



## **Prediksi Risiko Ketergantungan Smartphone Pada Anak Usia Dini Menggunakan Algoritma Random Forest**

Kamia Pertiwi<sup>1</sup>, Usman<sup>2</sup>, Ilyas<sup>3</sup>

<sup>1,2,3</sup> Program Studi Sistem Informasi, Fakultas Teknik dan Ilmu Komputer, Universitas Islam Indragiri

<sup>1</sup>[kamiapertiwi@gmail.com](mailto:kamiapertiwi@gmail.com), <sup>2</sup>[Usmanovsky13411@gmail.com](mailto:Usmanovsky13411@gmail.com)\*, <sup>3</sup>[daengilyas@gmail.com](mailto:daengilyas@gmail.com)\*

### **Abstrak**

Penggunaan smartphone pada anak usia dini semakin meningkat dan dapat menimbulkan risiko ketergantungan apabila tidak disertai pengawasan orang tua yang memadai. Penelitian ini bertujuan membangun model prediksi risiko ketergantungan smartphone pada anak usia dini menggunakan algoritma Random Forest serta membandingkan kinerjanya dengan Support Vector Machine dan K-Nearest Neighbor. Penelitian menerapkan pendekatan kuantitatif menggunakan data primer dari kuesioner yang diisi oleh 347 orang tua. Variabel penelitian meliputi umur anak, durasi dan frekuensi penggunaan smartphone, usia pertama mengenal smartphone, tujuan penggunaan, serta pola pengawasan orang tua. Tahapan penelitian mencakup pengumpulan data, pembersihan dan transformasi data, pelabelan tingkat risiko, pembagian data latih dan data uji dengan rasio 80:20, penanganan ketidakseimbangan kelas, pelatihan model, serta evaluasi menggunakan confusion matrix, accuracy, precision, recall, dan F1-score. Hasil pengujian menunjukkan bahwa Random Forest memperoleh akurasi 87,14%, lebih tinggi dibandingkan K-Nearest Neighbor sebesar 69% dan Support Vector Machine sebesar 47%. Analisis feature importance menunjukkan bahwa pengawasan orang tua merupakan faktor paling dominan, diikuti umur anak dan tujuan penggunaan smartphone. Temuan ini menunjukkan bahwa Random Forest mampu menghasilkan prediksi yang lebih baik pada data penelitian, meskipun kemampuan model dalam mengenali kelas risiko rendah masih perlu ditingkatkan. Model yang dihasilkan berpotensi digunakan sebagai pendukung deteksi dini bagi orang tua dan pihak terkait untuk merancang pengawasan, pendampingan, serta langkah pencegahan ketergantungan smartphone secara lebih tepat pada anak usia dini.

**Kata Kunci:** Anak Usia Dini, Ketergantungan Smartphone, Machine Learning, Prediksi Risiko, Random Forest

### **1. Pendahuluan**

Perkembangan teknologi informasi kini telah meningkatkan penggunaan smartphone pada anak usia dini. Walaupun bermanfaat sebagai media pembelajaran dan hiburan, penggunaan smartphone yang berlebihan tanpa pengawasan orang tua dapat meningkatkan risiko ketergantungan serta berdampak pada perkembangan sosial, emosional, dan perilaku anak. Berdasarkan data pada Badan Pusat Statistik tahun 2023, sebanyak 33,44% anak usia dini di Indonesia telah menggunakan gadget, dengan persentase 25,50% pada usia 0–4 tahun dan 52,76% pada usia 5–6 tahun. Kondisi ini menunjukkan betapa pentingnya deteksi dini terhadap risiko ketergantungan smartphone (Khairani et al., 2024).

Penelitian terdahulu telah banyak menerapkan *machine learning* pada berbagai bidang, akan tetapi penerapannya untuk memprediksi ketergantungan smartphone pada anak usia dini masih terbatas. Sebagian besar penelitian terdahulu hanya bersifat deskriptif dan belum mengintegrasikan faktor penggunaan smartphone serta pengawasan orang tua dalam model prediksi. (Angkasa et al., 2026).

Penelitian ini juga mengajukan hipotesis bahwa algoritma Random Forest mampu memprediksi risiko ketergantungan smartphone pada anak usia dini dengan baik berdasarkan karakteristik penggunaan smartphone dan pengawasan orang tua. Penelitian menggunakan data dari 347 responden dengan variabel umur anak, durasi penggunaan smartphone, frekuensi penggunaan, usia pertama mengenal smartphone, tujuan penggunaan smartphone, dan pola pengawasan orang tua. (Hr et al., 2026)

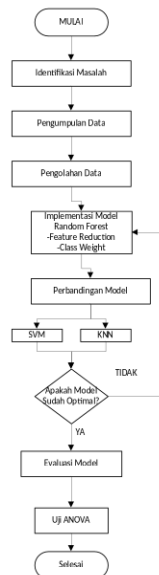
Tujuan penelitian ini adalah mengembangkan model prediksi risiko ketergantungan smartphone menggunakan algoritma random forest. Pengembangan penelitian ini terletak pada penerapan Random Forest dengan mengombinasikan variabel penggunaan smartphone dan pengawasan orang tua sebagai dasar prediksi, sehingga diharapkan dapat mendukung deteksi dini dan upaya pencegahan ketergantungan smartphone pada anak usia dini.

### **2. Metode Penelitian**

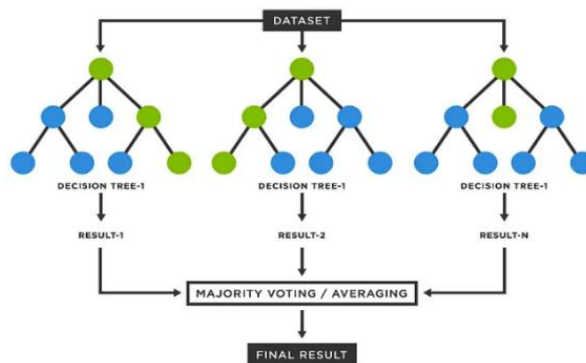
Penelitian ini menerapkan pendekatan kuantitatif dengan algoritma Random Forest untuk mengklasifikasikan Risiko ketergantungan smartphone pada anak usia dini (Fahria et al., 2026). Data yang digunakan dalam penelitian

ini merupakan data primer yang dikumpulkan melalui penyebaran kuesioner kepada orang tua yang memiliki anak usia dini.

Penelitian ini melalui beberapa tahapan, yaitu pengumpulan data, preprocessing dan pengolahan data, serta penerapan algoritma Random Forest sebagai model utama dalam proses prediksi. Selanjutnya, pada tahap evaluasi, performa model diukur menggunakan metrik accuracy, precision, recall, F1-score, serta confusion matrix untuk menilai kinerja model secara menyeluruh. Lalu, hasil prediksi dianalisis untuk mengidentifikasi variabel atau fitur yang paling berpengaruh terhadap risiko ketergantungan smartphone pada anak usia dini.



Gambar 1. Alur Penelitian



Gambar 2. Random Forest

Sumber: <https://medium.com>

Pemodelan Random Forest dimulai dengan penggunaan dataset sebagai data masukan. Dataset kemudian dibagi kedalam beberapa subset melalui penerapan Teknik bootstrap sampling untuk membentuk sejumlah pohon keputusan (*decision tree*) (Studi & Informatika, 2025). Lalu setiap pohon dilatih secara independen dan menghasilkan prediksi terhadap data uji. Seluruh hasil prediksi digabungkan menggunakan metode *majority voting*, yaitu kelas yang menghasilkan suara terbanyak dari seluruh pohon akan dipilih sebagai hasil klasifikasi. Proses tersebut menghasilkan prediksi akhir yang lebih stabil, akurat, dan mampu mengurangi risiko *overfitting* disbanding menggunakan satu pohon keputusan saja (Binarto & Wibowo, 2025).

### 3. Hasil dan Pembahasan

#### 3.1. Data Penelitian

Penelitian ini menggunakan data primer yang diperoleh melalui penyebaran kuesioner kepada orang tua anak usia dini. Penelitian ini melibatkan 347 responden dan data yang dikumpulkan terdiri atas satu variabel dependen, yaitu Tingkat Risiko ketergantungan smartphone, serta enam variabel independent, yaitu umur, durasi penggunaan smartphone, frekuensi penggunaan smartphone, usia pertama kali mengenal smartphone, tujuan penggunaan smartphone, dan pola pengawasan orang tua. Berdasarkan hasil pengumpulan data, mayoritas anak menggunakan smartphone dengan durasi lebih dari 2 jam per hari, serta jenis aplikasi yang dominan digunakan adalah aplikasi

hiburan seperti video dan game. Hal ini menunjukkan adanya potensi risiko ketergantungan yang cukup tinggi pada anak usia dini.

### 3.2. Preprocessing Data

Tahap prapemrosesan bertujuan meningkatkan kualitas data sebelum proses pemodelan. Tahapan yang dilakukan meliputi data cleaning, transformasi data, normalisasi data dan pembagian data.(Forest, 2023).

Pada tahap data cleaning, dilakukan pengecekan terhadap data yang tidak lengkap (*missing value*) dan duplikasi data. Hasil pengecekan menunjukkan bahwa seluruh data telah lengkap dan tidak ditemukan data duplikat, sehingga data dapat langsung digunakan untuk tahap berikutnya. Setelah itu dilakukan proses encoding, tahapan ini dilakukan untuk mengonversi data kategorikal menjadi nilai numerik sehingga dapat diproses oleh algoritma machine learning.

Tabel 1. Pelabelan Data

| Kelas  | Nilai |
|--------|-------|
| Rendah | 0     |
| Sedang | 1     |
| Tinggi | 2     |

### 3.3. Pembagian Data Training dan Data Testing

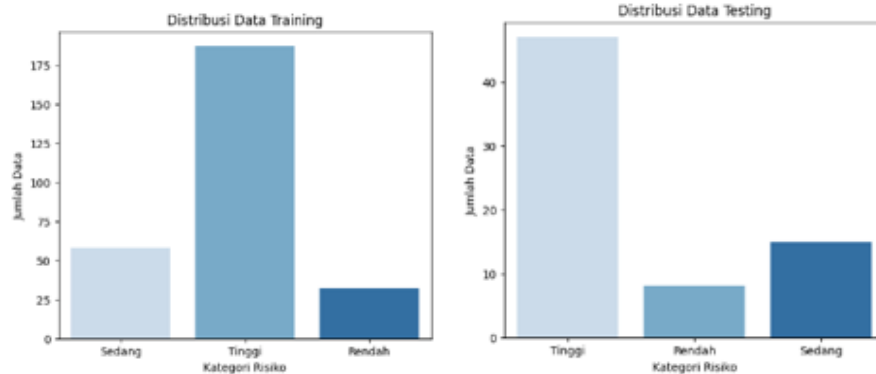
Sebelum pemodelan, data dibagi menjadi data training dan data testing menggunakan metode split data dengan rasio 80:20. Dari 347 data, sebanyak 277 data digunakan untuk pelatihan model, sedangkan 70 data digunakan untuk pengujian model.

Tabel 2. Distribusi Data Training

| Kelas        | Jumlah Data |
|--------------|-------------|
| Rendah       | 32          |
| Sedang       | 58          |
| Tinggi       | 187         |
| <b>Total</b> | <b>277</b>  |

Tabel 3. Distribusi Data Testing

| Kelas        | Jumlah Data |
|--------------|-------------|
| Rendah       | 8           |
| Sedang       | 15          |
| Tinggi       | 47          |
| <b>Total</b> | <b>70</b>   |

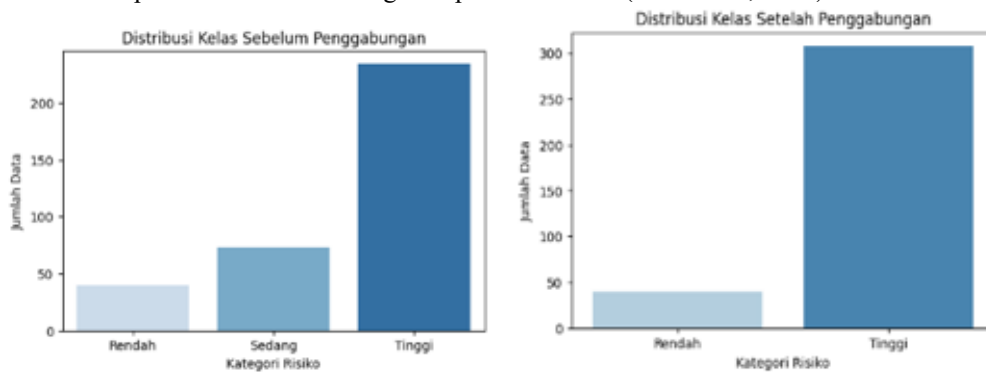


**Gambar 3.** Distribusi Data Training dan Data Testing

Berdasarkan hasil pembagian data tersebut, diperoleh sebanyak 277 data training dan 70 data testing. Pada data training, kelas risiko tinggi memiliki jumlah data paling banyak yaitu 187 data, sedangkan kelas risiko sedang sebanyak 58 data dan kelas risiko rendah sebanyak 32 data. Sedangkan, pada data testing terdapat 47 data untuk kelas risiko tinggi, 15 data untuk kelas risiko sedang, dan 8 data untuk kelas risiko rendah. Hasil distribusi data tersebut menunjukkan bahwa jumlah data pada setiap kelas masih tidak seimbang, di mana kelas risiko tinggi mendominasi dibandingkan kelas lainnya.

### 3.4. Model Random Forest

Sebelum proses pelatihan model diterapkan teknik penggabungan kelas (*class reduction*) dengan menggabungkan kelas Sedang ke dalam kelas Tinggi. Teknik ini digunakan untuk mengurangi ketidakseimbangan pada data serta meningkatkan kemampuan model dalam mengenali pola klasifikasi (Amri et al., 2026)

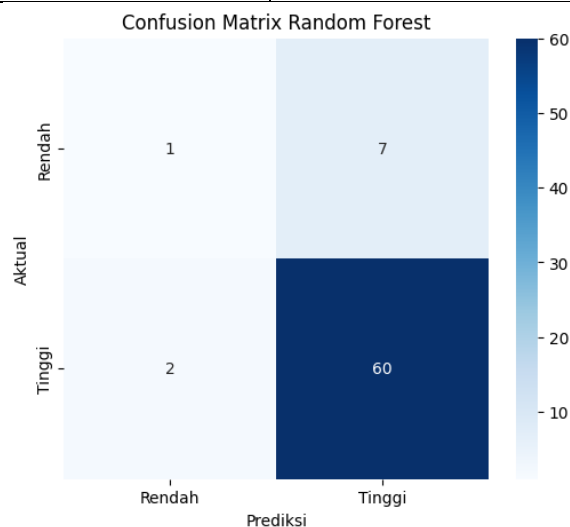


**Gambar 4.** Distribusi Kelas Sebelum Dan Setelah Penggabungan

Dari hasil penggabungan kelas tersebut dapat dilihat, jumlah data pada kategori Tinggi meningkat menjadi 307 data, sedangkan kategori Rendah tetap sebanyak 40 data.

**Tabel 4.** Parameter

| Parameter         | Nilai              |
|-------------------|--------------------|
| n_estimators      | 500                |
| max_depth         | 20                 |
| min_samples_split | 2                  |
| min_samples_leaf  | 1                  |
| max_features      | sqrt               |
| class_weight      | balanced_subsample |



**Gambar 5.** Confusion Matrix Random Forest

Pada confusion matrix tersebut dapat diketahui, sebanyak 1 data kategori Rendah dan 60 data kategori Tinggi berhasil diklasifikasikan dengan benar. Kemudian, terdapat 7 data kategori Rendah yang salah diprediksi sebagai Tinggi dan 2 data kategori Tinggi yang salah diprediksi sebagai Rendah.

**Tabel 5.** Classification Report Random Forest

| Kelas                   | Precision | Recall | F1-Score    | Support   |
|-------------------------|-----------|--------|-------------|-----------|
| Rendah                  | 0,33      | 0,12   | 0,18        | 8         |
| Tinggi                  | 0,90      | 0,97   | 0,93        | 62        |
| <b>Accuracy</b>         | -         | -      | <b>0,87</b> | <b>70</b> |
| <b>Macro Average</b>    | 0,61      | 0,55   | 0,56        | 70        |
| <b>Weighted Average</b> | 0,83      | 0,87   | 0,84        | 70        |

DOI: <https://doi.org/10.69693/ijmst.v4i2.12359>

Lisensi: Creative Commons Attribution 4.0 International (CC BY 4.0)

Hasil pengujian menggunakan data testing sebanyak 70 data, model Random Forest mencapai nilai akurasi 87,14%. Hasil tersebut menunjukkan bahwa model berhasil mengklasifikasikan 61 dari total 70 data uji dengan benar. Hasil evaluasi menunjukkan bahwa model *Random Forest* memiliki performa yang baik dalam memprediksi Tingkat Risiko ketergantungan smartphone pada anak usia dini. Tingginya nilai akurasi juga mengindikasikan bahwa penerapan teknik penggabungan kelas dan parameter yang digunakan mampu meningkatkan performa model dalam proses klasifikasi.

### 3.5. Model Support Vektor Machine (SVM)

Model Support Vector Machine digunakan sebagai algoritma pembanding dalam penelitian ini (Laili et al., 2025).

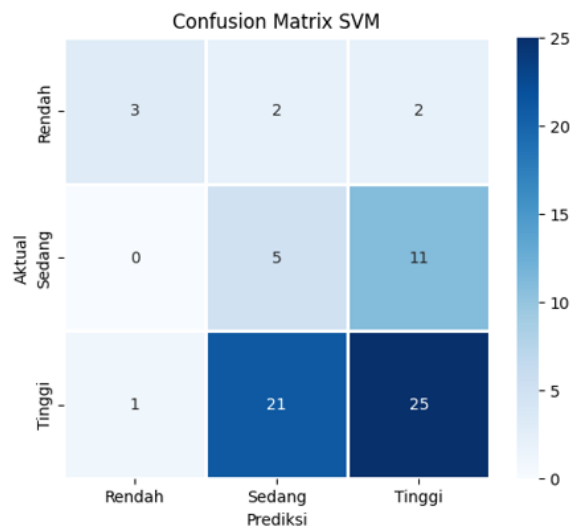
**Tabel 6.** Classification Report Model SVM

| Kelas                   | Precision | Recall | F1-Score    | Support |
|-------------------------|-----------|--------|-------------|---------|
| Rendah                  | 0,75      | 0,43   | 0,55        | 7       |
| Sedang                  | 0,18      | 0,31   | 0,23        | 16      |
| Tinggi                  | 0,66      | 0,53   | 0,59        | 47      |
| <b>Accuracy</b>         |           |        | <b>0,47</b> | 70      |
| <b>Macro Average</b>    | 0,53      | 0,42   | 0,45        | 70      |
| <b>Weighted Average</b> | 0,56      | 0,47   | 0,50        | 70      |

Pengujian menggunakan algoritma Support Vector Machine (SVM) menghasilkan akurasi sebesar 47%. Hasil tersebut menunjukkan bahwa performa model SVM dalam mengklasifikasikan Risiko ketergantungan smartphone pada anak usia dini masih relative rendah.

Pada kategori Rendah, model memperoleh nilai precision sebesar 0,75, recall sebesar 0,43, dan F1-score sebesar 0,55, yang menunjukkan bahwa model cukup baik dalam mengidentifikasi data pada kategori tersebut meskipun masih terjadi kesalahan klasifikasi. Pada kelas Sedang, nilai precision, recall, dan F1-score masing-masing sebesar 0,18, 0,31, dan 0,23, sehingga mengindikasikan bahwa model belum mampu mengenali pola pada kategori ini secara optimal. Sementara itu, pada kategori Tinggi, model memperoleh precision sebesar 0,66, recall sebesar 0,53, dan F1-score sebesar 0,59, yang menunjukkan kemampuan klasifikasi yang cukup baik, meskipun performanya masih perlu ditingkatkan.

Secara keseluruhan, hasil classification report menunjukkan bahwa algoritma SVM memiliki performa yang lebih rendah dibandingkan model Random Forest dalam mengklasifikasikan tingkat Risiko ketergantungan smartphone pada anak usia dini.



**Gambar 6.** Confusion Matrix SVM

Dari hasil confusion matrix pada algoritma SVM, diketahui bahwa model masih belum mampu melakukan klasifikasi dengan optimal. Pada kelas rendah, model berhasil mengklasifikasikan 3 data dengan benar, sedangkan 4 data lainnya mengalami kesalahan klasifikasi. Pada kelas Sedang, hanya 5 data yang benar diprediksi, sedangkan 11 data salah diklasifikasikan menjadi kelas Tinggi. Selain itu pada kelas Tinggi, model berhasil memprediksi 25 data dengan benar, tetapi masih terdapat 22 data yang salah klasifikasi.

### 3.6. K-Nearest Neighbor (KNN)

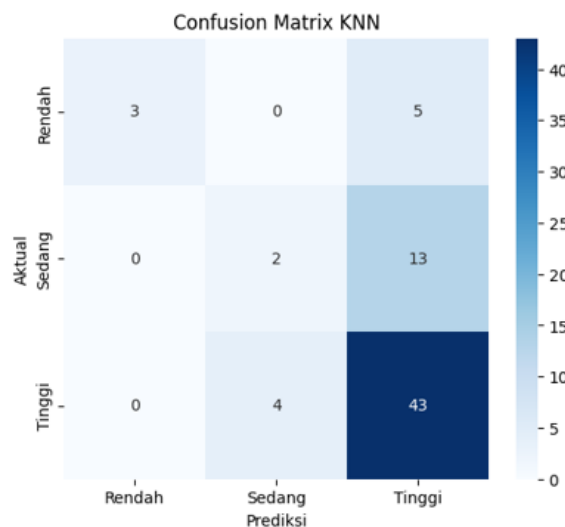
Berdasarkan hasil pengujian metode K-Nearest Neighbor (KNN), dilakukan percobaan menggunakan nilai K dari 1 hingga 5 untuk menentukan parameter terbaik pada model (Sabrina & Komarudin, 2024). Hasil pengujian

menunjukkan bahwa nilai akurasi tertinggi diperoleh pada  $K = 5$  dengan nilai akurasi sebesar 69%, sehingga nilai tersebut dipilih sebagai parameter optimal dalam proses klasifikasi.

**Tabel 7.** Classification Report Model KNN

| Kelas                   | Precision | Recall | F1-Score    | Support |
|-------------------------|-----------|--------|-------------|---------|
| Rendah                  | 1,00      | 0,38   | 0,55        | 8       |
| Sedang                  | 0,33      | 0,13   | 0,19        | 15      |
| Tinggi                  | 0,70      | 0,91   | 0,80        | 47      |
| <b>Accuracy</b>         |           |        | <b>0,69</b> | 70      |
| <b>Macro Average</b>    | 0,68      | 0,47   | 0,51        | 70      |
| <b>Weighted Average</b> | 0,66      | 0,69   | 0,64        | 70      |

Berdasarkan classification report, kelas Rendah memperoleh nilai precision sebesar 1,00, recall sebesar 0,38, dan F1-score sebesar 0,55. Hasil tersebut menunjukkan bahwa model masih memiliki keterbatasan dalam mengenali data pada kategori rendah. Pada kelas Sedang, dapat dilihat nilai precision sebesar 0,33, recall sebesar 0,13, dan f1-score sebesar 0,19. Nilai tersebut menunjukkan bahwa model belum cukup mampu mengidentifikasi data pada kelas Sedang dengan baik, karena sebagian besar data masih salah diklasifikasikan ke kelas lain. Pada kelas tinggi, model memperoleh nilai precision sebesar 0,70, recall sebesar 0,91, dan f1-score sebesar data yang termasuk dalam kategori tinggi secara tepat.



**Gambar 7.** Confusion Matrix KNN

Hasil confusion matrix KNN menunjukkan bahwa pada kelas rendah terdapat 8 data actual, dengan 3 data berhasil diprediksi secara tepat, sementara 5 data lainnya mengalami kesalahan klasifikasi ke kelas tinggi. Hal ini mengindikasikan bahwa model masih cukup kesulitan membedakan sebagian data Rendah dengan Tinggi. Pada kelas Sedang, dari total 15 data, hanya 2 data yang berhasil diprediksi dengan benar, sementara 13 data lainnya salah diklasifikasikan menjadi kelas Tinggi. Hal ini menunjukkan bahwa kelas Sedang masih belum dapat dikenali dengan baik oleh model KNN. Hasil confusion matrix menunjukkan bahwa pada kelas tinggi terdapat 47 data actual, dengan 43 data berhasil diprediksi secara tepat, sementara 4 data lainnya mengalami kesalahan klasifikasi ke kelas sedang. Temuan ini mengindikasikan bahwa model KNN lebih efektif dalam mengklasifikasikan data pada kelas Tinggi dibandingkan pada kelas lainnya.

Jika ditinjau secara keseluruhan, hasil confusion matrix menunjukkan bahwa model KNN cenderung lebih dominan dalam memprediksi kelas Tinggi, akan tetapi masih memiliki kelemahan dalam membedakan kelas Rendah dan Sedang.

### 3.7. Perbandingan Performa Model

Berdasarkan hasil pengujian, diperoleh perbandingan performa antar model sebagai berikut:

**Tabel 8.** Perbandingan Performa Model

| Model         | Accuracy | Precision (Macro Avg) | Recall (Macro Avg) | F1-Score (Macro Avg) |
|---------------|----------|-----------------------|--------------------|----------------------|
| Random Forest | 0.87     | 0.61                  | 0.55               | 0.56                 |
| SVM           | 0.47     | 0.53                  | 0.42               | 0.45                 |

|     |      |      |      |      |
|-----|------|------|------|------|
| KNN | 0.69 | 0.68 | 0.47 | 0.51 |
|-----|------|------|------|------|

Dari hasil perbandingan tiga model klasifikasi yaitu Random Forest, SVM, dan KNN, terdapat perbedaan performa yang cukup jelas pada setiap metrik evaluasi. Model Random Forest menghasilkan performa terbaik dengan akurasi sebesar 0,87, serta nilai precision 0,65, recall 0,60, dan F1-score 0,62. Hal tersebut menandakan bahwa model ini mampu mengklasifikasikan data dengan cukup baik dan seimbang dibandingkan model lainnya. Model KNN berada di posisi kedua dengan akurasi 0,69. Meskipun precision cukup tinggi yaitu 0,68, nilai recall yang lebih rendah (0,47) menunjukkan masih adanya data yang belum terdeteksi dengan baik, sehingga F1-score hanya mencapai 0,51. Sedangkan model SVM memiliki performa terendah dengan akurasi 0,47, serta nilai precision 0,53, recall 0,42, dan F1-score 0,45, yang menunjukkan bahwa model ini kurang optimal dalam mengenali pola data.

Secara keseluruhan, Random Forest menjadi model paling unggul untuk dataset ini karena memiliki keseimbangan performa terbaik di antara ketiga model.

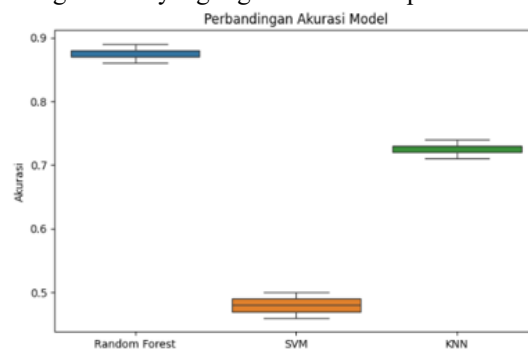
### 3.8. Uji Anova

Uji ANOVA digunakan untuk menganalisis apakah terdapat perbedaan yang signifikan dalam performa model *Random Forest*, *SVM*, *KNN*.

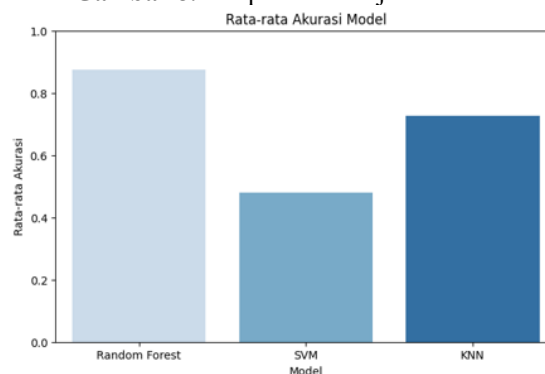
**Tabel 9.** Hasil Uji Anova Perbandingan Model

| Metode                  | F-value | P-value | Keterangan                    |
|-------------------------|---------|---------|-------------------------------|
| Random Forest, KNN, SVM | 1164.98 | < 0.001 | Terdapat Perbedaan Signifikan |

Hasil uji ANOVA pada model Random Forest, K-Nearest Neighbor (KNN), dan Support Vector Machie (SVM) menghasilkan f-value sebesar 1164,98 dan p-value <0,001. Nilai p-value yang lebih kecil dari Tingkat signifikansi ( $\alpha = 0,05$ ) menunjukkan bahwa hipotesis nol ( $H_0$ ) ditolak, sehingga terdapat perbedaan performa yang signifikan diantara ketiga model. Berdasarkan hasil uji ANOVA, dapat diketahui bahwa model *Random Forest*, *KNN*, dan *SVM* menunjukkan perbedaan performa yang signifikan dalam mengklasifikasikan data. Nilai F-value yang sangat besar juga menunjukkan bahwa perbedaan performa antar model cukup tinggi dan tidak terjadi secara kebetulan. Berdasarkan hasil pengujian model sebelumnya, Random Forest memiliki performa terbaik dibandingkan KNN dan SVM. Selain itu, Random Forest juga lebih stabil dalam mengklasifikasikan seluruh kelas data karena memiliki nilai evaluasi yang lebih seimbang. Dengan demikian, hasil uji ANOVA memperkuat bahwa terdapat perbedaan nyata pada performa ketiga model yang digunakan dalam penelitian ini.



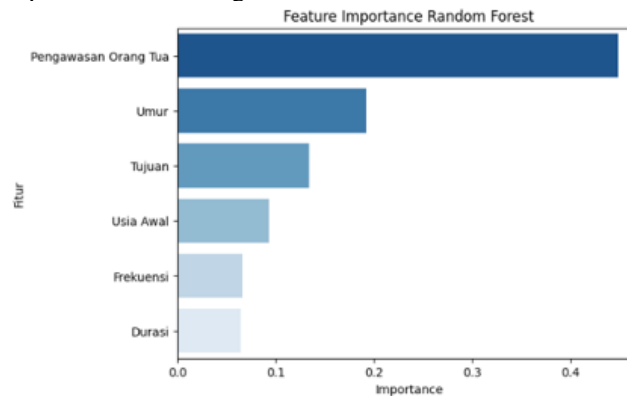
**Gambar 8.** Boxplot Hasil Uji ANOVA



**Gambar 9.** Rata-Rata Akurasi Ketiga Model

### 3.9. Analisis Feature Importance

Algoritma Random Forest memiliki kemampuan untuk mengukur Tingkat kepentingan setiap variabel melalui analisis feature importance. Hasil analisis menunjukkan bahwa variabel yang paling berpengaruh terhadap prediksi Risiko ketergantungan smartphone Adalah sebagai berikut:



**Gambar 10.** Diagram Feature Importance

Berdasarkan hasil analisis *feature importance* menggunakan algoritma *Random Forest*, diketahui bahwa masing-masing variabel memberikan kontribusi yang berbeda terhadap proses klasifikasi Tingkat Risiko ketergantungan smartphone pada anak usia dini. Variabel yang memiliki pengaruh paling besar adalah Pengawasan Orang Tua dengan nilai importance sekitar 0,45. Hasil analisis menunjukkan bahwa Tingkat pengawasan orang tua merupakan factor utama dalam memprediksi Risiko ketergantungan smartphone pada anak.

Variabel Umur menempati urutan kedua dengan nilai importance sekitar 0,19. Hal ini menunjukkan bahwa perbedaan usia anak turut berpengaruh terhadap tingkat risiko ketergantungan smartphone. Sementara itu, variabel Tujuan Penggunaan Smartphone memiliki nilai importance sekitar 0,14, yang mengindikasikan bahwa aktivitas yang dilakukan anak saat menggunakan smartphone juga berkontribusi pada hasil klasifikasi. Variabel Usia Awal Mengenal Smartphone memiliki nilai importance sekitar 0,09, sedangkan Frekuensi Penggunaan Smartphone dan Durasi Penggunaan Smartphone memiliki nilai importance yang lebih rendah, masing-masing sekitar 0,07 dan 0,06. Meskipun kedua variabel tersebut sering dianggap sebagai faktor utama dalam penggunaan smartphone, hasil penelitian menunjukkan bahwa pengaruhnya masih berada di bawah faktor pengawasan orang tua dan karakteristik anak. Untuk keseluruhan, hasil *feature importance* menunjukkan bahwa faktor yang berkaitan dengan peran dan pengawasan orang tua memiliki kontribusi paling dominan dalam menentukan tingkat risiko ketergantungan smartphone pada anak usia dini. Dengan demikian, peningkatan pengawasan dan pendampingan orang tua dapat menjadi salah satu upaya penting dalam mengurangi risiko ketergantungan smartphone pada anak.

## 4. Kesimpulan

Penelitian ini menghasilkan model prediksi risiko ketergantungan smartphone pada anak usia dini menggunakan Random Forest berdasarkan data 347 responden. Penerapan class reduction dan class weight membantu mengatasi ketidakseimbangan data. Random Forest mencapai akurasi 87%, lebih baik daripada Support Vector Machine dan K-Nearest Neighbor, sehingga hipotesis diterima. Model berpotensi mendukung deteksi dini bagi orang tua dan pihak terkait dalam mencegah ketergantungan smartphone. Keterbatasan penelitian mencakup jumlah responden dan variabel yang digunakan. Penelitian selanjutnya perlu memperluas wilayah, menambah responden dan variabel relevan, membandingkan algoritma lain, serta mengembangkan model menjadi aplikasi berbasis web atau perangkat seluler agar kemampuan generalisasi dan pemanfaatannya semakin baik.

## Reference

- ADDIN Mendeley Bibliography CSL\_BIBLIOGRAPHY Amri, S., Achhab, M. A. L., & Lazaar, M. (2026). *Comparative Analysis Of Machine Learning Based Algorithms For Predicting Injury Severity In Road Accidents*. 17(2), 512–523.
- Angkasa, C., Dharshinni, N. P., & Willgert, S. C. (2026). *Pengaruh Screen Time Berdasarkan Tujuan Penggunaan Terhadap Adiksi Gadget Menggunakan Support Vector Machine*. 10(1), 1353–1360.
- Binarto, A. J., & Wibowo, A. (2025). *Implementasi Majority Voting Pada Framework Cross-Industry Standard Process For Data Mining Untuk Prediksi Kepatuhan Wajib Pajak*. <https://doi.org/10.33364/Algoritma/V.22-1.2165>
- Fahria, D. N., Tania, K. D., & Kurnia, R. D. (2026). *Analisis Komparatif Algoritma Random Forest, Xgboost, Dan Catboost Untuk Klasifikasi Tingkat Stres Pengguna Media Sosial I*. 11(1), 1843–1853.
- Forest, R. (2023). *Klasifikasi Data Mining Dengan Algoritma Machine Larning Untuk Prediksi Penyakit Liver*. 14(2).
- Hr, J., No, B., & Utara, P. (2026). *Komparasi Naïve Bayes Dan Random Forest Untuk Prediksi Tingkat Stres Berdasarkan Pola Penggunaan Medsos*. 6(1), 10–15.
- Khairani, F., Naria, E., Lubis, I. K., Koka, E. M., Harahap, A. F., Ranguti, I. M., Enjelika, M. T., & Daulay, H. (2024). *Hubungan Peran Orang Tua Dengan Intensitas Penggunaan Gadget Pada Anak Usia Dini Di Kabupaten Tapanuli Selatan The Relationship Between*

DOI: <https://doi.org/10.69693/ijmst.v4i2.12359>

Lisensi: Creative Commons Attribution 4.0 International (CC BY 4.0)

- 
- Parental Roles With The Intensity Of Gadget Use In Early Childhood In South Tapanuli Regency*. 04(01), 52–58.
- Laili, E. F., Alawi, Z., Rohmah, R., Barata, M. A., Informatika, T., Nahdlatul, U., Sunan, U., Informasi, S., Nahdlatul, U., Sunan, U., Komputer, S., Nahdlatul, U., Sunan, U., & Bojonegoro, K. (2025). *Komparasi Algoritma Decision Tree Dan Support Vector Machine ( Svm ) Dalam*. 8(1), 67–76.
- Sabrina, P. N., & Komarudin, A. (2024). *Prediksi Penyakit Diabetes Dengan Metode K-Nearest Neighbor ( Knn ) Dan Seleksi Fitur Information Gain*. 8(6), 11320–11326.
- Studi, P., & Informatika, T. (2025). *Implementasi Algoritma Random Forest Berbasis Machine Learning Untuk Prediksi Klon Kopi Unggul*. 9(2), 230–239.
- Aslan, A., & Turgut, Y. E. (2024). Parental Mediation In Turkey: The Use Of Mobile Devices In Early Childhood. *E-Learning And Digital Media*, 21(5), 444–461. <https://doi.org/10.1177/20427530231167651>
- Fitzpatrick, C., Cristini, E., Bernard, J. Y., & Garon-Carrier, G. (2023). Meeting Preschool Screen Time Recommendations: Which Parental Strategies Matter? *Frontiers In Psychology*, 14, 1287396. <https://doi.org/10.3389/fpsyg.2023.1287396>
- Ko, Y., & Park, S. (2023). Impacts Of Problematic Smartphone Use On Children: Perspectives From Main Caregivers. *Archives Of Psychiatric Nursing*, 46, 59–64. <https://doi.org/10.1016/j.apnu.2023.08.007>
- Lee, J., & Kim, W. (2021). Prediction Of Problematic Smartphone Use: A Machine Learning Approach. *International Journal Of Environmental Research And Public Health*, 18(12), 6458. <https://doi.org/10.3390/ijerph18126458>
- Long, C., Liu, J., Wu, Y., & Liu, S. (2024). The Relationship Between Parental Smartphone Dependence And Elementary Students' Internet Addiction During The COVID-19 Lockdown In China: The Mediating Role Of Parent–Child Conflict And The Moderating Role Of Parental Roles. *Frontiers In Psychology*, 15, 1480151. <https://doi.org/10.3389/fpsyg.2024.1480151>
- Santos, R. M. S., Mendes, C. G., Miranda, D. M., & Romano-Silva, M. A. (2022). The Association Between Screen Time And Attention In Children: A Systematic Review. *Developmental Neuropsychology*, 47(4), 175–192. <https://doi.org/10.1080/87565641.2022.2064863>
- Susilowati, I. H., Nugraha, S., Alimoeso, S., & Hasiholan, B. P. (2021). Screen Time For Preschool Children: Learning From Home During The COVID-19 Pandemic. *Global Pediatric Health*, 8, 2333794X211017836. <https://doi.org/10.1177/2333794X211017836>
- Veldman, S. L. C., Altenburg, T. M., Chinapaw, M. J. M., & Gubbels, J. S. (2023). Correlates Of Screen Time In The Early Years (0–5 Years): A Systematic Review. *Preventive Medicine Reports*, 33, 102214. <https://doi.org/10.1016/j.pmedr.2023.102214>
- Wang, J., Liu, R.-D., Ding, Y., Hong, W., & Liu, J. (2024). How Parental Mediation And Parental Phubbing Affect Preschool Children's Screen Media Use: A Response Surface Analysis. *Cyberpsychology, Behavior, And Social Networking*, 27(9), 651–657. <https://doi.org/10.1089/Cyber.2023.0638>
- Yang, X., Jiang, P., & Zhu, L. (2023). Parental Problematic Smartphone Use And Children's Executive Function: The Mediating Role Of Technoference And The Moderating Role Of Children's Age. *Early Childhood Research Quarterly*, 63, 219–227. <https://doi.org/10.1016/j.ecresq.2022.12.017>