nyata. Nilai RMSE yang lebih tinggi menunjukkan bahwa tingkat akurasi semakin rendah, sedangkan nilai RMSE yang lebih rendah menunjukkan bahwa tingkat akurasi semakin tinggi. (Sautomo & Pardede, 2021).

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n \frac{(\hat{y}_i - y_i)^2}{n}} \quad (2.7)$$

Dimana,

$\hat{y}_i$: nilai hasil forecast

y_i : nilai observasi ke – i

N : banyaknya data

2.10.2 Mean Absolute Error (MAE)

Evaluasi *MAE* ini adalah nilai rata-rata dari kesalahan mutlak antara hasil prediksi dan nilai data yang sesungguhnya. (Sautomo & Pardede, 2021)

$$MAE = \frac{1}{n} \sum |f_i - y_i| \quad (2.8)$$

Dimana,

f_i : nilai hasil forecast

y_i : nilai observasi ke – i

n : banyaknya data

2.10.3 Coefficient of Determination (R^2)

Coefficient of Determination, atau yang lebih dikenal sebagai *R-Squared* (R^2), merupakan salah satu metrik evaluasi paling penting dalam regresi karena memberikan ukuran proporsi variansi dari data target yang bisa dijelaskan oleh model prediksi.

$$R^2 = 1 - \frac{\sum_{i=1}^n (X_i - Y_i)^2}{\sum_{i=1}^n (Y_i - \bar{Y})^2} \quad (2.9)$$

Coefficient of Determination dapat diartikan sebagai rasio dari variasi variabel tergantung yang bisa diprediksi melalui variabel-variabel bebas (Chicco et al., 2021).

2.11 Penelitian terdahulu

Tabel 2. 1 Alur Penelitian

NO	Nama Peneliti (Tahun)	judul	Hasil penelitian
1.	Khairunnisa, A. (2023)	Perbandingan Model Random forest dan Xgboost untuk prediksi kejahatan keusilaan di provinsi jawa barat	Dalam penelitian ini membandingkan kinerja model klasifikasi XGBoost dan Random Forest pada data tidak seimbang. Hasilnya, XGBoost yang dikombinasikan dengan SMOTE menunjukkan kinerja lebih baik. Meski akurasi tinggi, sensitivitas terhadap kelas minoritas rendah. Dengan menyesuaikan threshold dari 0,50 menjadi 0,49, balanced

			accuracy meningkat, membuktikan pentingnya evaluasi model yang menyeluruh, bukan hanya akurasi.
2.	Simamora, F. P., Purba, R., & Pasha, M. F. (2025)	Optimisasi Hyperparameter BiLSTM Menggunakan Bayesian Optimization untuk Prediksi Harga Saham	Penelitian ini mengungkapkan bahwa Bayesian Optimization merupakan metode yang efektif untuk mengoptimalkan hyperparameter pada model BiLSTM. Hasil terbaik diperoleh dari pengujian pada data saham BBKA dengan parameter: batch 78, dropout 0,2924, epoch 190, dan neuron 24; BYAN dengan batch 127, dropout 0,1234, epoch 180, dan neuron 19; serta TLKM dengan batch 9, dropout 0,2168, epoch 125, dan neuron 22. Ketika dibandingkan dengan Grid Search, model BiLSTM yang dioptimasi menggunakan Bayesian Optimization menunjukkan tingkat akurasi prediksi yang lebih tinggi, yang terlihat dari nilai MAPE yang lebih baik.
3.	Dwi Utami, Fathoni Dwiatmoko, & Nuari Anisa Sivi. (2025)	Analisis Pengaruh Bayesian Optimization Terhadap Kinerja SVM Dalam Prediksi Penyakit Diabetes.	Dalam Penelitian ini menunjukkan bahwa Bayesian Optimization dapat secara efektif meningkatkan kinerja model SVM dalam prediksi dini diabetes. Melalui metode ini, diperoleh kombinasi parameter optimal, seperti kernel, C, dan gamma. Model SVM yang telah dioptimalkan mencapai akurasi sebesar 95%, yang lebih tinggi dibandingkan dengan model SVM yang tidak dioptimalkan yang mencatat akurasi 94%. Selain itu, terdapat peningkatan presisi dan F1-Score sebesar 7%.
4.	Nugraha, A. C., & Irawan, M. I. (2023)	Komparasi Deteksi Kecurangan pada Data Klaim Asuransi Pelayanan	Penelitian ini mengungkapkan bahwa metode SVM dan XGBoost efektif digunakan

		Kesehatan Menggunakan Metode Support Vector Machine (SVM) dan Extreme Gradient Boosting (XGBoost)	untuk mengklasifikasikan penipuan dalam data klaim layanan kesehatan, setelah melalui serangkaian proses seperti pembersihan data, rekayasa fitur, one-hot encoding, oversampling, dan normalisasi Z-Score. Hasil klasifikasi menunjukkan bahwa XGBoost memberikan performa yang lebih baik dibandingkan dengan SVM, dengan nilai Balanced Accuracy dan Recall masing-masing mencapai 0,98 dan 0,984, sementara SVM hanya meraih 0,874 dan 0,854.
5.	Maulita, I., & Wahid, A. M. (2024)	Prediksi Magnitudo Gempa Menggunakan Random Forest, Support Vector Regression, XGBoost, LightGBM, dan Multi-Layer Perceptron Berdasarkan Data Kedalaman dan Geolokasi.	Penelitian ini membandingkan lima model machine learning dalam memprediksi magnitudo gempa berdasarkan data kedalaman dan geolokasi. Hasil menunjukkan bahwa LightGBM dan Random Forest memberikan performa terbaik dengan MAE dan RMSE rendah serta R^2 tinggi. XGBoost menunjukkan performa cukup baik, meskipun masih di bawah dua model tersebut. Sementara itu, SVR dan MLP menunjukkan hasil yang kurang optimal, dengan error tinggi dan R^2 rendah. Model ensemble seperti Random Forest dan LightGBM lebih unggul dalam memprediksi magnitudo gempa, khususnya dalam konteks data geofisika yang kompleks dan non-linear.
6.	Salsabil, M., Azizah, N. L., & Eviyanti, A. (2024)	Implementasi Data Mining dalam Melakukan Prediksi Penyakit Diabetes Menggunakan Metode Random Forest dan Xgboost	Penelitian ini menggunakan metode Random Forest dan XGBoost untuk memprediksi penyakit diabetes berdasarkan 768 data klinis dan biokimia yang diambil dari Kaggle, dengan melibatkan 9 indikator. Setelah melalui tahap

			preprocessing yang mencakup penanganan nilai hilang, penghilangan outlier, dan normalisasi, sebanyak 688 data dipilih untuk keperluan pelatihan dan pengujian dengan teknik cross-validation. Hasil evaluasi yang dilakukan dengan metrik akurasi, presisi, recall, dan F1-score menunjukkan bahwa model XGBoost menghasilkan performa yang lebih baik, dengan akurasi mencapai 76%, dibandingkan dengan Random Forest yang memperoleh akurasi sebesar 74%.
7.	Putra, B. S. C., Tahyudin, I., Kusuma, B. A., & Isnaini, K. N. (2024)	Efektivitas Algoritma Random Forest, XGBoost, dan Logistic Regression dalam Prediksi Penyakit Paru-paru	Berdasarkan evaluasi, XGBoost adalah algoritma terbaik untuk klasifikasi dataset ini. Sebelum tuning, XGBoost memiliki akurasi 94,4%, precision 100%, recall 88,38%, dan F1-score 93,83%. Setelah tuning, akurasi sedikit meningkat menjadi 94,44%, dengan precision 94,98%, recall 94,44%, dan F1-score 94,41%. Random Forest hampir sama baiknya, dengan akurasi meningkat dari 94,3% menjadi 94,33% setelah tuning, dan F1-score naik dari 93,70% menjadi 94,30%. Logistic Regression kurang memuaskan dengan akurasi hanya 87,40% setelah tuning. Hasil ini menunjukkan bahwa algoritma berbasis pohon keputusan seperti XGBoost dan Random Forest lebih efektif untuk klasifikasi data kompleks, karena seimbang antara precision dan recall.
8.	Hidayat, R. R., Larung, E. Y. P., Sinaga, E., CS, A., Ibrahim, I., Kardi, I. S., Putra, M. F. P., Samal, A. B., & Babingga, H. (2024)	Analitik prediktif sepakbola: Model Machine Learning Bri liga 1 Indonesia	Hasil penelitian menunjukkan bahwa model XGBoost dan Gradient Boosting Regressor memiliki akurasi tinggi dalam memprediksi hasil pertandingan

			BRI Liga 1 Indonesia. Model XGBoost untuk klasifikasi hasil pertandingan mencapai akurasi 98. 4%, dengan nilai precision, recall, dan F1-score yang hampir sempurna. Ini menunjukkan kemampuan model untuk mengenali pola kemenangan, kekalahan, dan hasil seri dengan baik. Model Gradient Boosting Regressor memiliki R-squared score 0,992, berarti dapat menjelaskan hampir semua variasi skor pertandingan dari variabel yang digunakan.
9.	Edi Jaya Kusuma, Ririn Nurmandhani, Lenci Aryani, Ika Pantiawati, & Guruh Fajar Shidik. (2025)	Optimasi Model Extreme Gradient Boosting Dalam Upaya Penentuan Tingkat Risiko Pada Ibu Hamil Berbasis Bayesian Optimization (BOXGB)	Pada penelitian ini, Implementasi algoritma Bayesian Optimization (BO) pada penentuan tingkat risiko kehamilan menggunakan model Decision Tree (DT) dan Extreme Gradient Boosting (XGB) menunjukkan hasil yang positif. Setelah proses normalisasi dan pembagian data menjadi data latih dan uji, BO digunakan untuk mengoptimasi hyper-parameter. Hasil evaluasi menunjukkan bahwa iterasi BO berpengaruh terhadap performa model, dengan XGB mencapai akurasi 87% serta presisi, recall, dan F1-score masing-masing sebesar 88%, 87%, dan 88%. BO terbukti mampu meningkatkan kinerja model dalam memprediksi risiko kehamilan berdasarkan data klinis. Namun, eksplorasi ruang hyper-parameter yang lebih luas dan teknik optimasi lain masih menjadi peluang pengembangan.

BAB III

METODOLOGI PENELITIAN

3.1 Jenis Penelitian

Penelitian ini adalah penelitian kuantitatif dengan pendekatan eksperimental yang bertujuan untuk mengembangkan model prediksi hasil belajar siswa berdasarkan aktivitas mereka di media sosial. Pendekatan ini dipilih karena dapat mengukur secara objektif efektivitas algoritma *Machine Learning*, khususnya *XGBoost* yang telah dioptimalkan melalui *Bayesian Optimization*, dalam memprediksi hasil akademik siswa dengan akurasi yang tinggi. Selain itu, penelitian ini bersifat komputasional karena mencakup pengolahan data, pemodelan, dan evaluasi algoritma selama prosesnya. Penelitian ini di fokuskan pada penggunaan data numerik yang diperoleh dari pre-test, kuis, dan post-test siswa untuk membangun model prediksi yang akurat.

3.2 Metode Penelitian

Metode yang digunakan dalam studi ini adalah eksperimen dengan pendekatan kuantitatif, yang didukung oleh pengujian dan evaluasi model menggunakan algoritma *XGBoost* serta metode *Bayesian Optimization*. Dalam hal ini, *XGBoost* berfungsi sebagai model utama untuk melakukan prediksi, sementara *Bayesian Optimization* digunakan untuk menyempurnakan *hyperparameter* model demi mencapai hasil prediksi yang optimal.

3.3 Teknik Pengumpulan Data

Dalam penelitian ini, teknik mengumpulkan data dilakukan dari pendekatan kuantitatif. Data diperoleh melalui tes tertulis dan aktivitas belajar yang menggunakan media sosial. Terdapat tiga jenis data utama yang dikumpulkan dari siswa, yaitu nilai pre-test, kuis, dan post-test. Pengumpulan data dilakukan secara langsung di sekolah selama kegiatan pembelajaran IPA berlangsung.

1. *Pre-Test*

Tes awal diberikan kepada siswa sebelum mereka mengakses materi melalui video pembelajaran yang diunggah di media sosial, seperti YouTube atau TikTok. Tujuan dari pre-test ini adalah untuk mengevaluasi sejauh mana pemahaman awal siswa terhadap topik IPA yang akan mereka pelajari.

2. Kuis Pembelajaran (Kuis 1–3)

Setelah menyaksikan setiap video pembelajaran, siswa diharapkan untuk mengerjakan kuis yang telah disediakan secara daring. Kuis ini berperan sebagai alat evaluasi untuk mengukur pemahaman mereka terhadap materi yang disampaikan melalui media sosial. Skor yang diperoleh dari kuis tersebut menjadi salah satu indikator untuk memprediksi hasil belajar siswa.

3. *Post-Test*

Setelah seluruh rangkaian video dan kuis selesai, siswa akan menjalani post-test sebagai evaluasi akhir untuk mengukur sejauh mana pemahaman mereka meningkat setelah mengikuti pembelajaran berbasis media sosial. Hasil dari post-test ini akan menjadi variabel terikat (*target*) dalam pemodelan prediksi.

Data yang diperoleh dari ketiga sumber tersebut kemudian direkap dalam sebuah tabel dan dimasukkan ke dalam file *spreadsheet*. Selanjutnya, data tersebut diolah menggunakan bahasa pemrograman *Python* dengan bantuan pustaka seperti *pandas*, *scikit-learn*, dan *Xgboost* untuk melakukan analisis dan prediksi. Selain itu, kami juga menerapkan metode *Bayesian Optimization* sebagai teknik untuk mencari hyperparameter terbaik, yang bertujuan mengoptimalkan performa model *XGBoost* agar hasil prediksi menjadi lebih akurat dan stabil.

3.4 Alat Bantu Penelitian

Penelitian ini memakai alat bantuan yang ada dalam proses penelitiannya, berikut ini beberapa alat yang akan digunakan dalam penelitian ini adalah sebagai berikut:

1. Perangkat Keras (*Hardware*)

Satu unit komputer dengan spesifikasi yaitu :

- *Processor : Intel(R) Core(TM) i5-10400 CPU @ 2.90GHz*
- *RAM : 16 GB*
- *VGA : NVIDIA GeForce RTX 2060 Super Ventus GP OC*
- *Storage : 236 GB SSD, 931 GB HHD*

2. Perangkat Lunak (*Software*)

- *Sistem Operasi : Windows 11 64-bit*
- *Aplikasi Pengolah Kata : Microsoft Word 2019*
- *Bahasa Pemrograman : Python*
- *Integrated Development Environment (IDE) : Visual Studio code*

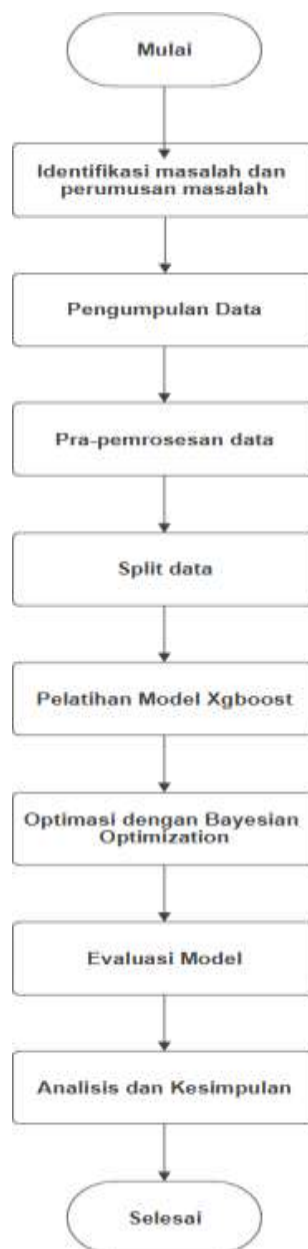
Dalam penelitian ini, proses pemodelan prediksi dilakukan menggunakan bahasa pemrograman *Python* melalui *Visual Studio Code* sebagai IDE. Beberapa pustaka (*library*) *Python* yang digunakan antara lain:

- *Xgboost* untuk membangun model prediksi
- *Bayesian-optimization* untuk melakukan *tuning hyperparameter* secara otomatis
- *Scikit-learn* untuk evaluasi model menggunakan metrik *MAE*, *RMSE*, dan *R²*.

3.5 Alur penelitian

Untuk mencapai tujuan dalam membangun model prediksi hasil belajar siswa yang optimal, penelitian ini mengintegrasikan algoritma *XGBoost* dengan metode optimasi hyperparameter *Bayesian Optimization*. Pemilihan kedua metode ini bertujuan untuk meningkatkan akurasi prediksi secara signifikan dan stabil, berdasarkan data yang mencakup nilai pre-test, hasil kuis, dan post-test. Agar proses implementasi metode dalam

penelitian ini lebih jelas dan sistematis, peneliti menyusun alur penelitian dalam bentuk flowchart. Diagram tersebut menggambarkan langkah demi langkah mulai dari pengumpulan data hingga evaluasi model, dengan harapan dapat memberikan pemahaman visual kepada pembaca tentang cara pembangunan dan pengoptimalan model prediksi menggunakan bahasa pemrograman *Python*. Di bawah ini, dapat dilihat flowchart yang menunjukkan alur penerapan metode *XGBoost* dan *Bayesian Optimization* dalam memprediksi hasil belajar siswa.



Gambar 3. 1 Alur Penelitian

Adapun keterangan mengenai perancangan analisis pada flowchart di atas ini adalah sebagai berikut:

1. Proses Pertama Pada tahap ini, peneliti melakukan pengamatan terhadap fenomena atau permasalahan yang muncul di lapangan, terutama yang berkaitan dengan efektivitas pembelajaran IPA yang memanfaatkan media sosial. Ditemukan bahwa masih sedikit pendekatan yang mengkaji hasil belajar siswa dengan metode prediksi yang berbasis data digital. Setelah masalah diidentifikasi, peneliti kemudian merumuskan pertanyaan penelitian dan menetapkan tujuan yang ingin dicapai. Tujuan utama dari penelitian ini adalah untuk membangun sebuah model prediksi hasil belajar siswa dengan memanfaatkan metode *XGBoost* dan *Bayesian Optimization*, yang didasarkan pada interaksi siswa dengan media pembelajaran digital.
2. Proses Kedua Peneliti mengumpulkan data dengan cara mengadakan pre-test untuk mengetahui nilai siswa sebelum pembelajaran dimulai. Selain itu, mereka juga mencatat skor kuis yang diperoleh dari video pembelajaran IPA yang diunggah di media sosial, serta nilai post-test setelah pembelajaran selesai. Data ini diambil dari siswa SMP yang berperan sebagai responden.
3. Proses Ketiga Data yang telah dikumpulkan akan melalui pemeriksaan awal. Proses ini mencakup beberapa langkah, seperti memeriksa adanya data yang kosong, melakukan normalisasi jika diperlukan, serta memisahkan fitur (X) dari target (Y). Semua langkah ini bertujuan untuk memastikan bahwa data siap digunakan dalam tahap pelatihan model.
4. Proses Keempat Data biasanya dibagi menjadi dua bagian, yaitu data pelatihan (training) dan data pengujian (testing), dengan rasio umum 80:20. Pembagian ini bertujuan agar model dapat mempelajari pola dari sebagian data dan kemudian diuji menggunakan data yang belum pernah dia lihat sebelumnya.

5. Proses Lima Pada tahap ini, model *XGBoost* dilatih dengan menggunakan data pelatihan. Model ini akan mempelajari pola-pola yang menghubungkan nilai pre-test dan kuis dengan nilai post-test siswa.
6. Proses Enam Untuk meningkatkan akurasi model, kami melakukan penyetelan hyperparameter dengan menggunakan metode *Bayesian Optimization*. Metode ini memungkinkan kami untuk menemukan kombinasi parameter terbaik secara lebih efisien dibandingkan dengan pendekatan grid search.
7. Proses Ketujuh Model yang telah dioptimasi dievaluasi dengan menggunakan berbagai metrik, antara lain *MAE (Mean Absolute Error)*, *RMSE (Root Mean Square Error)*, dan *R² (Coefficient of Determination)*. Proses evaluasi ini bertujuan untuk menilai seberapa akurat model dalam memprediksi nilai siswa.
8. Proses Delapan Tahap akhir melibatkan analisis terhadap hasil prediksi model, dengan membandingkannya pada nilai aktual. Dari sini, kita dapat menarik kesimpulan mengenai efektivitas media sosial dalam mendukung pembelajaran IPA serta menilai akurasi model yang telah dibangun.

BAB IV

HASIL DAN PEMBAHASAN

4.1 Pengumpulan data

Data yang diterapkan dalam studi ini adalah nilai siswa pada kelas VIII dalam pembelajaran IPA yang menggunakan media sosial. Pengumpulan data dilakukan melalui tes Pre-Test, Quis 1, Quis 2, Quis 3, dan Post-Test. Semua data tersebut disusun dalam sebuah file *spreadsheet* Excel yang berjudul "Data Siswa. *xlsx*" dan selanjutnya dianalisis dengan menggunakan bahasa pemrograman *Python*.

4.2 Pra-pemrosesan data

Pra-pemrosesan data sangat krusial dalam studi ini karena bertujuan untuk mengubah data mentah menjadi format yang siap digunakan untuk melatih model machine learning. Data yang dianalisis terdiri dari nilai pre-test, kuis 1, kuis 2, kuis 3, dan post-test dari 30 siswa yang berpartisipasi dalam penelitian ini. Tujuan dari penelitian ini adalah untuk meramalkan nilai post-test siswa berdasarkan nilai pre-test serta hasil kuis yang telah dilakukan sebelumnya. Oleh sebab itu, setiap langkah dalam pra-pemrosesan dilakukan untuk memastikan bahwa model dapat menerima input dalam kondisi yang bersih, lengkap, dan terstruktur dengan baik.

4.2.1 Import Library

Langkah pertama dalam pra-pemrosesan adalah mengimpor library atau pustaka yang dibutuhkan dalam proses pelatihan dan evaluasi model.

```
import pandas as pd
from xgboost import XGBRegressor
from sklearn.model_selection import train_test_split
from sklearn.metrics import mean_absolute_error, mean_squared_error, r2_score
from bayes_opt import BayesianOptimization
import numpy as np
```

Gambar 4. 1 Import Library

Berdasarkan dari gambar 4.1 dapat dijelaskan fungsi dari library tersebut

- a. *pandas* berfungsi untuk membaca dan mengatur dataset siswa.
- b. *xgboost* dimanfaatkan untuk menciptakan model prediktif yang berbasis *XGBoost*.
- c. *bayes_opt* diterapkan untuk secara otomatis menemukan hyperparameter yang optimal.
- d. *train_test_split* digunakan untuk memisahkan data menjadi set pelatihan dan pengujian.
- e. *mean_absolute_error*, *mean_squared_error*, dan *r2_score* digunakan untuk mengevaluasi kinerja model secara kuantitatif.
- f. *numpy* dipakai untuk melakukan perhitungan matematis, khususnya dalam penilaian model.

4.2.2 Import Dataset

Dataset didapatkan dalam format file Excel dengan nama "Data Siswa. xlsx".

Dataset ini berisi nilai-nilai siswa yang mencakup pre-test, quis 1 hingga 3, serta post-test.

```
# Baca data dari file Excel
df = pd.read_excel("Data Siswa.xlsx")
```

Gambar 4. 2 Import Dataset

4.2.3 Pengecekan dan Penanganan Nilai Kosong

Sebagian siswa memiliki informasi yang tidak penuh. Oleh karena itu, dilakukan pemeriksaan dan pengisian nilai yang hilang dengan menggunakan metode imputasi rata-rata.

```
# Isi nilai kosong dengan rata-rata kolom
df.fillna(df.mean(numeric_only=True), inplace=True)
```

Gambar 4. 3 Mengisi Nilai Kosong

4.2.4 Pemisahan Fitur dan Target

```
# Pisahkan fitur dan target
x = df[['Pre-Test', 'Quis 1', 'Quis 2', 'Quis 3']]
y = df['Post-test']
```

Gambar 4. 4 Memisahkan fitur prediksi dan target

Model yang digunakan dalam kajian ini memerlukan input (fitur) dan keluaran (target). Fitur yang diambil adalah nilai pre-test dan quis, sedangkan target yang ingin diprediksi adalah nilai post-test.

4.2.5 Split data

Setelah data dipersiapkan, langkah berikutnya adalah membagi data menjadi dua bagian: data pelatihan dan data pengujian. Sesuai dengan standar dalam pembelajaran mesin, rasio yang digunakan adalah 80% untuk data pelatihan dan 20% untuk data pengujian. Yang dimana *test_size* = 0.2 Artinya, 20% dari total data akan dipakai sebagai data untuk pengujian. Sementara itu, 80% sisanya akan secara otomatis digunakan sebagai data untuk pelatihan. Sementara itu *random_state* = 42 digunakan untuk memastikan bahwa pembagian data dilakukan dengan cara yang sama setiap kali program dijalankan, sehingga hasil yang didapat dapat diulang.

```
# Bagi data menjadi training dan testing
x_train, x_test, y_train, y_test = train_test_split(x, y, test_size=0.2, random_state=42)
```

Gambar 4. 5 Pembagian Data *Training* Dan *Testing*

4.2.6 Deskripsi Dataset

Dataset mencakup total 30 murid. Sebagian murid memiliki data yang lengkap, sementara lainnya hanya memiliki beberapa nilai. Setelah proses imputasi, seluruh data menjadi valid dan siap untuk digunakan.

Tabel 4. 1 Dataset Awal

No	Nama Siswa	Pre-Test	Quis 1	Quis 2	Quis 3	Post-test
1.	Aura Putri	94	70	80	75	65
2.	Alissya az zahra	92	80	90		75
3.	Adinda Mharani Dabutar	78	78	65	50	55
4.	Amal Luddin Sitompul	57	20	30	30	30
5.	Ahmad Luthfi Lubis	82	85		82	80
6.	Anggun fitriani	55	60	80		60
7.	Chalisa Zahira	74	70	60	90	80
8.	Daffa Zairi Lubis	83	60	50	65	75
9.	Elparisah Putri Ayu	74	90	65	55	70
10.	Fandly Syahputra Dlm	65	50	10		65

Pada Tabel 4.1 merupakan gambaran struktur awal dataset sebelum dilakukan pra-pemrosesan. Dataset yang digunakan dalam penelitian ini merupakan data mentah yang diperoleh dari hasil pengumpulan nilai pre-test, quis 1, quis 2, quis 3, dan post-test dari 31 siswa. Pada data tersebut terdapat beberapa nilai kosong pada kolom Quis maupun Post-test.

Tabel 4. 2 Dataset Hasil Pra-Pemrosesan

No	Nama Siswa	Pre-Test	Quis 1	Quis 2	Quis 3	Post-test
1.	Aura Putri	94	70	80	75	65
2.	Alissya az zahra	92	80	90	75	75
3.	Adinda Mharani Dabutar	78	78	65	50	55
4.	Amal Luddin Sitompul	57	20	30	30	30

5.	Ahmad Luthfi Lubis	82	85	60	82	80
6.	Anggun fitriani	55	60	80	75	60
7.	Chalisa Zahira	74	70	60	90	80
8.	Daffa Zairi Lubis	83	60	50	65	75
9.	Elparisah Putri Ayu	74	90	65	55	70
10.	Fandly Syahputra Dlm	65	50	10	75	65

Pada Tabel 4.2 merupakan hasil dari tahapan pra-pemrosesan data untuk memastikan data layak digunakan dalam proses pelatihan model machine learning. Nilai-nilai kosong diisi menggunakan metode imputasi rata-rata kolom bersangkutan.

4.3 Pelatihan model *xgboost*

Dalam penelitian ini, model *XGBoost* dikembangkan untuk memahami keterkaitan antara variabel input (*Pre-Test*, Quis 1, Quis 2, dan Quis 3) dan variabel yang diinginkan (*Post-test*). Model itu dibuat berdasarkan data pelatihan dengan menggunakan parameter dasar sebelum proses optimasi hyperparameter dilakukan. Pada tahap ini, model *XGBoost* dilatih menggunakan data pelatihan yang sudah diproses. Tujuannya adalah untuk membangun model awal sebelum dilakukan optimasi. Proses pelatihan ini menggunakan parameter default atau parameter yang telah ditentukan sebelumnya.

```
# latih model Xgboost
best_model = XGBRegressor(
    max_depth=int(best_params['max_depth']),
    learning_rate=best_params['learning_rate'],
    n_estimators=int(best_params['n_estimators']),
    random_state=42
)
best_model.fit(X_train, y_train)

# Prediksi
y_pred = best_model.predict(X_test)
```

Gambar 4. 6 Pelatihan Model *Xgboost*

4.4 Optimasi dengan Bayesian optimization

Setelah model awal selesai dibuat, tahap berikutnya adalah melakukan optimasi hyperparameter untuk memperbaiki kinerja model. Untuk ini, digunakan *Bayesian Optimization* karena metode ini dapat menemukan parameter terbaik secara efektif.

```
# Jalankan optimasi
optimizer = BayesianOptimization(f=xgb_evaluate, pbounds=pbounds, random_state=42)
optimizer.maximize(init_points=5, n_iter=10)

# Dapatkan parameter terbaik
best_params = optimizer.max['params']
```

Gambar 4. 7 Optimasi *Bayesian Optimization*

4.4.1 Fungsi evaluasi dan ruang parameter

Pada proses pengoptimalan hyperparameter menggunakan metode Bayesian Optimization, dibutuhkan fungsi evaluasi dan area pencarian parameter. Fungsi ini akan dimanfaatkan oleh algoritma *Bayesian Optimization* untuk menilai kinerja model *XGBoost* berdasarkan variasi kombinasi parameter yang diuji.

```
# Fungsi untuk optimasi Bayesian
def xgb_evaluate(max_depth, learning_rate, n_estimators):
    model = XGBRegressor(
        max_depth=int(max_depth),
        learning_rate=learning_rate,
        n_estimators=int(n_estimators),
        random_state=42,
        verbosity=0
    )
    model.fit(X_train, y_train)
    preds = model.predict(X_test)
    return -mean_squared_error(y_test, preds)

# Ruang pencarian parameter
pbounds = {
    'max_depth': (3, 10),
    'learning_rate': (0.01, 0.3),
    'n_estimators': (50, 300)
}
```

Gambar 4. 8 Fungsi Evaluasi dan Ruang Parameter

4.4.1 Fungsi evaluasi dan ruang parameter

Pada stadion ini, dilakukan pemetaan untuk menampilkan hasil dari prediksi model serta proses penyempurnaan hyperparameter. Pemetaan ini bertujuan agar temuan penelitian dapat lebih jelas dimengerti melalui grafik.

```
# (A) Scatter Actual vs Predicted
plt.figure(figsize=(6, 6))
plt.scatter(y_test, y_pred, alpha=0.85)
min_val = min(y_test.min(), y_pred.min())
max_val = max(y_test.max(), y_pred.max())
plt.plot([min_val, max_val], [min_val, max_val], linestyle='--')
plt.xlabel("Actual Post-test")
plt.ylabel("Predicted Post-test")
plt.title("Actual vs Predicted")
# teks metrik di sudut
textstr = f"MAE = {mae:.3f}\nRMSE = {rmse:.3f}\nR2 = {r2:.3f}"
plt.gcf().text(0.65, 0.15, textstr, bbox=dict(facecolor='white', alpha=0.85))
plt.tight_layout()
plt.savefig("fig/plot_actual_vs_predicted.png", dpi=300)

# (B) Riwayat Bayesian Optimization
plt.figure(figsize=(7, 5))
plt.plot(bo_iters, bo_scores, marker='o')
plt.xlabel("Iterasi")
plt.ylabel("Skor (RMSE) - lebih besar lebih baik")
plt.title("Riwayat Bayesian Optimization")
plt.tight_layout()
plt.savefig("fig/plot_bo_history.png", dpi=300)
```

Gambar 4. 9 Visualisasi hasil prediksi dan optimasi

4.5 Uji validasi dan Reliabilitas

Tabel 4. 3 Tabel uji validasi

Pertanyaan	R	R tabel	Kesimpulan
1	0,842	0,3610	Valid
2	0,771	0,3610	Valid
3	0,861	0,3610	Valid
4	0,721	0,3610	Valid
5	0,749	0,3610	Valid

Hasil uji validitas menunjukkan bahwa seluruh butir pertanyaan pada kuesioner memiliki nilai korelasi (r hitung) lebih besar daripada nilai pembanding (r tabel = 0,3610). Dengan demikian, setiap pertanyaan dapat dikatakan valid, artinya pertanyaan yang diberikan benar-benar mampu mengukur hal yang ingin diteliti. Hal ini menegaskan bahwa instrumen yang digunakan sudah tepat sasaran dan sesuai dengan tujuan penelitian.

Tabel 4. 4 Tabel Reliabilitas

<i>Cronbach's Alpha</i>	<i>N of Items</i>
.794	6

Selanjutnya, hasil uji reliabilitas menggunakan Cronbach's Alpha memperoleh nilai sebesar 0,794. Angka ini berada di atas standar minimum 0,70, sehingga kuesioner dinyatakan reliabel. Reliabel di sini berarti jawaban responden konsisten; jika instrumen yang sama diberikan lagi dalam kondisi serupa, maka kemungkinan besar akan memberikan hasil yang sama.

Meskipun ada beberapa nilai korelasi antar pertanyaan yang tidak setinggi yang lain, hal ini lebih disebabkan oleh jumlah responden yang terbatas (30 siswa), sehingga variasi jawaban kurang beragam. Faktor tersebut tidak menunjukkan kelemahan pada instrumen, melainkan lebih pada keterbatasan sampel penelitian. Secara keseluruhan, hasil uji validitas dan reliabilitas ini mengonfirmasi bahwa kuesioner yang dipakai dalam penelitian telah memenuhi standar yang baik, sesuai, dan dapat dipilih sebagai sarana pengumpulan data.. Dengan adanya instrumen yang valid dan reliabel, maka data yang dihasilkan juga dapat dipertanggungjawabkan dan mendukung akurasi hasil penelitian ini.

4.6 Evaluasi model

Evaluasi terhadap model dilaksanakan untuk menilai seberapa besar tingkat akurasi prediksi yang dihasilkan oleh metode XGBoost yang telah dioptimasi melalui Optimasi Bayesian. Uji coba ini dilaksanakan dengan memanfaatkan data uji yang diperoleh dari pembagian dataset menjadi 80% untuk pelatihan dan 20% untuk pengujian.

Tabel 4. 5 Data Aktual dan Prediksi

No	Nama Siswa	Pre-Test	Quis 1	Quis 2	Quis 3	Actual Post-test	Predicted Post-test	Selisih
1	Wahyu Kurnia	75	60	40.000000	80.000	60.0	73.546051	13.546051
2	M. Erlangga Gifhari	70	30	60.185185	75.000	45.0	45.436100	0.436100
3	Ramon Cleo	70	50	20.000000	70.000	40.0	33.356106	-6.643894
4	M. Anwar Siregar	20	20	30.000000	73.125	30.0	37.986778	7.986778
5	Elparisah Putri Ayu	74	90	65.000000	55.000	70.0	56.543159	-13.456841
6	Fandly Syahputra Dlm	65	50	10.000000	73.125	65.0	46.115311	-18.884689

Pada table 4.3 Hasil evaluasi model pada data uji adalah sebagai berikut:

1. MAE : 10.159058888753256
2. $RMSE$: 11.756797228431319
3. R^2 : 0.32756728127889023

a. MAE (Mean Absolute Error)

Untuk memastikan hasil nilai dari data uji dapat dilakukan perhitungan dengan rumus MAE

(Mean Absolute Error) yaitu :

$$MAE = \frac{1}{n} \sum |f_i - y_i| \quad (4.1)$$

$$= \frac{1}{6} [13.546051 + 0.436100 + 6.643894 + 7.986778 + 13.45684 + 18.884689]$$

$$= \frac{60.953353}{6}$$

$$= 10.1590$$

Nilai *MAE* yang mencapai 10. 1590 menunjukkan bahwa kesalahan rata-rata absolut antara nilai yang sebenarnya dan yang diprediksi adalah sekitar 10,16 poin. *MAE* lebih mudah dipahami karena menggunakan satuan yang identik dengan data (nilai ujian siswa), dan nilai yang semakin kecil menunjukkan bahwa prediksi semakin akurat.

b. *RMSE (Root Mean Squared Error)*

Untuk memastikan hasil nilai dari data uji dapat dilakukan perhitungan dengan rumus *RMSE (Root Mean Squared Error)* yaitu :

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n \frac{(\hat{y}_i - y_i)^2}{n}} \quad (4.2)$$

$$= \sqrt{\frac{1}{6} [(73.546051 - 60.0)^2 + (45.436100 - 45.0)^2 + (33.356106 - 40.0)^2 + (37.9866778 - 30.0)^2 + (56.543159 - 70.0)^2 + (46.115311 - 65.0)^2]}$$

$$= \sqrt{\frac{1}{6} [183.4634 + 0.1902 + 44.1283 + 63.7884 + 181.1005 + 356.5941]}$$

$$= \sqrt{\frac{828.2649}{6}}$$

$$= \sqrt{138.04415}$$

$$= 11.7568$$

Nilai Nilai *RMSE* yang tercatat sebesar 11,7568, sedangkan nilai *MAE* sebesar 10,1590. Karena nilai *MAE* < *RMSE*, hal ini menunjukkan adanya beberapa perbedaan prediksi yang cukup besar pada sebagian siswa. Perbedaan yang besar ini disebut outlier, yaitu data yang nilainya jauh berbeda dibandingkan sebagian besar data lainnya. Kondisi ini wajar, karena *RMSE* memberikan penalti lebih besar terhadap kesalahan ekstrem, yaitu kesalahan prediksi yang nilainya jauh lebih besar daripada rata-rata kesalahan. Dengan demikian, meskipun rata-rata kesalahan absolut (*MAE*) masih sekitar 10 poin, keberadaan beberapa hasil prediksi yang jauh melenceng dari nilai sebenarnya membuat *RMSE*

meningkat menjadi sekitar 11,75 angka nilai ujian. Hal ini mengindikasikan bahwa model masih perlu ditingkatkan agar hasil prediksi bisa lebih mendekati nilai sebenarnya.

c. R^2 (Coefficient of Determination)

Untuk memastikan hasil nilai dari data uji dapat dilakukan perhitungan dengan rumus R^2 (Coefficient of Determination) yaitu :

$$R_2 = 1 - \frac{\sum_{i=1}^n (X_i - Y_i)^2}{\sum_{i=1}^n (\bar{Y}_i - Y_i)^2} \quad (4.3)$$

$$\bar{Y}_i = \frac{60.0 + 45.0 + 40.0 + 30.0 + 70.0 + 65.0}{6} = 51.6667$$

$$R^2 = 1 - \frac{828.2649}{1229.1666}$$

$$= 1 - 0.6724$$

$$= 0.3276$$

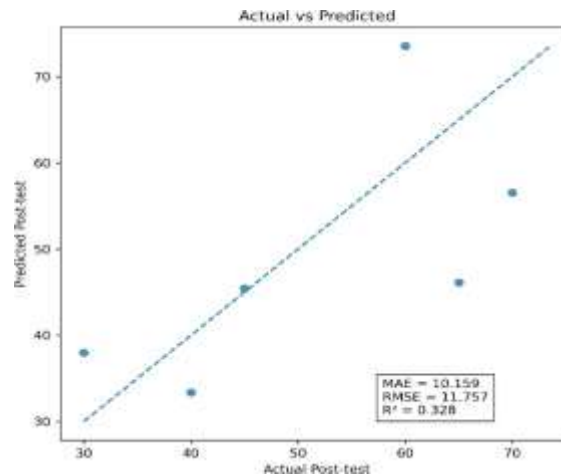
Nilai R^2 sebesar 0.3276 menunjukkan bahwa sekitar 32,76% variasi pada nilai post-test siswa dapat dijelaskan oleh model, sedangkan sisanya (67,24%) dipengaruhi oleh faktor lain di luar model ini. Faktor-faktor di luar model tersebut dapat mencakup aspek internal siswa seperti motivasi belajar, minat terhadap mata pelajaran IPA, kemampuan dasar akademik, serta kedisiplinan dalam belajar mandiri. Selain itu, faktor eksternal seperti tingkat kehadiran siswa, dukungan guru dan orang tua, kualitas jaringan internet saat mengakses media sosial, hingga intensitas dan cara siswa berinteraksi dengan konten pembelajaran juga berpengaruh terhadap hasil belajar namun tidak terekam dalam dataset penelitian ini.

Nilai R^2 termasuk rendah, menunjukkan bahwa model belum mampu menjelaskan seluruh variasi data dengan baik. Rendahnya nilai R^2 disebabkan oleh beberapa hal, antara lain jumlah data yang terbatas (hanya 30 siswa) sehingga pola yang terbentuk kurang baik, variabel input yang digunakan hanya nilai pre-test dan kuis sehingga belum mewakili semua faktor penentu hasil belajar, serta adanya variasi (nilai siswa) antar individu yang menimbulkan noise pada data. Dengan kata lain, model tidak memiliki cukup informasi

untuk menangkap keseluruhan kompleksitas faktor yang mempengaruhi hasil belajar siswa.

4.7 Visualisasi Hasil

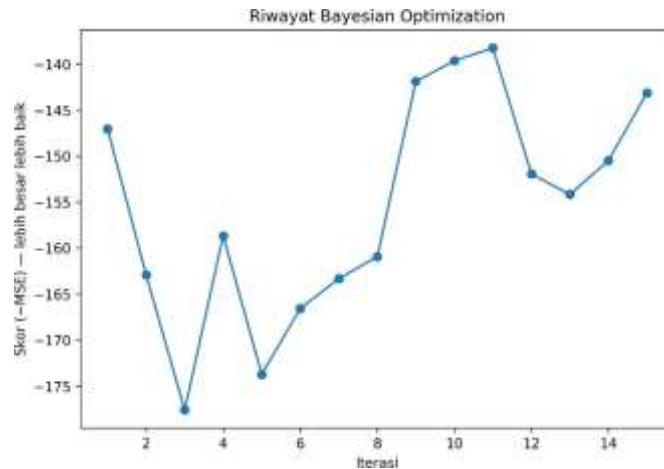
4.7.1 Grafik nilai *post-test* aktual dengan hasil prediksi



Gambar 4. 10 Hasil Aktual Dengan Prediksi

Gambar 4. 10 menunjukkan grafik ini menggambarkan hubungan antara nilai *post-test* yang sebenarnya dan hasil estimasi dari model. Dapat terlihat bahwa hanya sebagian kecil titik data yang benar-benar mendekati garis regresi, sementara sebagian besar titik lainnya menyebar cukup jauh dari garis. Kondisi ini sesuai dengan nilai R^2 sebesar 0,328 yang menunjukkan bahwa model hanya mampu menjelaskan sekitar 32,8% variasi nilai *post-test* siswa. Variasi yang dimaksud adalah perbedaan nilai *post-test* antar siswa yang ada di dalam data, misalnya ada siswa yang memperoleh nilai tinggi, sedang, dan rendah. Dengan kata lain, meskipun model *XGBoost* yang dioptimalkan dengan Bayesian Optimization mampu menghasilkan prediksi dengan rata-rata kesalahan *MAE* sebesar 10,159 dan *RMSE* sebesar 11,757, model belum mampu sepenuhnya menangkap pola perbedaan nilai tinggi–rendah tersebut sehingga banyak prediksi yang masih menyimpang dari nilai aktual.

4.7.2 Grafik riwayat bayesian optimization



Gambar 4. 11 Hasil riwayat *Bayesian Optimization*

Gambar 4.11 menunjukkan grafik riwayat proses pengoptimalan hyperparameter menggunakan metode Bayesian Optimization. Sumbu horizontal menggambarkan jumlah iterasi, sedangkan sumbu vertikal menunjukkan skor evaluasi.

Pada iterasi awal (1–3), garis grafik terlihat menurun tajam, menandakan kualitas model masih rendah dengan tingkat kesalahan yang cukup besar. Selanjutnya, mulai dari iterasi ke-4 hingga ke-10, garis grafik mengalami kenaikan bertahap yang signifikan, menunjukkan bahwa *Bayesian Optimization* mulai menemukan kombinasi hyperparameter yang lebih baik sehingga kualitas prediksi meningkat. Pada iterasi ke-11, garis mencapai puncak tertinggi, yang menandakan titik kombinasi hyperparameter terbaik. Setelah itu, pada iterasi ke-12 hingga ke-13, garis grafik sedikit menurun kembali, menggambarkan adanya fluktuasi kualitas prediksi akibat variasi parameter yang diuji. Namun, pada iterasi berikutnya, garis grafik kembali naik dan cenderung stabil, yang menandakan bahwa proses optimasi berhasil menemukan kombinasi parameter yang cukup optimal dan konsisten. Secara keseluruhan, pola garis pada grafik ini memperlihatkan bahwa *Bayesian Optimization* mampu meningkatkan performa model secara bertahap. Fluktuasi yang terjadi merupakan hal wajar karena setiap iterasi menguji kombinasi parameter yang

berbeda, namun tren kenaikan secara umum membuktikan bahwa metode ini efektif dalam menemukan hyperparameter dengan kualitas prediksi yang lebih baik.

4.8 Analisis Hasil

Berdasarkan penilaian model yang terlihat pada Tabel 4.5, nilai Rata-rata Kesalahan Absolut (*MAE*) yang dicapai adalah 10,16, yang menunjukkan bahwa rata-rata perbedaan absolut antara hasil prediksi dan nilai nyata post-test siswa sekitar 10 poin. Sementara itu, nilai Kesalahan Kuadrat Rata-rata Akar (*RMSE*) sebesar 11,75 sedikit lebih tinggi dibandingkan *MAE*. Hal ini disebabkan karena dalam perhitungannya, *RMSE* memberikan penalti yang lebih besar terhadap kesalahan prediksi yang ekstrem atau *outlier*. Kesalahan ekstrem atau *outlier* ini merujuk pada beberapa kasus siswa yang nilai prediksi modelnya berbeda cukup jauh dari nilai sebenarnya, sehingga membuat nilai *RMSE* meningkat lebih tinggi dibandingkan *MAE*. Dengan kata lain, rata-rata kesalahan prediksi model masih berada pada selisih sekitar 10–12 poin dari nilai post-test siswa, namun adanya sebagian data dengan selisih yang besar menunjukkan bahwa model belum sepenuhnya mampu menangkap pola hasil belajar seluruh siswa dengan baik.

Di sisi lain, nilai Koefisien Determinasi (R^2) yang diperoleh adalah 0,327, yang menunjukkan bahwa model dapat menjelaskan sekitar 32,7% variasi dalam data nilai pasca-tes siswa, sementara sisanya (sekitar 67,3%) dipengaruhi oleh faktor-faktor lain yang tidak terakomodasi dalam model. Nilai R^2 yang relatif rendah ini bisa disebabkan oleh:

- A. Jumlah data yang terbatas (hanya 30 siswa), sehingga hubungan antar variabel belum dapat terlihat dengan baik.
- B. Variabel masukan yang hanya mencakup nilai pre-tes dan kuis, tanpa mempertimbangkan faktor eksternal seperti latar belakang siswa, tingkat kehadiran, atau interaksi lewat media sosial.
- C. Tingginya variasi antara siswa, yang membuat model kesulitan untuk menemukan pola yang umum.

Namun demikian, model *XGBoost* yang telah dioptimasi dengan *Bayesian Optimization* masih mampu memberikan prediksi yang dekat dengan nilai aktual pada sebagian besar siswa, terlihat dari perbedaan yang tidak signifikan antara nilai nyata dan prediksi pada beberapa contoh. Dengan kata lain, meskipun model ini belum sepenuhnya akurat, ia menunjukkan potensi sebagai alat bantu analisis hasil belajar berdasarkan data dari pembelajaran siswa.

Nilai R^2 dalam studi ini masih tergolong rendah. Hal ini dapat disebabkan oleh beberapa faktor, antara lain:

1. Keterbatasan variabel input yang digunakan dalam model. Hanya nilai pre-test dan kuis yang dilibatkan, sementara hasil belajar siswa juga dipengaruhi faktor lain seperti motivasi, kedisiplinan, kehadiran, dan interaksi dengan media sosial yang tidak tercakup dalam data.
2. Hubungan antar variabel yang bersifat tidak lurus (*non-linear*), artinya peningkatan pada satu variabel (misalnya nilai kuis) tidak selalu berbanding lurus dengan hasil *post-test*. Ada kondisi tertentu di mana nilai kuis tinggi belum tentu menghasilkan nilai *post-test* tinggi, sehingga sulit dijelaskan hanya dengan fitur sederhana yang tersedia.
3. Adanya ketidakaturan (*noise*) pada data, misalnya siswa yang menjawab kuis secara asal-asalan, tidak fokus saat tes, atau kondisi lingkungan saat ujian yang kurang mendukung. Hal ini menimbulkan data yang tidak konsisten, sehingga menambah kesalahan (*error*) pada model.
4. Karakteristik algoritma, di mana *XGBoost* meskipun kuat, tetap memiliki keterbatasan ketika pola data sangat bervariasi dengan informasi yang terbatas.

BAB IV

KESIMPULAN

5.1 Kesimpulan

Berdasarkan hasil studi yang telah dilakukan tentang peningkatan model prediksi hasil belajar siswa dalam mata pelajaran IPA dengan memanfaatkan media sosial melalui *XGBoost* dan *Bayesian Optimization*, beberapa kesimpulan dapat ditarik sebagai berikut:

1. Proses pra-pemrosesan data, seperti imputasi nilai kosong menggunakan rata-rata dan pembagian data training-testing dengan rasio 80:20, berhasil mempersiapkan data untuk pelatihan model.
2. Berdasarkan hasil evaluasi, nilai *MAE* sebesar 10,16 menunjukkan bahwa rata-rata kesalahan prediksi model terhadap nilai post-test siswa berada sekitar 10 poin. Sementara itu, nilai *RMSE* sebesar 11,75 sedikit lebih tinggi dibandingkan *MAE*. Perbedaan ini muncul karena *RMSE* lebih menekankan kesalahan yang besar pada beberapa data siswa yang prediksinya cukup jauh dari nilai sebenarnya. Dengan kata lain, meskipun sebagian besar prediksi sudah cukup dekat dengan nilai aktual, ada beberapa kasus yang selisihnya lebih besar dari rata-rata. Jika dibandingkan, *MAE* lebih cocok digunakan untuk menilai performa model dalam penelitian ini karena lebih mudah dipahami dan tidak terlalu dipengaruhi oleh beberapa kesalahan yang besar.

5.2 Saran

1. Untuk penelitian selanjutnya, disarankan untuk menggunakan jumlah data yang lebih besar agar model dapat belajar lebih banyak variasi pola dan menghasilkan akurasi yang lebih tinggi.
2. Peneliti juga disarankan untuk mencoba metode evaluasi lainnya, seperti cross-validation, serta eksperimen dengan algoritma lain, untuk membandingkan performa metode yang digunakan.
3. Diharapkan penelitian ini dapat menjadi referensi bagi guru atau pengembang pembelajaran dalam memanfaatkan teknologi machine learning untuk mengevaluasi efektivitas metode pembelajaran modern, khususnya yang memanfaatkan media sosial.
4. Bagi peneliti lain, evaluasi model tidak selalu harus dilakukan dengan library bawaan. Implementasi manual terhadap metrik evaluasi (seperti *MAE*, *RMSE*, dan *R²*) dapat meningkatkan pemahaman terhadap proses evaluasi dan memperkuat validitas hasil penelitian.

DAFTAR PUSTAKA

- Aisa, H. (2021). PENGARUH AKTIVITAS BELAJAR DAN LINGKUNGAN BELAJAR SISWA TERHADAP HASIL BELAJAR BAHASA INDONESIA SISWA KELAS VIII SMP NEGERI 3 KUSAMBI. *Jurnal Bastra (Bahasa Dan Sastra)*, 6(1), 40–48.
- Chicco, D., Warrens, M., & Jurman, G. (2021). The coefficient of determination R-squared is more informative than SMAPE, MAE, MAPE, MSE and RMSE in regression analysis evaluation. *PeerJ Computer Science*, 7, e623. <https://doi.org/10.7717/peerj-cs.623>
- Dwi Utami, Fathoni Dwi atmoko, & Nuari Anisa Sivi. (2025). Analisis Pengaruh Bayesian Optimization Terhadap Kinerja SVM Dalam Prediksi Penyakit Diabetes. *Infotek: Jurnal Informatika Dan Teknologi*, 8(1 SE-Articles), 140–150. <https://doi.org/10.29408/jit.v8i1.28468>
- Helal Mridha, M. (2023). The effect of social-media on students' learning: a case study of a selected government college. *I-Manager s Journal on School Educational Technology*, 17, 50–62. <https://doi.org/10.26634/jsch.17.4.18868>
- Khairunnisa, A. (2023). PERBANDINGAN MODEL RANDOM FOREST DAN XGBOOST UNTUK PREDIKSI KEJAHATAN KESUSILAAN DI PROVINSI JAWA BARAT. *JIKO (Jurnal Informatika Dan Komputer)*, 7(2), 202–208.
- Kushariyadi, K., Apriyanto, H., Herdiana, Y., Asy'ari, F. H., Judijanto, L., Pasrun, Y. P., & Mardikawati, B. (2024). *Artificial Intelligence: Dinamika Perkembangan AI Beserta Penerapannya*. PT. Sonpedia Publishing Indonesia.
- Nugraha, A. C., & Irawan, M. I. (2023). Komparasi Deteksi Kecurangan pada Data Klaim Asuransi Pelayanan Kesehatan Menggunakan Metode Support Vector Machine (SVM) dan Extreme Gradient Boosting (XGBoost). *Jurnal Sains Dan Seni ITS*, 12(1), A40–A46.
- Optimasi Model Extreme Gradient Boosting Dalam Upaya Penentuan Tingkat Risiko Pada Ibu Hamil Berbasis Bayesian Optimization (BOXGB). (2025). *Jurnal Teknologi*

- Informasi Dan Ilmu Komputer*, 12(1 SE-Ilmu Komputer), 111–120.
<https://doi.org/10.25126/jtiik.20251219001>
- Retnoningsih, E., & Pramudita, R. (2020). Mengenal Machine Learning Dengan Teknik Supervised Dan Unsupervised Learning Menggunakan Python. *BINA INSANI ICT JOURNAL*; Vol 7 No 2 (2020): *Bina Insani ICT Journal (Desember 2020)*; 156-165 ; *Bahasa Indonesia*; Vol 7 No 2 (2020): *Bina Insani ICT Journal (Desember 2020)*; 156-165 ; 2527-9777 ; 2355-3421 ; 10.51211/Biict.V7i2. <http://ejournal-binainsani.ac.id/index.php/BIICT/article/view/1422>
- Sautomo, S., & Pardede, H. F. (2021). Prediksi Belanja Pemerintah Indonesia Menggunakan Long Short-Term Memory (LSTM). *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 5(1 SE-Information Technology Articles).
<https://doi.org/10.29207/resti.v5i1.2815>
- Simamora, F. P., Purba, R., & Pasha, M. F. (2025). Optimisasi Hyperparameter BiLSTM Menggunakan Bayesian Optimization untuk Prediksi Harga Saham. *Jambura Journal of Mathematics*, 7(1), 8–13.
- Simanullang, A. M. (2021). *Makalah Machine Learning*.
- Siregar, N. A., Harahap, N. R., & Harahap, H. S. (2023). Hubungan antara pretest dan posttest dengan hasil belajar siswa kelas VII B di MTs Alwashliyah Pantai Cermin. *Jurnal Ilmiah Edunomika*, 7(1).
- Sitio, A., & Sianturi, F. (2024). Penerapan Algoritma Machine Learning dalam Analisis Pola Perilaku Penggunaan Internet. *DÍKÉ*. <https://doi.org/10.69688/dike.v2i2.102>
- Sudiantini, D. (2023). Pengaruh Media Sosial Dalam Prestasi Pendidikan. *Madani: Jurnal Ilmiah Multidisiplin*, 1(5), 184-188, (2023-06-05).
<https://doi.org/10.5281/zenodo.8005661>
- Suendarti, M., & Hasbullah. (2020). *Pemahaman Konsep Ilmu Pengetahuan Alam ditinjau dari Motivasi Belajar Siswa*.
- Suma, B. (2020). *Implementasi Machine Learning Di Dalam Prediksi Cuaca*.

<https://doi.org/10.13140/RG.2.2.16086.47680>

- Salsabil, M., Azizah, N. L., & Eviyanti, A. (2024). Implementasi Data Mining Dalam Melakukan Prediksi Penyakit Diabetes Menggunakan Metode Random Forest Dan Xgboost. *Jurnal Ilmiah Komputasi*; Vol. 23 No. 1 (2024): *Jurnal Ilmiah Komputasi : Vol. 23 No 1, Maret 2024*; 51-58 ; 2549-7227 ; 1412-9434. <https://ejournal.jakstik.ac.id/index.php/komputasi/article/view/3507>
- Maulita, I., & Wahid, A. M. (2024). Prediksi Magnitudo Gempa Menggunakan Random Forest, Support Vector Regression, XGBoost, LightGBM, dan Multi-Layer Perceptron Berdasarkan Data Kedalaman dan Geolokasi. *Jurnal Pendidikan Dan Teknologi Indonesia*; Vol 4 No 5 (2024): *JPTI - Mei 2024*; 221-232 ; 2775-4219 ; 2775-4227. <https://jpti.journals.id/index.php/jpti/article/view/470>
- Putra, B. S. C., Tahyudin, I., Kusuma, B. A., & Isnaini, K. N. (2024). Efektivitas Algoritma Random Forest, XGBoost, dan Logistic Regression dalam Prediksi Penyakit Paru-paru. *Techno.Com*; Vol. 23 No. 4 (2024): *November 2024*; 909-922 ; 2356-2579 ; 1412-2693. <https://publikasi.dinus.ac.id/index.php/technoc/article/view/11705>
- Hidayat, R. R., Larung, E. Y. P., Sinaga, E., CS, A., Ibrahim, I., Kardi, I. S., Putra, M. F. P., Samal, A. B., & Babingga, H. (2024). Analitik Prediktif Sepakbola: Model Machine Learning BRI Liga 1 Indonesia. *Multilateral : Jurnal Pendidikan Jasmani Dan Olahraga*; Vol 23, No 4 (2024): *Special Issue National Conference: Mengembalikan Marwah PJOK Dan Penerapan Teknologi Dalam Pembelajaran PJOK*; 386 - 399 ; 2549-1415 ; 1412-3428 ; 10.20527/Multilateral.V23i4. <https://ppjp.ulm.ac.id/journal/index.php/multilateralpjk/article/view/21889>

Coding

```
import os

import pandas as pd

import numpy as np

import matplotlib.pyplot as plt

from xgboost import XGBRegressor

from sklearn.model_selection import train_test_split

from sklearn.metrics import mean_absolute_error, mean_squared_error, r2_score

from bayes_opt import BayesianOptimization


# -----

# 1) Load & Prep Data

# -----

df = pd.read_excel("Data Siswa.xlsx")


# Isi nilai kosong numerik dengan rata-rata kolom

df.fillna(df.mean(numeric_only=True), inplace=True)


# Fitur & target (sesuaikan dengan nama kolom Anda)

X = df[['Pre-Test', 'Quis 1', 'Quis 2', 'Quis 3']]

y = df['Post-test']


# (opsional) simpan nama siswa untuk tabel hasil

nama_siswa = df['Nama Siswa'].values if 'Nama Siswa' in df.columns else None


# Split data (tetap konstan di semua laptop dengan random_state)

X_train, X_test, y_train, y_test, nama_train, nama_test = train_test_split(
```

```

X, y, nama_siswa, test_size=0.2, random_state=42
)

# -----
# 2) Bayesian Optimization
# -----

def xgb_evaluate(max_depth, learning_rate, n_estimators):

    model = XGBRegressor(

        max_depth=int(max_depth),

        learning_rate=learning_rate,

        n_estimators=int(n_estimators),

        objective='reg:squarederror',

        random_state=42,

        verbosity=0

    )

    # evaluasi langsung pada test set (sesuai skripsi Anda)

    model.fit(X_train, y_train, eval_set=[(X_test, y_test)], verbose=False)

    preds = model.predict(X_test)

    # BO memaksimalkan skor -> kembalikan -MSE

    return -mean_squared_error(y_test, preds)

# ruang pencarian

pbounds = {

    'max_depth': (3, 10),

    'learning_rate': (0.01, 0.3),

    'n_estimators': (50, 300)

}

```

```

optimizer = BayesianOptimization(f=xgb_evaluate, pbounds=pbounds, random_state=42)
optimizer.maximize(init_points=5, n_iter=10)

# simpan riwayat untuk grafik BO

bo_iters = list(range(1, len(optimizer.res) + 1))

bo_scores = [res['target'] for res in optimizer.res] # -MSE (semakin besar semakin baik)

best_params = optimizer.max['params']

# -----

# 3) Train Model Final

# -----

best_model = XGBRegressor(

    max_depth=int(best_params['max_depth']),

    learning_rate=best_params['learning_rate'],

    n_estimators=int(best_params['n_estimators']),

    objective='reg:squarederror',

    random_state=42

)

best_model.fit(X_train, y_train)

# Prediksi & metrik

y_pred = best_model.predict(X_test)

mae = mean_absolute_error(y_test, y_pred)

rmse = np.sqrt(mean_squared_error(y_test, y_pred))

r2 = r2_score(y_test, y_pred)

```

```

# -----

# 4) Tabel Hasil & Simpan

# -----

results = X_test.copy()

if nama_test is not None:

    results['Nama Siswa'] = nama_test

results['Actual Post-test'] = y_test.values

results['Predicted Post-test'] = y_pred

results['Selisih'] = results['Predicted Post-test'] - results['Actual Post-test']


results = results.reset_index(drop=True)

results.index += 1

results.insert(0, 'No', results.index)


print("\nHasil Perbandingan Prediksi vs Asli:")

cols_show = ['No'] + ([ 'Nama Siswa'] if 'Nama Siswa' in results.columns else []) + \

    ['Pre-Test', 'Quis 1', 'Quis 2', 'Quis 3',

     'Actual Post-test', 'Predicted Post-test', 'Selisih']

print(results[cols_show].to_string(index=False))


print("\nEvaluasi Model:")

print(f"MAE : {mae}")

print(f"RMSE : {rmse}")

print(f" $R^2$  : {r2}")


# simpan hasil

results.to_excel("Hasil_Prediksi_Lengkap.xlsx", index=False)

```

```

# -----

# 5) Visualisasi (2 Grafik)

# -----

os.makedirs("fig", exist_ok=True)

# (A) Scatter Actual vs Predicted

plt.figure(figsize=(6, 6))

plt.scatter(y_test, y_pred, alpha=0.85)

min_val = min(y_test.min(), y_pred.min())

max_val = max(y_test.max(), y_pred.max())

plt.plot([min_val, max_val], [min_val, max_val], linestyle='--')

plt.xlabel("Actual Post-test")

plt.ylabel("Predicted Post-test")

plt.title("Actual vs Predicted")

# teks metrik di sudut

textstr = f"MAE = {mae:.3f}\nRMSE = {rmse:.3f}\nR2 = {r2:.3f}"

plt.gcf().text(0.65, 0.15, textstr, bbox=dict(facecolor='white', alpha=0.85))

plt.tight_layout()

plt.savefig("fig/plot_actual_vs_predicted.png", dpi=300)

# (B) Riwayat Bayesian Optimization

plt.figure(figsize=(7, 5))

plt.plot(bo_iters, bo_scores, marker='o')

plt.xlabel("Iterasi")

plt.ylabel("Skor (−MSE) — lebih besar lebih baik")

plt.title("Riwayat Bayesian Optimization")

```

```
plt.tight_layout()

plt.savefig("fig/plot_bo_history.png", dpi=300)

print("\nFile tersimpan:")

print(" - fig/plot_actual_vs_predicted.png")

print(" - fig/plot_bo_history.png")

print(" - Hasil_Prediksi_Lengkap.xlsx")
```

Output Codingan

Hasil Perbandingan Prediksi vs Asli:									
No	Nama Siswa	Pre-Test	Quis 1	Quis 2	Quis 3	Actual Post-test	Predicted Post-test	Selisih	
1	Wahyu Kurnia	75	60	40.000000	80.000	60.0	73.546051	13.546051	
2	M. Erlangga Gifhari	70	30	60.185185	75.000	45.0	45.436100	0.436100	
3	Ramon Cleo	70	50	20.000000	70.000	40.0	33.356106	-6.643894	
4	M.Anwar Siregar	20	20	30.000000	73.125	30.0	37.986778	7.986778	
5	Elparisah Putri Ayu	74	90	65.000000	55.000	70.0	56.543159	-13.456841	
6	Fandly Syahputra Dlm	65	50	10.000000	73.125	65.0	46.115311	-18.884689	
Evaluasi Model:									
MAE : 10.159058888753256									
RMSE : 11.756797228431319									
R ² : 0.32756728127889023									

Hasil Perbandingan Prediksi vs Asli:

No	Nama Siswa	Pre-Test	Quis 1	Quis 2	Quis 3	Actual Post-test	Predicted Post-test	Selisih
1	Wahyu Kurnia	75	60	40.000000	80.000	60.0	73.546051	13.546051
2	M. Erlangga Gifhari	70	30	60.185185	75.000	45.0	45.436100	0.436100
3	Ramon Cleo	70	50	20.000000	70.000	40.0	33.356106	-6.643894
4	M.Anwar Siregar	20	20	30.000000	73.125	30.0	37.986778	7.986778
5	Elparisah Putri Ayu	74	90	65.000000	55.000	70.0	56.543159	-13.456841
6	Fandly Syahputra Dlm	65	50	10.000000	73.125	65.0	46.115311	-18.884689

Evaluasi Model:

MAE : 10.159058888753256
 RMSE : 11.756797228431319
 R² : 0.32756728127889023