

Klasifikasi Kinerja Konten YouTube Mobile Legends Menggunakan LightGBM dengan Pelabelan Berbasis K-Means dan Interpretasi SHAP

Wayan Praka^{1*}, Ella Kristiantini Susan², Lidya Nurmala Eva³

^{1,2} Magister Teknologi Informasi, Fakultas Pascasarjana, Universitas Teknologi Yogyakarta, Indonesia

Email: ^{1*}wayanpraka123@gmail.com, ²sekarlangit02@gmail.com, ³evanurmalalidya21@gmail.com

(* wayanpraka123@gmail.com)

Abstrak- YouTube merupakan platform berbagi video yang dimanfaatkan kreator konten, termasuk pada kategori gaming seperti *Mobile Legends*. Tingginya persaingan antar kreator *Mobile Legends* menjadikan analisis kinerja konten penting sebagai dasar penyusunan strategi pengembangan konten berbasis data. Namun, Sebagian besar penelitian terdahulu masih menggunakan pelabelan data secara manual sehingga rentan terhadap subjektivitas dan kurang efisien. Selain itu, model klasifikasi dengan akurasi tinggi cenderung bersifat *black-box*, sehingga sulit menjelaskan faktor-faktor yang memengaruhi hasil prediksi. Penelitian ini bertujuan mengembangkan model klasifikasi kinerja konten *YouTube Mobile Legends* melalui integrasi *K-Means*, *LightGBM*, dan *Shapley Additive exPlanations* (SHAP). Dataset terdiri dari 3.052 video yang diperoleh melalui *YouTube API v3*. Algoritma *K-Means* digunakan untuk menghasilkan pelabelan otomatis, sedangkan *LightGBM* digunakan sebagai model klasifikasi dan SHAP diterapkan untuk menginterpretasikan keputusan model. Hasil menunjukkan bahwa jumlah kluster paling optimal adalah $k=2$, yang menghasilkan dua kelas, yaitu Kinerja Tinggi sebanyak 72,44% dan Kinerja Rendah sebanyak 27,56% serta pembagian data 80% menjadi pelatihan dan 20% menjadi pengujian. Hasil dari model *LightGBM* yaitu *accuracy* mencapai 93,78%, *precision* mencapai 93,87%, *recall* mencapai 93,78%, dan *F1-score* mencapai 93,81% pada data pengujian. Hasil *10-Fold Cross Validation* dengan *accuracy* mencapai 92,75% dan $\pm 0,0148$ pada standar deviasi. Penerapan SHAP berhasil membuktikan bahwa *Subscriber* adalah faktor paling berpengaruh, disusul oleh Durasi, Usia Video dan Panjang Judul. Temuan ini menunjukkan bahwa integrasi *K-Means*, *LightGBM*, dan SHAP merupakan pendekatan yang efektif untuk mengklasifikasikan kinerja konten *Mobile Legends* di YouTube secara objektif, akurat, dan transparan serta berpotensi mendukung kreator konten dalam menyusun strategi pengembangan konten berbasis data.

Kata Kunci: Klasifikasi YouTube; Mobile Legends; K-Means; LightGBM; Interpretasi SHAP.

Abstract- YouTube is a video-sharing platform widely used by content creators, including those in the gaming category such as *Mobile Legends*. The intense competition among *Mobile Legends* content creators has made content performance analysis essential for developing data-driven content strategies. However, most previous studies have relied on manual data labeling, which is prone to subjectivity and lacks efficiency. Furthermore, high-accuracy classification models are often treated as *black-box* models, making it difficult to explain the factors influencing prediction outcomes. This study aims to develop a YouTube *Mobile Legends* content performance classification model by integrating *K-Means*, *LightGBM*, and *Shapley Additive exPlanations* (SHAP). The dataset consists of 3,052 videos collected using the YouTube Data API v3. The *K-Means* algorithm was employed to generate automatic labels, while *LightGBM* was used as the classification model and SHAP was applied to interpret the model's predictions. The results indicate that the optimal number of clusters was $K=2$, producing two performance classes: High Performance (72.44%) and Low Performance (27.56%). The dataset was subsequently divided into 80% for training and 20% for testing. On the testing dataset, the *LightGBM* model achieved an accuracy of 93.78%, precision of 93.87%, recall of 93.78%, and an *F1-score* of 93.81%. Furthermore, 10-fold cross-validation yielded a mean accuracy of 92.75% with a standard deviation of ± 0.0148 , demonstrating the model's robustness and stability. SHAP analysis revealed that Subscriber Count was the most influential feature, followed by Duration, Video Age, and Title Length. These findings demonstrate that the integration of *K-Means*, *LightGBM*, and SHAP provides an effective approach for objectively, accurately, and transparently classifying the performance of *Mobile Legends* content on YouTube, while offering valuable insights to support content creators in developing data-driven content strategies.

Keywords: YouTube classification; Mobile Legends; K-Means; LightGBM; SHAP.

1. PENDAHULUAN

YouTube telah berkembang menjadi salah satu platform berbagi video terbesar yang tidak hanya sekedar media hiburan, tetapi bisa menjadi sumber data penting dalam berbagai penelitian mengenai media sosial dan perilaku pengguna [1]. Seiring dengan meningkatnya jumlah kreator konten, persaingan untuk menghasilkan video yang mampu menarik perhatian audiens juga semakin tinggi. Dalam konteks tersebut, *user engagement* menjadi sebuah parameter penting yang menjadi dasar pengukuran kinerja suatu video karena mencerminkan tingkat interaksi pengguna terhadap konten yang dipublikasikan [2]. Selain itu, perkembangan format konten, seperti video reguler dan video pendek, menunjukkan adanya perbedaan karakteristik interaksi pengguna yang memengaruhi performa sebuah video [3]. Penelitian terdahulu juga menunjukkan bahwa faktor-faktor seperti durasi video, judul, isi konten, serta aspek emosional yang disampaikan memiliki pengaruh terhadap tingkat keterlibatan audiens di YouTube [4]. Fenomena tersebut juga terlihat pada kategori konten *Mobile Legends*, yang menjadi *game online* dengan jumlah pemain dan komunitas besar di Indonesia. Tingginya minat masyarakat terhadap permainan ini mendorong kreator untuk memproduksi konten *Mobile Legends* di YouTube sebagai media hiburan maupun sumber informasi bagi para pemain [5]. Beberapa konten *Mobile Legends* bahkan mampu mencapai daftar video *trending* di YouTube, yang menunjukkan tingginya perhatian dan interaksi dari pengguna [6]. Selain melalui video, tingginya aktivitas komunitas *Mobile Legends* juga tercermin pada berbagai penelitian mengenai

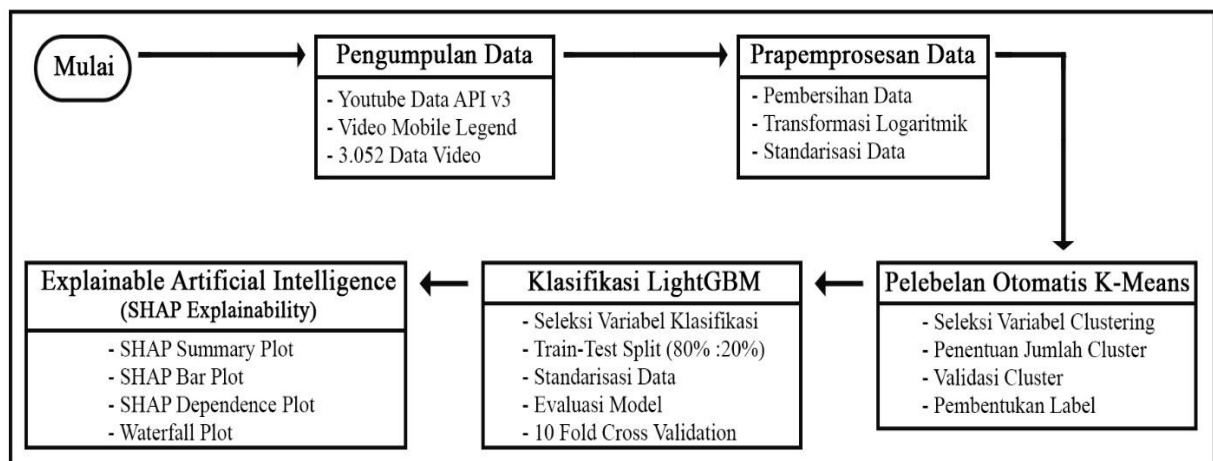
analisis sentimen ulasan aplikasi [7] serta preferensi pemain dalam memilih tipe hero [8]. Kondisi tersebut menunjukkan bahwa analisis terhadap kinerja konten *Mobile Legends* di *YouTube* menjadi penting untuk memahami karakteristik video yang mampu memperoleh kinerja tinggi.

Berbagai penelitian telah memanfaatkan teknik *machine learning* untuk menganalisis data media sosial. Namun, sebagian besar penelitian masih menggunakan pendekatan *supervised learning* dengan proses pelabelan yang dilakukan secara manual atau berdasarkan ambang batas tertentu. Pendekatan tersebut berpotensi menimbulkan subjektivitas karena penentuan label sangat bergantung pada pertimbangan peneliti. Untuk mengurangi permasalahan tersebut, algoritma *K-Means* telah diterapkan secara luas pada berbagai penelitian dalam pengelompokan data media sosial didasarkan pada karakteristik seperti analisis komunikasi opini publik [9], pengelompokan aktivitas dan respons pengguna [10], pemetaan pola interaksi pengguna [11], identifikasi awal komentar yang mengandung kejahatan verbal [12], serta analisis waktu optimal publikasi konten berdasarkan tingkat *engagement* pengguna [13]. Selain itu, penelitian mengenai klasifikasi video *YouTube* juga menunjukkan bahwa hasil klusterisasi menggunakan *K-Means* dapat dimanfaatkan sebagai dasar pembentukan label sebelum proses klasifikasi dilakukan [14]. Di sisi lain, pemilihan algoritma klasifikasi juga menjadi faktor penting dalam menghasilkan model prediksi yang akurat. Sejumlah penelitian membuktikan bahwa *LightGBM* menghasilkan kinerja yang baik pada sejumlah permasalahan klasifikasi berbasis data tabular, seperti sistem rekomendasi [15], prediksi kecanduan media sosial [16], dan prediksi customer churn [17]. Keandalan model tersebut umumnya dievaluasi dengan metode *K-Fold Cross Validation* dalam mengukur stabilitas kinerja model [18]. Meskipun demikian, penelitian yang mengintegrasikan pelabelan otomatis menggunakan *K-Means*, klasifikasi menggunakan *LightGBM*, dan interpretasi model menggunakan *Explainable Artificial Intelligence* (XAI) pada analisis kinerja konten *YouTube* masih relatif terbatas. Selain itu, sifat *LightGBM* sebagai model *black-box* menyebabkan proses pengambilan keputusan model sulit dipahami apabila tidak disertai dengan metode interpretasi yang memadai.

Berdasarkan celah penelitian tersebut, penelitian ini mengusulkan pendekatan untuk mengklasifikasikan kinerja konten *Mobile Legends* di *YouTube* melalui integrasi pelabelan otomatis menggunakan algoritma *K-Means*, klasifikasi menggunakan *LightGBM*, dan *Explainable Artificial Intelligence* (XAI) memakai *Shapley Additive exPlanations* (SHAP) [19], [20]. Pendekatan ini bertujuan menghasilkan proses pembentukan label yang lebih objektif tanpa bergantung pada label manual, sekaligus membangun model klasifikasi yang memiliki tingkat akurasi yang baik dan mudah diinterpretasikan. Pelabelan otomatis dilakukan dengan mengelompokkan video berdasarkan metrik kinerja yang terdiri atas *Jumlah View*, *Jumlah Like*, *Jumlah Comment* dan *Engagement Rate*. Selanjutnya, hasil pelabelan tersebut digunakan sebagai target pada proses klasifikasi menggunakan fitur *Panjang Judul*, *Usia Video*, *Durasi*, dan *Jumlah Subscriber*. Dalam meningkatkan transparansi, penelitian menggunakan SHAP untuk menjelaskan kontribusi setiap fitur terhadap hasil klasifikasi, baik pada tingkat global maupun lokal. Melalui pendekatan tersebut, penelitian yang dilakukan diharapkan bisa memberikan kontribusi untuk pengembangan metode klasifikasi kinerja konten *YouTube* yang tidak hanya menghasilkan model klasifikasi dengan performa yang baik, serta bisa memberikan interpretasi yang mudah dipahami terhadap setiap faktor yang berpengaruh pada kinerja video.

2. METODOLOGI PENELITIAN

Alur penelitian dirancang secara sistematis melalui lima tahapan yaitu pengumpulan data, prapemrosesan data, pelabelan otomatis menggunakan *K-Means*, klasifikasi menggunakan *LightGBM* beserta evaluasi model dan *Explainable Artificial Intelligence* menggunakan SHAP. Tahapan penelitian divisualisasikan pada Gambar 1.



Gambar 1. Alur Penelitian Klasifikasi Kinerja *Mobile Legends* di *Youtube*

2.1 Pengumpulan Data

Dataset penelitian diperoleh melalui *YouTube Data API v3* dengan bahasa *python* menggunakan *google colab* [1]. Data diambil dari periode 20 juli 2025 hingga 20 juli 2026 untuk merepresentasikan kinerja video *Mobile Legends* dalam rentang waktu satu tahun dengan format video horizontal. Menggunakan kata kunci sebagai berikut *Mobile Legends*, *MLBB*, *Mobile Legends Indonesia*, *MLBB Gameplay*, *Top Global Miya*, *Jess No Limit Mobile Legends*, dan *Gameplay Proplayer MLBB*. Dataset terdiri atas 3.052 hasil dari *Scraping* menggunakan 6 *API_Keys*. Video memiliki sebelas variabel, yaitu *Id_Video*, *Judul*, *Panjang_Judul*, *Tanggal_Publish*, *Usia_Video*, *Durasi*, *Jumlah_Subscriber*, *Jumlah_View*, *Jumlah_Like*, *Jumlah_Comment*, dan *Engagement_Rate*. Variabel-variabel tersebut dipilih karena penelitian sebelumnya menunjukkan bahwa karakteristik video dan tingkat keterlibatan pengguna memiliki hubungan dengan kinerja suatu konten di *YouTube* [2], [3], [5], [6], [11]. Contoh dataset disajikan pada Tabel 1.

Tabel 1. Contoh Dataset Mobile Legends

Variabel	Video A	Video B	Video C
Id_Video	jHnMFjhY1uE	EDJJXovtHQM	rnuaz48Wa0E
Judul	Ketemu Top Global Pake Meta Aneh, Ga Dikasih Kill Sama Sekali! - Mobile Legends	Namatin Mobile Legends tapi Hero Sahabat Only	MIYA HEALING BUILD TUTORIAL 2026!?? (HACK REGEN EFFECT!??) - Mobile Legends
Panjang_Judul	80	45	74
Tanggal_Publish	13/08/2025 13:44	22/05/2026 14:03	05/03/2026 18:01
Usia_Video	343 Hari	61 Hari	138 Hari
Durasi	671 Detik	1.694 Detik	1.438 Detik
Jumlah_Subscriber	6.880.000	188.000	58.700
Jumlah_View	1.068.750	47.516	2.403
Jumlah_Like	11.858	847	66
Jumlah_Comment	714	78	4
Engagement_Rate	1.176	1.947	2.913

2.2 Prapemrosesan Data

Tahap prapemrosesan dilakukan bertujuan untuk mengoptimalkan kualitas data sebelum dilakukan analisis. Tahapan dilakukan pemeriksaan duplikat data berdasarkan *Id_Video*, pemeriksaan *missing value*, penanganan nilai yang hilang menggunakan imputasi median apabila diperlukan, serta analisis *outlier*. Berikutnya dilakukan proses transformasi logaritmik terhadap variabel *Jumlah_View*, *Jumlah_Like*, dan *Jumlah_Comment* untuk mengurangi tingkat kemencengan distribusi data. Setelah itu, proses standarisasi menggunakan *StandardScaler* diterapkan pada variabel *Jumlah_View*, *Jumlah_Like*, *Jumlah_Comment*, dan *Engagement_Rate* agar seluruh variabel mempunyai skala yang seragam sebelum digunakan pada proses klusterisasi menggunakan *K-Means* [11], [12], [13], [14].

Transformasi logaritmik menggunakan fungsi $\log(1 + x)$ agar nilai nol tetap dapat diproses tanpa menghasilkan nilai tak terdefinisi. Persamaan yang digunakan adalah:

$$X' = \log(1 + X) \quad (1)$$

dengan (X) merupakan nilai asli suatu variabel dan (X) merupakan nilai setelah transformasi.

Setelah transformasi, standarisasi dilakukan menggunakan metode *Z-score* dengan persamaan:

$$Z = \frac{X - \mu}{\sigma} \quad (2)$$

dengan (X) adalah nilai asli, (μ) adalah rata-rata variabel, dan (σ) adalah standar deviasi. Standarisasi dilakukan agar fitur dengan rentang nilai besar tidak mendominasi proses perhitungan jarak pada *K-Means*.

2.3 K-Means

Pelabelan otomatis dilakukan menggunakan algoritma *K-Means*. Variabel yang digunakan dalam proses klusterisasi terdiri atas *Jumlah_View*, *Jumlah_Like*, *Jumlah_Comment*, dan *Engagement_Rate* karena seluruh variabel tersebut merepresentasikan kinerja video berdasarkan tingkat interaksi pengguna. Penggunaan *K-Means* dalam penelitian media sosial telah banyak dimanfaatkan untuk mengelompokkan pola interaksi pengguna, aktivitas media sosial, maupun karakteristik keterlibatan *audiens* [10], [11], [12], [13]. Penentuan jumlah klaster dilakukan dengan mengevaluasi nilai *WCSS* (*Within Cluster Sum of Squares*) menggunakan metode *Silhouette Score*, *Elbow*, *Davies-Bouldin Index* dan *Calinski-Harabasz Index*. Jumlah klaster terbaik kemudian digunakan untuk membentuk label kinerja yang selanjutnya

menjadi target pada proses klasifikasi. Pendekatan ini mengacu pada penelitian yang memanfaatkan hasil klusterisasi sebagai dasar pembentukan label sebelum dilakukan proses klasifikasi [14].

Jumlah cluster diuji pada rentang $K=2$ sampai $K=10$. Salah satu ukuran yang digunakan adalah *WCSS (Within Cluster Sum of Squares)*, yang memiliki bentuk:

$$WCSS = \sum_{j=1}^K \sum_{x_i \in C_j} ||x_i - c_j||^2 \quad (3)$$

Nilai *WCSS* akan menurun ketika jumlah cluster bertambah. *Metode Elbow* digunakan untuk memilih titik ketika penurunan *WCSS* mulai melandai.

Silhouette Score digunakan untuk mengukur kualitas pemisahan cluster. Nilai *silhouette* suatu data dirumuskan sebagai:

$$S(i) = \frac{b(i) - a(i)}{\max(a(i), b(i))} \quad (4)$$

dengan $a(i)$ adalah rata-rata jarak data (i) terhadap anggota dalam cluster yang sama, sedangkan $b(i)$ adalah rata-rata jarak minimum terhadap anggota pada cluster terdekat. Nilai *silhouette* berada pada rentang (-1) hingga (1) , dimana nilai yang semakin tinggi menunjukkan pemisahan cluster yang semakin baik.

Calinski-Harabasz Index dihitung berdasarkan rasio dispersi antarcluster dan dispersi dalam cluster:

$$CH = \frac{Tr(B_k)/(K-1)}{Tr(W_k)/(n-K)} \quad (5)$$

dengan (B_k) merepresentasikan dispersi antarcluster, (W_k) merupakan dispersi dalam cluster, (K) adalah jumlah cluster, dan (n) adalah jumlah data. Nilai *CH* yang lebih tinggi menunjukkan struktur cluster yang lebih baik.

Davies-Bouldin Index digunakan untuk mengukur kemiripan antarcluster:

$$DB = \frac{1}{K} \sum_{i=1}^K \max_{j \neq i} \left(\frac{S_i + S_j}{M_{ij}} \right) \quad (6)$$

dengan (S_i) dan (S_j) merupakan rata-rata jarak anggota terhadap *centroid* masing-masing cluster, sedangkan (M_{ij}) adalah jarak antarcentroid. Nilai *Davies-Bouldin* yang lebih rendah menunjukkan pemisahan cluster yang lebih baik.

2.2 LightGBM

Tahap klasifikasi dilakukan menggunakan algoritma *LightGBM* dengan memanfaatkan label hasil klusterisasi sebagai variabel target. Variabel prediktor yang digunakan terdiri atas *Panjang_Judul*, *Usia_Video*, *Durasi*, dan *Jumlah_Subscriber*. Sedangkan variabel yang digunakan pada proses klusterisasi tidak lagi digunakan sebagai fitur klasifikasi untuk menghindari data *leakage*. Dataset dibagi menjadi 80% menjadi data pelatihan dan 20% menjadi data pengujian dengan metode *Stratified Train-Test Split* supaya proporsi setiap kelas tetap terjaga. Algoritma *LightGBM* dipilih berdasarkan pada kemampuannya untuk menangani data tabular secara efisien dengan akurasi yang tinggi pada berbagai penelitian sebelumnya [15], [16], [17]. Secara umum, prediksi model gradient boosting dapat dituliskan sebagai:

$$y_i = \sum_{m=1}^M f_m(x_i) \quad (7)$$

dengan (y_i) merupakan prediksi untuk data ke- (i) , (M) merupakan jumlah tree, dan (f_m) merupakan fungsi pohon keputusan ke- (m) .

2.2 Evaluasi Model

Evaluasi untuk kinerja metode ini dilakukan dengan empat parameter pengujian, yaitu *Accuracy*, *Precision*, *Recall*, dan *F1-Score* yang diperoleh dari *Classification Report* serta divisualisasikan melalui *Confusion Matrix*. Selain itu, dilakukan juga pengujian *10-Fold Cross Validation* terhadap data pelatihan dengan *StratifiedKFold* untuk menghitung konsistensi dan kemampuan generalisasi model. Penggunaan *k-Fold Cross Validation* telah banyak digunakan sebagai metode evaluasi untuk memperoleh estimasi kinerja model yang lebih stabil dibandingkan pengujian menggunakan satu kali pembagian data [18].

Accuracy digunakan untuk mengukur proporsi prediksi yang benar terhadap seluruh data:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (8)$$

Precision menunjukkan proporsi prediksi positif yang benar:

$$Precision = \frac{TP}{TP + FP} \quad (9)$$

Recall menunjukkan kemampuan model dalam mengenali data positif yang sebenarnya:

$$Recall = \frac{TP}{TP + FN} \quad (10)$$

Sedangkan *F1-Score* merupakan rata-rata harmonik antara *precision* dan *recall*:

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (11)$$

dengan (TP), (TN), (FP), dan (FN) masing-masing menunjukkan true positive, true negative, false positive, dan false negative.

10-Fold Cross Validation berbasis *StratifiedKFold* pada data pelatit diperoleh melalui rata-rata performa setiap fold. Rata-rata metrik dihitung dengan:

$$M = \frac{1}{K} \sum_{i=1}^k M_i \quad (12)$$

sedangkan standar deviasi dihitung dengan:

$$SD = \sqrt{\frac{\sum_{i=1}^k (M_i - M)^2}{k - 1}} \quad (13)$$

dengan (M_i) merupakan nilai metrik pada fold ke-(i), (M) adalah rata-rata metrik, dan ($k=10$). Semakin kecil standar deviasi, semakin konsisten performa model antarfold.

2.2 Shapley Additive exPlanations (SHAP)

Dalam meningkatkan transparansi pada model, penelitian ini memanfaatkan *Shapley Additive exPlanations* (SHAP) sebagai pendekatan *Explainable Artificial Intelligence* (XAI). SHAP dimanfaatkan untuk menganalisis kontribusi setiap fitur terhadap hasil prediksi model, baik pada tingkat global maupun lokal. Interpretasi model dilakukan menggunakan *SHAP Bar Plot*, *SHAP Summary Plot*, *SHAP Waterfall Plot* dan *SHAP Dependence Plot*. Pendekatan ini memungkinkan proses pengambilan keputusan pada model *LightGBM* diinterpretasikan secara lebih jelas sehingga faktor-faktor yang memengaruhi klasifikasi kinerja video dapat dianalisis secara komprehensif [19], [20]. Nilai SHAP didasarkan pada konsep nilai Shapley yang menghitung kontribusi marginal setiap fitur terhadap prediksi model. Secara umum, fungsi SHAP dapat dituliskan sebagai:

$$\phi_j = \sum_{S \subseteq N \setminus \{j\}} \frac{|S|! (M - |S| - 1)!}{M!} [f(S \cup \{j\}) - f(S)] \quad (14)$$

dengan (ϕ_j) merupakan nilai kontribusi fitur ke-(j), (S) adalah subset fitur, (M) adalah jumlah seluruh fitur, dan ($f(S)$) merupakan prediksi model ketika hanya subset fitur (S) digunakan. Prediksi model kemudian dapat dijelaskan sebagai:

$$f(x) = \phi_0 + \sum_{j=1}^M \phi_j \quad (15)$$

dengan (ϕ_0) merupakan nilai dasar dan (ϕ_j) merupakan kontribusi masing-masing fitur terhadap prediksi.

3. HASIL DAN PEMBAHASAN

3.1 Karakteristik Dataset

Dataset yang digunakan sebanyak 3.052 video bertema *Mobile Legends* yang memiliki 11 variabel, yaitu *Id_Video*, *Judul*, *Panjang_Judul*, *Tanggal_Publish*, *Usia_Video*, *Durasi*, *Jumlah_Subscriber*, *Jumlah_View*, *Jumlah_Like*, *Jumlah_Comment*, dan *Engagement_Rate*. Variabel *Video_ID* dan *Title* bertipe data object, sedangkan sembilan variabel lainnya bertipe numerik sehingga dapat digunakan dalam proses analisis statistik maupun pemodelan *machine learning*.

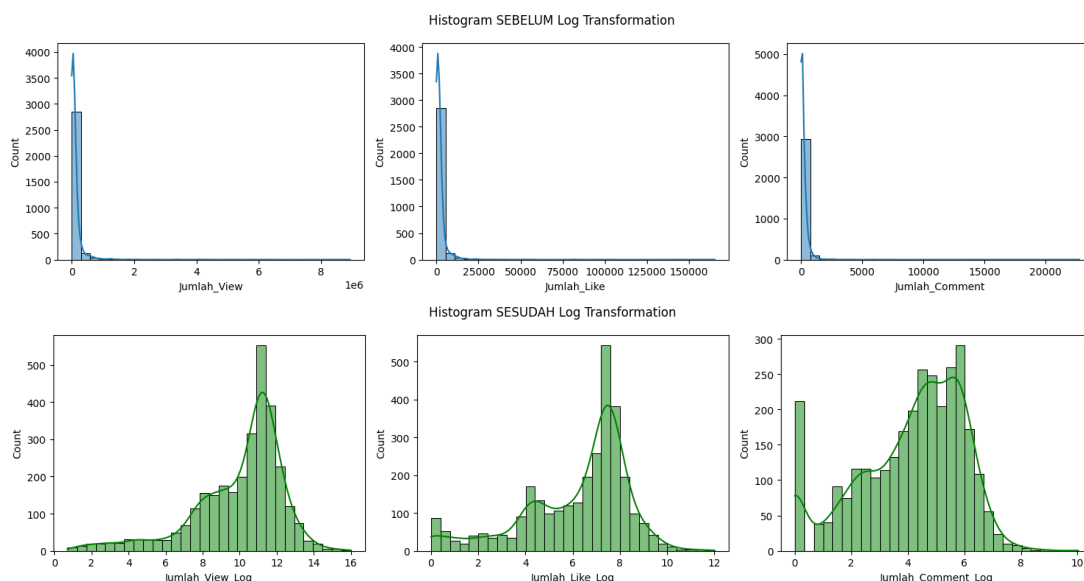
Tabel 2. Statistik Deskriptif Dataset

Variabel	Min	Max	Mean	Std. Deviasi
Panjang_Judul	12	117	75,57	19,62
Usia_Video	1	367	125,08	98,35
Durasi	180	1800	937,06	345,73
Jumlah_Subscriber	0	54.600.000	2.995.462	7.644.048
Jumlah_View	1	8.926.216	108.329	323.201
Jumlah_Like	0	165.248	2.059	5.746
Jumlah_Comment	0	22.773	204	541
Engagement_Rate	0	233.333	2.650	5.021

Berdasarkan Tabel 2, rata-rata panjang judul video adalah 75,57 karakter, usia video 125,08 hari, durasi 937,06 detik, dan jumlah subscriber kanal sekitar 2,99 juta. Sementara itu, rata-rata jumlah view sebanyak 108.329, jumlah like sebanyak 2.059, dan jumlah komentar sebanyak 204. Sehingga dataset memiliki karakteristik yang beragam dan bisa digunakan untuk proses klusterisasi serta klasifikasi kinerja video.

3.2 Pemrosesan Data

Tahap pembersihan data dilakukan sebagai proses awal untuk memastikan kualitas, konsistensi, dan kelayakan data sebelum digunakan untuk proses pemodelan. Hasil menunjukkan tidak ada data duplikat berdasarkan *Id_Video* maupun *missing value* terhadap semua variabel, sehingga jumlah data tetap sebanyak 3.052 video. Hasil ini menunjukkan bahwa dataset memiliki kualitas yang baik dan siap digunakan pada tahap transformasi data serta pembentukan model.

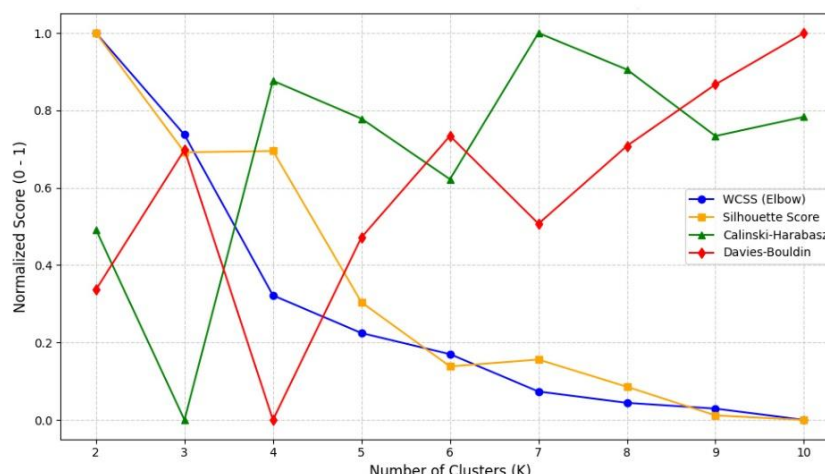


Gambar 2. Histogram Sebelum dan Sesudah Transformasi Log

Berdasarkan Gambar 2. Transformasi log diterapkan pada variabel *Jumlah_View*, *Jumlah_Like*, dan *Jumlah_Comment* sebelum proses klusterisasi menggunakan *K-Means*. Proses ini menghasilkan empat variabel yang digunakan sebagai masukan *K-Means*, yaitu *Jumlah_View_Log*, *Jumlah_Like_Log*, *Jumlah_Comment_Log*, serta *Engagement_Rate*, dengan ukuran data sebesar 3.052×4 . Transformasi dan standarisasi tersebut dilakukan untuk meningkatkan stabilitas proses klusterisasi serta mengurangi dominasi variabel yang memiliki rentang nilai sangat besar.

3.3 Hasil K-Means Clustering

Proses penentuan jumlah kluster dilakukan menggunakan skenario K=2 hingga K=10 menggunakan empat metrik evaluasi internal, yaitu *Silhouette Score*, *Elbow Method*, *Davies-Bouldin Index*, dan *Calinski-Harabasz Index*. Divisualisasikan pada Gambar 3.



Gambar 3. Grafik *Silhouette Score*, *Elbow Method*, *Davies-Bouldin Index* dan *Calinski-Harabasz Index*

Berdasarkan Gambar 3, hasil kurva WCSS mengalami penurunan tajam hingga K=4 kemudian mulai melandai, yang mengindikasikan titik *elbow* berada pada kisaran K=3-4. Sementara itu, *Silhouette Score* mencapai nilai tertinggi pada K=2, menunjukkan jumlah kluster tersebut menghasilkan tingkat kohesi intra-kluster dan separasi antar-kluster yang paling baik. *Davies-Bouldin Index* menunjukkan performa terbaik pada K=4 sedangkan *Calinski-Harabasz Index* memperoleh nilai maksimum pada K=7. Meskipun beberapa metrik mengusulkan jumlah kluster yang berbeda, penelitian ini menetapkan K=2 sebagai jumlah kluster optimal karena menghasilkan nilai *Silhouette Score* tertinggi serta sesuai dengan tujuan penelitian dalam membentuk dua kategori kinerja, yaitu Kinerja Tinggi dan Kinerja Rendah, sehingga label yang dihasilkan lebih mudah diinterpretasikan dan digunakan pada tahap klasifikasi selanjutnya.

Tabel 3. Hasil Validasi Cluster

Pengujian	Nilai	Interpretasi
Silhouette Score	0,5805	Struktur cluster cukup baik
Mann-Whitney U Test	$p = 1,147 \times 10^{-7}$	Perbedaan cluster signifikan

Berdasarkan Tabel 3, hasil validasi klusterisasi dievaluasi melalui pengujian statistik dan kualitas cluster. Nilai *Silhouette Score* mencapai 0,5805 menunjukkan hasil klusterisasi mempunyai pemisahan antar cluster yang cukup baik. Selanjutnya, *Mann-Whitney U Test* menghasilkan *p-value* sebesar $1,147 \times 10^{-7}$ jauh lebih kecil dari 0,05. Hasil tersebut menandakan bahwa adanya perbedaan signifikan antara cluster Rendah dan Tinggi yang mengindikasikan pembagian cluster merepresentasikan perbedaan karakteristik kinerja video, sehingga bisa digunakan sebagai pembentukan label otomatis pada tahap klasifikasi.

Tabel 4. Profil Hasil Cluster K-Means

Cluster	Jumlah Video	Jumlah View	Jumlah Like	Jumlah Comment	Engagement Rate
0	841	3.405,10	53,80	8,54	3.042,30
1	2.211	148.239,59	2.822,79	278,84	2.501,86

Berdasarkan Tabel 4, hasil Cluster *K-Means* dengan K=2 berhasil mengelompokkan 3.052 video ke dalam dua kluster, yang terdiri atas *Cluster 0* sebanyak 841 video dan *Cluster 1* sebanyak 2.211 video. Berdasarkan nilai rata-rata *Jumlah View*, *Jumlah Like*, dan *Jumlah Comment*, Cluster 1 menunjukkan kinerja yang lebih tinggi sehingga diberi label Tinggi, sedangkan Cluster 0 diberi label Rendah. Meskipun rata-rata *Engagement Rate* pada Cluster 0 sedikit lebih tinggi, penamaan cluster didasarkan pada keseluruhan metrik kinerja, terutama *Jumlah View* sebagai indikator utama video di *YouTube*. Oleh karena itu, pelabelan yang menggunakan kombinasi seluruh metrik memberikan representasi kinerja video yang lebih objektif.

Tabel 5. Distribusi Label Kinerja

Label	Jumlah Data	Persentase (%)
Kinerja Tinggi	2.211	72,44
Kinerja Rendah	841	27,56
Total	3.052	100,00

Berdasarkan Tabel 5, hasil klusterisasi kemudian diubah menjadi dua label klasifikasi, yaitu Kinerja Rendah dan Kinerja Tinggi. Sebanyak 2.211 video atau 72,44% termasuk kategori Kinerja Tinggi, sedangkan 841 video atau 27,56% termasuk kategori Kinerja Rendah. Meskipun distribusi kelas tidak seimbang, proporsi tersebut tetap dipertahankan melalui

Stratified Train-Test Split sehingga proses pelatihan model tidak terpengaruh. Pelabelan otomatis ini juga mengurangi subjektivitas karena seluruh kelas dibentuk berdasarkan karakteristik data aktual, bukan menggunakan nilai ambang atau *threshold* yang ditentukan secara manual.

3.4 Perbandingan 6 Model

Sebelum tahap pemodelan dipilih, dilakukan perbandingan kinerja enam model klasifikasi menggunakan parameter default, yaitu *LightGBM*, *XGBoost*, *Random Forest*, *Logistic Regression*, *Decision Tree*, dan *Support Vector Machine* (SVM). Semua algoritma dilatih dan diuji dengan dataset yang sama serta skema pembagian 80% sebagai data pelatihan dan 20% sebagai data pengujian. Selanjutnya, kinerja setiap model dievaluasi menggunakan metrik *Accuracy*, *Recall*, *Precision*, dan *F1-score* untuk memastikan perbandingan dilakukan secara konsisten. Hasil perbandingan model disajikan pada Tabel 6.

Tabel 6. Perbandingan Model dengan Parameter Bawaan

Rank	Model	Accuracy	Precision	Recall	F1-Score
1	LightGBM	0,9378	0,9387	0,9378	0,9381
2	Random Forest	0,9362	0,9363	0,9362	0,9362
3	XGBoost	0,9231	0,9243	0,9231	0,9236
4	Decision Tree	0,9149	0,9170	0,9149	0,9156
5	Logistic Regression	0,8412	0,8380	0,8412	0,8392
6	SVM	0,8265	0,8218	0,8265	0,8233

Berdasarkan Tabel 6, Hasil pengujian menunjukkan bahwa *LightGBM* memperoleh kinerja terbaik dengan *Accuracy* mencapai 93,78%, *Precision* mencapai 93,87%, *Recall* mencapai 93,78%, dan *F1-Score* mencapai 93,81%, diikuti oleh *Random Forest* dengan selisih yang sangat kecil. Sebaliknya, algoritma *Logistic Regression* dan SVM menghasilkan kinerja yang lebih rendah dibandingkan metode berbasis *ensemble tree*. Hasil tersebut menunjukkan bahwa karakteristik dataset yang bersifat nonlinier lebih sesuai dimodelkan menggunakan algoritma berbasis pohon keputusan daripada metode linier. Kemampuan *LightGBM* dalam membangun pohon keputusan secara bertahap memungkinkan model menangkap hubungan yang kompleks antarfitur secara lebih efektif dibandingkan algoritma lainnya.

3.5 Hasil Model LightGBM

Kinerja model *LightGBM* dievaluasi menggunakan data pengujian yang berjumlah 611 video untuk mengetahui kemampuan model dalam mengklasifikasikan kinerja konten *Mobile Legends* di *YouTube*. Hasil Model *LightGBM* disajikan pada Tabel 7.

Tabel 7. Hasil Model LightGBM pada Data Pengujian

Metrik	Accuracy	Precision	Recall	F1-Score
Nilai	93,78%	93,87%	93,78%	93,81%

Berdasarkan Tabel 7, model *LightGBM* menunjukkan kinerja klasifikasi pada data pengujian dengan *accuracy* sebesar 93,78% menunjukkan bahwa model mampu mengklasifikasikan hampir seluruh data video *Mobile Legends* ke dalam kategori kinerja yang sesuai. Nilai *precision* sebesar 93,87% mengindikasikan bahwa sebagian besar prediksi yang dihasilkan model telah sesuai dengan kelas sebenarnya, sehingga tingkat kesalahan dalam mengklasifikasikan video ke kelas yang tidak tepat relatif rendah. Nilai *recall* sebesar 93,78% menunjukkan bahwa model bisa mengidentifikasi sebagian besar video pada setiap kelas kinerja, sehingga hanya sebagian kecil data yang tidak berhasil diidentifikasi dengan benar. Keseimbangan antara *precision* dan *recall* tercermin dari *F1-score* sebesar 93,81%, yang memiliki kemampuan konsisten untuk menjaga ketepatan sekaligus kelengkapan hasil klasifikasi. Selisih nilai antarmetrik yang relatif kecil menunjukkan bahwa model memiliki kinerja yang seimbang pada setiap kelas serta mampu melakukan generalisasi dengan baik terhadap data baru.

3.6 Evaluasi Model LightGBM

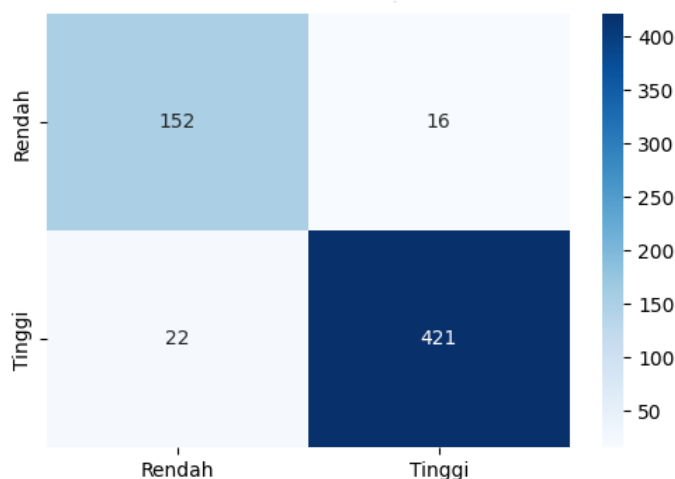
Kinerja model *LightGBM* dievaluasi melalui *Classification Report*, *Confusion Matrix* serta *10-Fold Cross Validation* guna mengevaluasi kapabilitas model dalam memetakan data ke dalam kelas Kinerja Rendah dan Kinerja Tinggi, yang disajikan pada Tabel 8 untuk *Classification Report*, Gambar 4 untuk *Confusion Matrix* dan Tabel 9 untuk *10-Fold Cross Validation*.

Tabel 8. Classification Report per kelas (Kinerja Rendah dan Kinerja Tinggi)

Kelas	Precision	Recall	F1-Score	Support
Kinerja Rendah	0,87	0,90	0,89	168

Kinerja Tinggi	0,96	0,95	0,96	443
Accuracy	-	-	0,94	611
Macro Average	0,92	0,93	0,92	611
Weighted Average	0,94	0,94	0,94	611

Berdasarkan Tabel 8, pada kelas Kinerja Rendah, model menghasilkan *precision* mencapai 0,87, *recall* mencapai 0,90, dan *F1-score* mencapai 0,89, menunjukkan sebagian besar video berkinerja rendah berhasil dikenali dengan baik meskipun masih terdapat sejumlah kecil kesalahan prediksi. Sedangkan pada kelas Kinerja Tinggi, model menghasilkan *precision* mencapai 0,96, *recall* mencapai 0,95, dan *F1-score* mencapai 0,96, mengindikasikan bahwa model memiliki tingkat ketepatan dan kemampuan identifikasi yang sangat tinggi terhadap video berkinerja tinggi. Perbedaan nilai metrik antara kedua kelas dipengaruhi oleh distribusi data yang tidak seimbang, pada kelas Kinerja Tinggi memiliki jumlah data yaitu 443 video lebih banyak dibandingkan kelas Kinerja Rendah yaitu 168 video. Disisi lain, nilai *macro average* yaitu 0,92 hingga 0,93 mengindikasikan bahwa kinerja model tetap baik ketika setiap kelas diberi bobot yang sama, sementara *weighted average* mencapai 0,94 mengindikasikan model tetap bisa mempertahankan kinerja yang konsisten dengan mempertimbangkan proporsi setiap kelas. Secara keseluruhan, hasil pada *classification report* membuktikan model *LightGBM* mempunyai kemampuan yang seimbang untuk mengklasifikasikan kedua kategori kinerja video serta mampu menghasilkan prediksi yang akurat dan andal pada data pengujian.



Gambar 4. *Confusion Matrix LightGBM*

Berdasarkan Gambar 4, pada *Confusion Matrix* menunjukkan bahwa model *LightGBM* memiliki keandalan klasifikasi yang baik dalam membedakan video berkinerja Rendah dan Tinggi. Dari 611 data pengujian, sebanyak 573 data berhasil diklasifikasikan dengan benar, terdiri atas 152 *true negative* dan 421 *true positive*. Sementara itu, hanya terdapat 38 kesalahan klasifikasi, yaitu 16 *false positive* dan 22 *false negative*. Minimnya kesalahan mengonfirmasi keandalan model dalam membedakan kedua kategori kelas dengan baik. Hasil ini konsisten dengan nilai *precision*, *recall*, dan *F1-score* yang tinggi pada kedua kelas, sehingga model tidak hanya akurat, tetapi juga mampu menjaga keseimbangan antara pengenalan kelas serta minimisasi kesalahan klasifikasi. Dengan demikian, model *LightGBM* mempunyai kemampuan generalisasi yang baik dan layak digunakan dalam mengklasifikasikan kinerja konten *Mobile Legends* di *YouTube*.

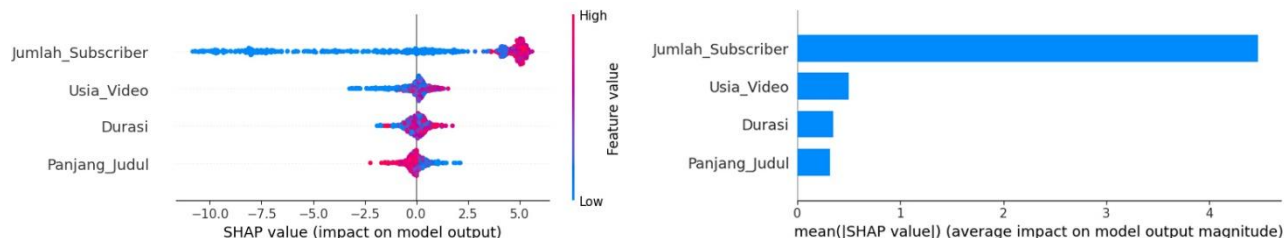
Tabel 9. Hasil 10-Fold Cross Validation

Metrik	Mean	Standar Deviasi ($\pm$)
Accuracy	0,9275	0,0148
Precision	0,9274	0,0148
Recall	0,9275	0,0148
F1-Score	0,9270	0,0150

Berdasarkan Tabel 9, model *LightGBM* juga dievaluasi menggunakan *10-Fold Cross Validation* dengan metode *StratifiedKfold* untuk mengukur kemampuan generalisasi model. Hasil evaluasi menunjukkan bahwa *mean* pada *accuracy* yaitu 92,75%, *precision* yaitu 92,74%, *recall* yaitu 92,75%, dan *F1-score* yaitu 92,70%. Seluruh metrik memiliki standar deviasi sekitar 0,015. Hal ini menunjukkan bahwa variasi kinerja antar fold relatif kecil dan kestabilan model yang baik. Selisih yang kecil antara hasil pengujian dan *cross validation* menunjukkan bahwa model memiliki kemampuan generalisasi yang baik terhadap data baru. Hasil ini menandakan bahwa *LightGBM* mampu menghasilkan kinerja yang stabil dan akurat.

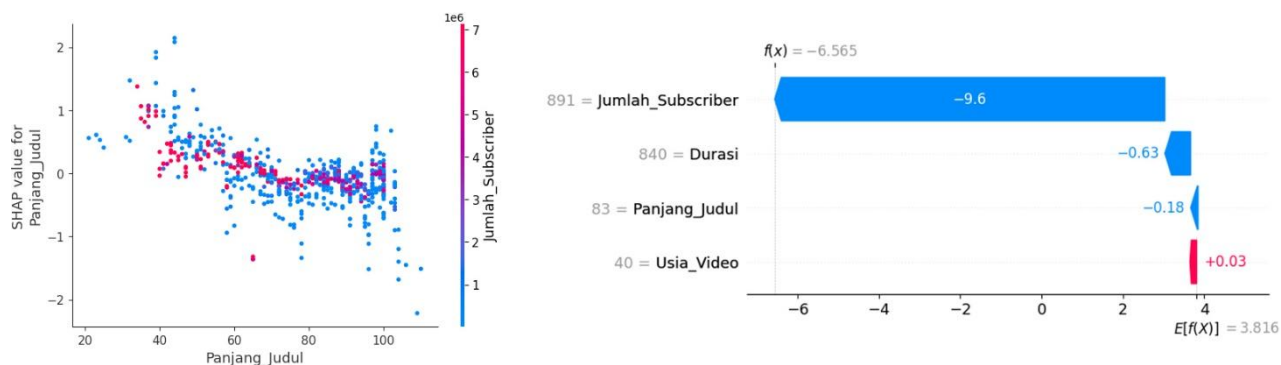
3.7 SHAP Explainability

Penelitian yang dilakukan menerapkan *Shapley Additive exPlanations* (SHAP) untuk menginterpretasikan keputusan *LightGBM* secara lebih transparan. Interpretasi dilakukan melalui *SHAP Bar Plot*, *SHAP Summary Plot*, *SHAP Waterfall Plot* dan *SHAP Dependence Plot*. Yang divisualisasikan melalui Gambar 6 dan Gambar 7.



Gambar 6. *SHAP Summary Plot dan SHAP Bar Plot*

Berdasarkan Gambar 6, hasil interpretasi pada *SHAP Summary Plot* serta *SHAP Bar Plot* mengonfirmasi bahwa *Jumlah_Subscriber* dan *Usia_Video* merupakan dua variabel yang memiliki kontribusi terbesar terhadap prediksi model. Nilai *Jumlah_Subscriber* yang tinggi cenderung menghasilkan nilai SHAP positif sehingga meningkatkan video diklasifikasikan ke dalam kategori Kinerja Tinggi, sementara *Jumlah_Subscriber* yang rendah diklasifikasikan pada prediksi Kinerja Rendah. Sementara pada *Usia_Video*, di mana video yang telah lama dipublikasikan memiliki peluang lebih besar memperoleh akumulasi interaksi pengguna sehingga meningkatkan kinerjanya. Sebaliknya, *Durasi* dan *Panjang_Judul* memberikan pengaruh yang relatif lebih kecil terhadap hasil prediksi. SHAP tidak hanya mengidentifikasi tingkat kepentingan setiap fitur, tetapi juga menjelaskan arah pengaruhnya terhadap keputusan model, sehingga mendukung perannya sebagai pendekatan *Explainable Artificial Intelligence* yang efektif untuk menginterpretasikan model berbasis *ensemble tree* pada tingkat global maupun lokal.



Gambar 7. *SHAP Dependence Plot dan Waterfall Plot*

Berdasarkan Gambar 7, pada *SHAP Dependence Plot* peningkatan *Jumlah_Subscriber* umumnya diikuti oleh peningkatan nilai SHAP yang mengarah pada prediksi Kinerja Tinggi. Namun, hubungan tersebut tidak sepenuhnya linier karena beberapa video dari kanal dengan banyak pelanggan tetap memiliki kinerja rendah. Dimana keberhasilan video tidak hanya ditentukan oleh *Jumlah_Subscriber*, tetapi juga dipengaruhi oleh faktor lain, seperti *Usia_Video*, *Durasi* dan ketertarikan *audiens*. Sementara itu, *SHAP Waterfall Plot* digunakan untuk menjelaskan keputusan model pada tingkat individual dengan memperlihatkan kontribusi setiap fitur terhadap hasil prediksi. Video dengan *Jumlah_Subscriber* yang tinggi dan *Usia_Video* yang lebih lama memberikan kontribusi positif terhadap prediksi Kinerja Tinggi, sedangkan nilai yang rendah mendorong prediksi Kinerja Rendah. Hasil ini menunjukkan bahwa SHAP mampu mengungkap hubungan kompleks antarfitur sekaligus memberikan interpretasi yang rinci dan transparan terhadap setiap keputusan model.

4. KESIMPULAN

Penelitian ini menghasilkan model klasifikasi kinerja konten *Mobile Legends* di *YouTube* melalui integrasi *K-Means*, *LightGBM*, dan *Shapley Additive exPlanations* (SHAP). *K-Means* menghasilkan pelabelan otomatis berdasarkan *Jumlah_View*, *Jumlah_Like*, *Jumlah_Comment*, dan *Engagement_Rate*, sehingga pembentukan label lebih objektif dan mengurangi subjektivitas. Hasil evaluasi internal menetapkan dua kluster yang diinterpretasikan sebagai kelas Kinerja Rendah dan Kinerja Tinggi. Pada tahap klasifikasi, *LightGBM* dipilih sebagai model akhir karena memberikan performa terbaik dibandingkan algoritma pembandingan, dengan *accuracy* mencapai 93,78%, *precision* mencapai 93,87%, *recall*

mencapai 93,78%, dan *F1-score* mencapai 93,81% pada data pengujian. Pengujian *10-Fold Cross Validation* menghasilkan rata-rata *accuracy* sebesar 92,75% dengan standar deviasi rendah, yang menunjukkan kemampuan generalisasi model yang baik dan stabil. Penerapan SHAP meningkatkan interpretabilitas melalui penjelasan kontribusi setiap fitur terhadap hasil prediksi sehingga keputusan model dapat dipahami secara lebih transparan. Dengan demikian, kombinasi *K-Means*, *LightGBM*, dan SHAP efektif sebagai klasifikasi kinerja konten *YouTube*. Penelitian ini masih terbatas pada dataset *Mobile Legends* dan empat variabel prediktor, yaitu *Panjang Judul*, *Usia Video*, *Durasi*, dan *Jumlah Subscriber*. Penelitian selanjutnya disarankan memperluas dataset, menambahkan variabel relevan seperti waktu unggah, jenis konten, atau karakteristik kanal, serta membandingkan *LightGBM* dengan pendekatan *deep learning* supaya memperoleh model yang lebih komprehensif dan memiliki kemampuan prediksi yang lebih baik.

REFERENCES

- [1] A. Efstratiou, "On YouTube Search API Use in Research," *Association for Computing Machinery (ACM)*, pp. 919–927, Oct. 2025, doi: 10.1145/3730567.3764492.
- [2] S. Yang, D. Brossard, D. A. Scheufele, and M. A. Xenos, "The science of YouTube: What factors influence user engagement with online science videos?," *Journal PLoS ONE*, vol. 17, no. 5, pp. 1–19, May 2022, doi: 10.1371/journal.pone.0267697.
- [3] C. Violot, T. Elmas, I. Bilogrevic, and M. Humbert, "Shorts vs. Regular Videos on YouTube: A Comparative Analysis of User Engagement and Content Creation Trends," *WEBSCI 24: Proceedings of the 16th ACM Web Science Conference*, no. 1, pp. 213–223, May 2024, doi: 10.1145/3614419.3644023.
- [4] N. Charolin and A. Wedhasmara, "Evaluasi Keterlibatan Masyarakat di YouTube Pemerintah: Durasi, Judul, Dialog, Konten, Emosi," *Journal of Business and Information System*, vol. 6, no. 1, pp. 131–144, Jun. 2024, doi: 10.36067/jbis.v6i1.234.
- [5] S. W. Tampubolon and P. Dirgantara, "Pengaruh Konten Youtube Oura Gaming Terhadap Pemenuhan Kebutuhan Informasi Mobile Legends," *Jurnal Ilmu Komunikasi UHO : Jurnal Penelitian Kajian Ilmu Sosial dan Informasi*, vol. 8, no. 4, pp. 684–694, Oct. 2023, doi: 10.52423/jikuho.v8i4.142.
- [6] C. O. G. A. D. H. Uray, Vika, and C. Cahyaningtias, "Menelaah Kesuksesan Konten YouTube yang Trending Analisis Data Video Trending di YouTube," *Informatik: Jurnal Ilmu Komputer*, vol. 20, no. 2, pp. 84–92, Aug. 2024, doi: 10.52958/iftk.v20i2.7517.
- [7] M. F. A. Mufid, E. Daniati, and A. Nugroho, "Analisis Sentimen Ulasan Aplikasi Mobile Legends Berbasis Pelabelan IndobERT dan SMOTE dengan Komparasi Algoritma Klasifikasi Naive Bayes dan SVM," *Management of Information System Journal*, vol. 4, no. 3, pp. 151–159, Jul. 2026, doi: 10.47065/mis.v2i3.2783.
- [8] M. Alfudola, N. Suarna, and I. Ali, "Klasifikasi Pemilihan Tipe Hero Mobile Legends Terhadap Minat Pemain Menggunakan Algoritma Naive Bayes Studi Kasus: Komunitas Game Mobile Legends Kota Cirebon," *JATI: Jurnal Mahasiswa Teknik Informatika*, vol. 7, no. 2, pp. 1269–1273, Apr. 2023, doi: 10.36040/jati.v7i2.6541.
- [9] M. Qi, J. Q. Zhao, and Y. Feng, "An optimized public opinion communication system in social media networks based on K-means cluster analysis," *Heliyon*, vol. 10, no. 24, pp. 1–12, Nov. 2024, doi: 10.1016/j.heliyon.2024.e40033.
- [10] M. N. Afrilia, N. Rahaningsih, R. D. Dana, and N. D. Nuris, "Optimasi Analisis Clustering Untuk Aktivitas Dan Respon Pengguna Media Sosial Dengan K-Means," *JATI: Jurnal Mahasiswa Teknik Informatika*, vol. 8, no. 1, pp. 148–155, Feb. 2024, doi: 10.36040/jati.v8i1.8334.
- [11] R. Y. Simanullang and M. Iqbal, "Pengelompokan Pola Interaksi Pengguna Media Sosial Menggunakan Algoritma K-Means untuk Pemetaan Aktivitas Online," *Journal of Informatics Management and Information Technology*, vol. 6, no. 1, pp. 37–43, Jan. 2026, doi: 10.47065/jimat.v6i1.954.
- [12] A. R. Lahitani, A. N. Zhafarina, N. S. W. Oktavia, and N. Jariyah, "Pemetaan Topik Pembicaraan Pada Komentar Live Youtube Menggunakan K-Means Clustering sebagai Identifikasi awal Kejahatan Verbal Cyberbullying," *JTE UNIBA: Jurnal Teknik Elektro Uniba*, vol. 8, no. 2, pp. 399–403, Apr. 2024, doi: 10.36277/jteuniba.v8i2.253.
- [13] Fathoni, N. A. Basulina, S. Gustiani, S. Amalia, A. Sabila, and A. Ibrahim, "Identifikasi Waktu Optimal Posting Terhadap Pola Engagement Sosial Media Menggunakan Metode K-Means Clustering," *JATI: Jurnal Mahasiswa Teknik Informatika*, vol. 9, no. 4, pp. 6237–6244, Aug. 2025, doi: 10.36040/jati.v9i4.14027.
- [14] S. K. Sitanggang, F. R. Umbara, and H. Ashaury, "Klasifikasi Video Pada Media Sosial Youtube Dengan Menggunakan Metode K-Means Dan Support Vector Machine," *Jurnal Locus Penelitian dan Pengabdian*, vol. 2, no. 10, pp. 1027–1032, Oct. 2023, doi: 10.58344/locus.v2i10.1732.
- [15] W. A. Aristawidya and M. Rahardi, "A Hybrid Approach to Music Recommendations Based on Audio Similarity Using Autoencoder and LightGBM," *Journal of Applied Informatics and Computing (JAIC)*, vol. 9, no. 6, pp. 3191–3197, Dec. 2025, doi: 10.30871/jaic.v9i6.10516.
- [16] R. Jannah and R. Prathivi, "Analisa Performa Metode Lightgbm Untuk Prediksi Kecanduan Media Sosial," *Jurnal Transformatika*, vol. 23, no. 2, pp. 164–173, Jan. 2026, doi: 10.26623/transformatika.v23i2.13165.
- [17] A. N. Royana, Y. V. Via, and C. A. Putra, "Evaluasi Kinerja Lightgbm Dan Catboost Untuk Prediksi Churn Berdasarkan Dataset Pelanggan Layanan Streaming Musik," *JATI: Jurnal Mahasiswa Teknik Informatika*, vol. 9, no. 4, pp. 6977–6982, Aug. 2025, doi: 10.36040/jati.v9i4.14358.
- [18] Moch. A. Aprihartha and I. Idham, "Optimization of Classification Algorithms Performance with k-Fold Cross Validation," *Eigen Mathematics Journal*, vol. 7, no. 2, pp. 61–66, Dec. 2024, doi: 10.29303/emj.v7i2.212.
- [19] A. M. Salih *et al.*, "A Perspective on Explainable Artificial Intelligence Methods: SHAP and LIME," *Advanced Intelligent Systems*, vol. 7, no. 1, pp. 1–8, Jan. 2025, doi: 10.1002/aisy.202400304.
- [20] Y. F. Saputra *et al.*, "Ensemble Learning untuk Model Prediksi Risiko Preeklamsia dan Explainable Ai Berbasis SHAP," *Metik Jurnal*, vol. 10, no. 1, p. 2026, Jun. 2026, doi: 10.47002/kp377403.