

Deteksi Aplikasi Digital Berisiko Melanggar Regulasi Privasi Anak Menggunakan CatBoost

DEVINA SAWITRI, YULIANI PUJI ASTUTI, YUNI ROSITA DEWI*

Program Studi Sarjana Sains Data, Universitas Negeri Surabaya, Indonesia

Abstrak

Perkembangan teknologi meningkatkan eksposur anak-anak terhadap aplikasi digital yang berpotensi melanggar regulasi privasi anak. Penelitian ini bertujuan membangun dan mengevaluasi model klasifikasi berbasis *CatBoost* untuk mendeteksi aplikasi digital yang berisiko melanggar regulasi privasi anak. Data yang digunakan bersumber dari kompetisi Kaggle Find IT 2025, terdiri dari 7.000 entri dengan 17 kolom atribut. Tahap persiapan data mencakup pembersihan nilai tidak konsisten, penghapusan nilai redundan, dan penanganan nilai tidak valid. Pembagian data dilakukan secara *stratified* dengan rasio 80:10:10 untuk data latih, data validasi, dan data uji. Mengingat karakteristik *dataset* yang tidak seimbang dengan rasio kelas 1:9, optimasi model dilakukan melalui kombinasi *hyperparameter optimization* menggunakan *Optuna*, implementasi *Mutually Disjoint Datasets*, dan optimasi *threshold*. Hasil evaluasi menunjukkan bahwa model yang dirancang mampu mencapai *Macro Recall* sebesar 0,84. *Recall* pada kelas minoritas mencapai 0,90, yang berarti 63 dari 70 aplikasi berisiko berhasil terdeteksi. *Recall* kelas mayoritas mencapai 0,78, yakni 491 dari 629 aplikasi aman berhasil terdeteksi dengan benar. Hasil tersebut lebih baik dibandingkan dengan model *CatBoost* tanpa optimasi yang hanya mampu mendeteksi 14 dari 139 aplikasi berisiko, dengan *recall* kelas minoritas sebesar 0,10. Strategi memaksimalkan *recall* diprioritaskan karena dalam hal perlindungan privasi anak, kesalahan melewatkan aplikasi berbahaya lebih merugikan dibandingkan kesalahan menandai aplikasi yang sebenarnya aman. Berdasarkan hasil tersebut, model yang dirancang menunjukkan peningkatan dalam kemampuan deteksi aplikasi berbahaya bagi privasi anak.

Kata Kunci *CatBoost; klasifikasi; dataset tidak seimbang; privasi anak; aplikasi digital*

Abstract

Technological advancements have increased children's exposure to digital applications that may violate children's privacy regulations. This study aims to develop and evaluate a CatBoost-based classification model to detect digital applications that pose a risk of violating children's privacy regulations. The dataset used was obtained from the Kaggle Find IT 2025 competition and consists of 7,000 records with 17 attributes. The data preparation stage included cleaning inconsistent values, removing redundant values, and handling invalid values. The dataset was divided using a stratified approach with a ratio of 80:10:10 for the training, validation, and testing sets, respectively. Considering the imbalance nature of the dataset, which has a class ratio of 1:9, the model was optimized using a combination of hyperparameter optimization with Optuna, the implementation of *Mutually Disjoint Datasets*, and threshold optimization. The evaluation results show that the proposed model achieved a Macro Recall of 0.84. The recall

*Corresponding author. Email: yulianipuji@unesa.ac.id or yunidewi@unesa.ac.id.

for the minority class reached 0.90, indicating that 63 out of 70 high-risk applications were successfully detected. The recall for the majority class reached 0.78, meaning that 491 out of 629 safe applications were correctly identified. These results outperform the baseline CatBoost model without optimization, which was only able to detect 14 out of 139 high-risk applications, corresponding to a minority-class recall of 0.10. Maximizing recall was prioritized because, in the context of children's privacy protection, failing to identify a harmful application is more detrimental than incorrectly classifying a safe application as risky. Based on these findings, the proposed model demonstrates improved capability in detecting applications that may threaten children's privacy.

Keywords *CatBoost; classification; imbalanced dataset; children's privacy; digital application*

1 Pendahuluan

Pesatnya perkembangan teknologi digital telah mengubah pola interaksi anak-anak dengan perangkat digital. Data Badan Pusat Statistik ([Badan Pusat Statistik, 2025](#)) mencatat bahwa 42,25% anak usia dini di Indonesia telah menggunakan telepon seluler dan 41,02% sudah mengakses internet, sementara laporan UNICEF ([Stoilova et al., 2021](#)) menyebutkan bahwa lebih dari sepertiga pengguna internet global adalah anak-anak dan remaja. Akses yang luas ini membuka potensi risiko privasi yang belum sepenuhnya dipahami, baik oleh anak maupun orang tua.

Berbagai negara telah mengadopsi aturan hukum untuk melindungi privasi anak secara lebih mendalam. Di Amerika Serikat, perlindungan tersebut diatur melalui Children's Online Privacy Protection Act (COPPA), sedangkan di Indonesia diterapkan melalui Undang-Undang Perlindungan Data Pribadi (UU PDP). Kedua regulasi ini memiliki kesamaan tujuan, yaitu memberikan perlindungan hukum terhadap data pribadi, termasuk bagi anak ([Putri et al., 2024](#); [Afina and Prastyanti, 2025](#)).

Berdasarkan ketentuan COPPA, suatu aplikasi dapat dianggap berpotensi melanggar regulasi apabila tidak menyediakan kebijakan privasi yang jelas dan mudah diakses, mengumpulkan data pribadi anak tanpa memperoleh persetujuan orang tua yang dapat diverifikasi, membagikan data kepada pihak ketiga tanpa izin yang sesuai, tidak menyediakan mekanisme bagi orang tua untuk mengakses maupun menghapus data anak, tidak menerapkan langkah-langkah yang memadai untuk menjaga keamanan data, atau mengumpulkan informasi pribadi yang tidak diperlukan untuk tujuan operasional layanan ([Federal Trade Commission, 2025](#)).

Undang-Undang Pelindungan Data Pribadi (UU PDP) telah mengklasifikasikan data anak sebagai data pribadi yang memerlukan perlindungan khusus, namun implementasinya pada platform digital tetap menghadapi berbagai kendala ([Andiani and Wiraguna, 2025](#)). Keterbatasan dalam pengawasan manual menjadi salah satu kendala dalam identifikasi pelanggaran privasi anak, terlebih dengan jumlah aplikasi digital yang terus bertambah. Dengan demikian, diperlukan solusi yang efektif dan efisien untuk mengidentifikasi pelanggaran privasi pada aplikasi digital secara lebih sistematis.

Machine learning memberikan solusi terhadap permasalahan tersebut. [Caushaj and Sugumarman \(2023\)](#) membuktikan efektivitas *machine learning* dalam klasifikasi risiko privasi aplikasi, dengan model Decision Tree. [Zakariya and Ramli \(2023\)](#) turut merancang model penilaian risiko privasi aplikasi dengan fitur izin dan atribut aplikasi menggunakan *ensemble learning* berbasis *Decision Tree*, *KNN*, dan *Random Forest*. Senada dengan [Hakami \(2025\)](#), mengklasifikasikan aplikasi berdasarkan tingkat risiko remaja pada fitur teks ulasan pengguna menggunakan *Attentional CNN* berbasis *BERT embeddings*. Namun ketiga penelitian tersebut belum menyoroti

aspek perlindungan data anak secara spesifik, celah yang menjadi titik tolak penelitian ini.

Penelitian menggunakan *CatBoost*, yakni algoritma berbasis *gradient boosting decision tree* yang dioptimalkan untuk data kategorikal melalui *ordered target statistics* dan *ordered boosting* (Prokhorenkova et al., 2018). *CatBoost* relevan untuk mendeteksi pelanggaran privasi anak pada aplikasi digital karena sebagian besar fitur dari data yg digunakan bersifat kategorikal. Algoritma *CatBoost* juga telah menunjukkan performa yang baik pada berbagai tugas klasifikasi (Nayak et al., 2022).

Berdasarkan uraian tersebut, penelitian berfokus pada deteksi aplikasi digital berisiko melanggar regulasi privasi anak menggunakan *CatBoost*. Data yang digunakan berupa data ribuan aplikasi yang mengacu pada indikator kepatuhan regulasi *Children's Online Privacy Protection Act* (COPPA). Kontribusi penelitian ini meliputi pembangunan model klasifikasi berbasis *CatBoost* untuk mendeteksi aplikasi digital yang berisiko melanggar regulasi privasi anak. Penelitian dapat berkontribusi pada perlindungan data pribadi anak, membantu regulator mengawasi kepatuhan aplikasi, serta meningkatkan kesadaran pengembang dan masyarakat terhadap keamanan data anak.

2 Metode

2.1 Alur Penelitian

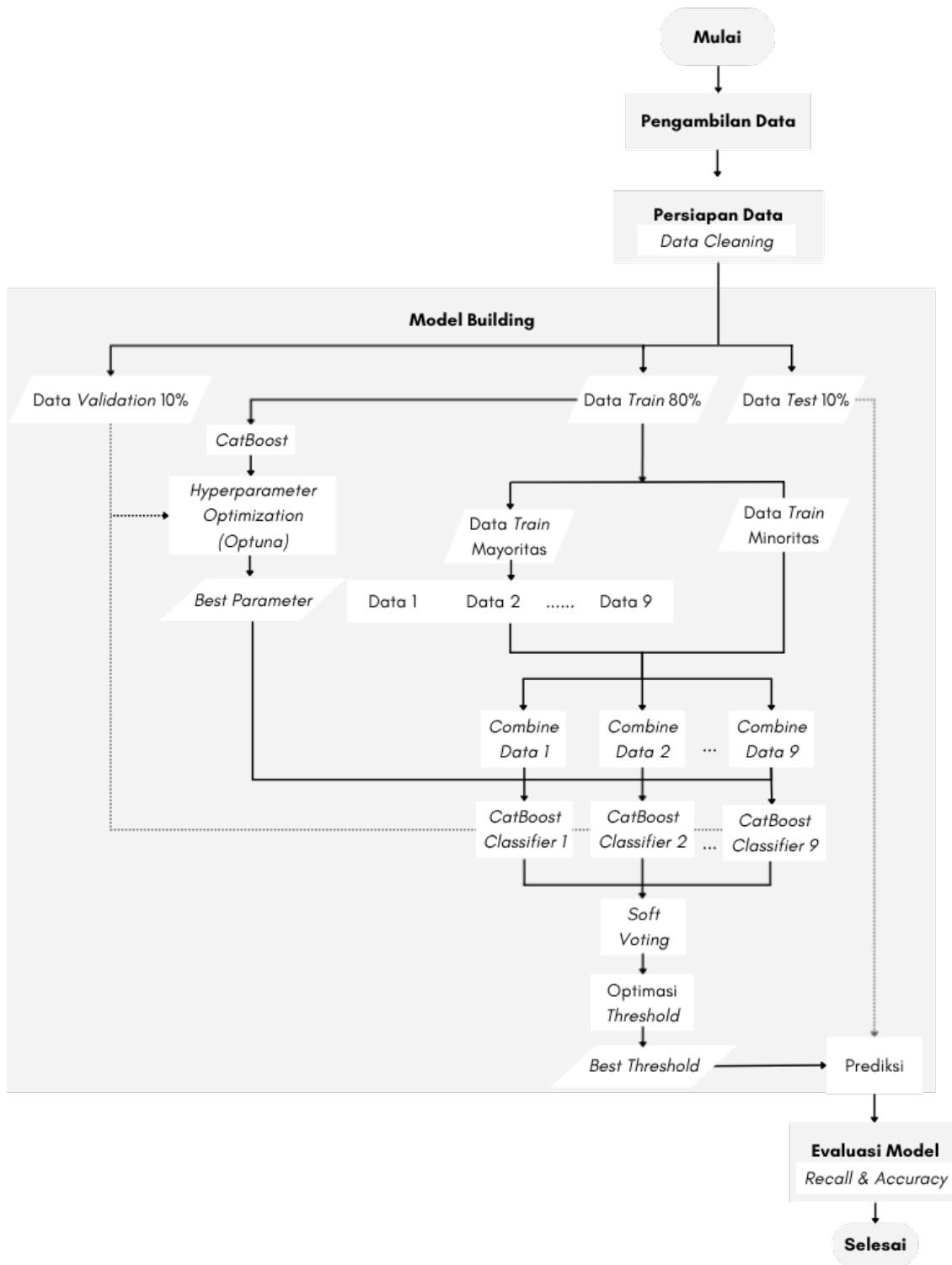
Penelitian ini terdiri atas lima tahap utama, yakni pengambilan data, persiapan data, pembagian *dataset*, pembentukan model, dan evaluasi model, sebagaimana ditunjukkan pada Gambar 1. Dimulai dari pengambilan data dan dilanjutkan dengan persiapan data melalui proses *data cleaning*. Setelah data dipersiapkan, *dataset* dibagi secara stratifikasi menjadi data latih sebesar 80%, data validasi sebesar 10%, dan data uji sebesar 10%. Data validasi digunakan pada tahap optimasi *hyperparameter* dan penentuan *threshold*, sedangkan data uji disimpan untuk evaluasi akhir model.

Pada tahap *model building*, data latih dipisahkan berdasarkan kelas mayoritas dan minoritas untuk membentuk sembilan *dataset* pelatihan dengan pendekatan *Mutually Disjoint Datasets*. Data dari kelas mayoritas dibagi menjadi sembilan bagian, kemudian masing-masing bagian digabungkan dengan seluruh data dari kelas minoritas sehingga diperoleh sembilan *dataset* dengan distribusi kelas yang lebih seimbang. Selanjutnya, masing-masing *dataset* digunakan untuk melatih satu model *CatBoost Classifier*, sehingga terbentuk sembilan model klasifikasi.

Probabilitas prediksi dari kesembilan model kemudian digabungkan menggunakan metode *soft voting*. Hasil penggabungan tersebut selanjutnya digunakan untuk menentukan *threshold* optimal berdasarkan data validasi. *Threshold* terbaik digunakan untuk menghasilkan prediksi pada data uji, yang kemudian dievaluasi menggunakan metrik *recall* dan *accuracy*.

2.2 Pengambilan Data

Data diperoleh dari kompetisi pada platform Kaggle, berukuran 7.000 baris dengan 17 kolom termasuk kolom target, sebagaimana ditunjukkan pada Tabel 1. *Dataset* terdiri atas fitur kategorikal, biner, dan numerik dengan tingkat *missing value* yang bervariasi. Berdasarkan pemeriksaan *missing value*, beberapa fitur tidak memiliki *missing value*, sedangkan fitur lainnya memiliki proporsi *missing value* yang cukup tinggi. Proporsi tertinggi terdapat pada penilaian keamanan konten aplikasi sebesar 88,03%.



Gambar 1: Alur Penelitian

Data merupakan hasil *web scraping* dari salah satu platform distribusi aplikasi disertai fitur turunan terkait kepatuhan terhadap *Children's Online Privacy Protection Act* (COPPA). Variabel target bertipe *boolean*, dengan nilai *True* menandakan risiko tinggi pelanggaran COPPA dan *False* menandakan risiko rendah. Dari 16 fitur lainnya, 12 fitur bersifat kategorikal dan 5 fitur bersifat numerik.

Tabel 1: Karakteristik Fitur pada *Dataset* Mentah

Nama Kolom	Tipe Data	Missing (%)
Negara pengembang	Kategorikal	0
Kode wilayah pasar	Kategorikal	0,91
Jumlah ulasan pengguna	Numerik	0
Kategori aplikasi	Kategorikal	0
Rentang jumlah unduhan	Kategorikal	30,70
Jenis perangkat	Kategorikal	0
Tautan kebijakan privasi	Biner	10,71
Tautan syarat ketentuan	Biner	66,21
Kualitas tautan syarat ketentuan	Kategorikal	66,21
Skor email korporat	Numerik	16,11
Biaya iklan	Numerik	81,13
Usia aplikasi	Numerik	0,71
Rata-rata penilaian pengguna	Numerik	17,60
Penilaian keamanan konten aplikasi	Kategorikal	88,03
Penilaian keamanan deskripsi aplikasi	Kategorikal	0
Penilaian orientasi iklan	Kategorikal	0
Risiko pelanggaran COPPA	Biner	0

Hasil pemeriksaan kelengkapan data menunjukkan variasi tingkat *missing value* antarfitur. Mulai dari 0% pada beberapa fitur hingga mencapai 88,03% pada penilaian keamanan konten aplikasi. Pada penelitian ini, *missing value* tidak diimputasi karena proses imputasi berpotensi menghasilkan nilai yang tidak merepresentasikan kondisi sebenarnya sehingga dapat menimbulkan bias pada proses pemodelan. Analisis distribusi variabel target juga dilakukan pada data. Hasil menunjukkan bahwa *dataset* bersifat *imbalanced*, dengan rasio kelas *False* terhadap *True* sebesar 1:9 sebagaimana ditunjukkan pada Gambar 2. Kondisi ini dipertimbangkan pada tahap *model building* (Subbab 2.5).

2.3 Persiapan Data

Data cleaning dilakukan melalui tiga langkah. Pertama, mengatasi inkonsistensi penulisan. Kedua, mengatasi redundansi. Ketiga, mengatasi data tidak valid. Inkonsistensi penulisan ditemukan pada kolom negara pengembang, melalui pemeriksaan nilai yang tidak dapat diidentifikasi sebagai nama negara. Nilai tersebut kemudian diseragamkan menjadi “unknown”.

Redundansi ditemukan antara kolom tautan syarat ketentuan dengan kolom kualitas tautan syarat ketentuan melalui pemeriksaan pola nilai pada kedua kolom.. Kedua kolom memiliki pola nilai yang saling berpasangan secara konsisten, nilai *low* berpasangan dengan *True*, nilai *high* berpasangan dengan *False*, dan *missing value* pada kedua kolom muncul secara bersamaan. Sebagai solusi, kolom kualitas tautan syarat ketentuan dihapus.

Penghapusan baris duplikat juga dilakukan pada data berdasarkan kesamaan seluruh atribut. Adapun data tidak valid ditemukan pada kolom rentang jumlah unduhan, berupa nilai yang tidak merepresentasikan format minimum–maksimum yang seharusnya. Nilai tersebut ditangani dengan mengubahnya menjadi *missing value*.

2.4 CatBoost

CatBoost (*Categorical Boosting*) dipilih sebagai algoritma klasifikasi karena dirancang khusus untuk menangani fitur kategorikal melalui dua mekanisme berbasis *ordering*, *ordered target statistics* dan *ordered boosting* (Prokhorenkova et al., 2018; Yandex, 2025). *Ordered target statistics* mengubah fitur kategorikal menjadi representasi numerik dengan menghitung statistik target dari data sebelumnya dalam urutan acak, sehingga mencegah *target leakage*, sebagaimana ditunjukkan pada Persamaan (1).

$$\hat{x}_{ik} = \frac{\sum_{x_j \in D_k} \mathbf{1}(x_{ij} = x_{ik}) \cdot y_j + ap}{\sum_{x_j \in D_k} \mathbf{1}(x_{ij} = x_{ik}) + a} \quad (1)$$

Pada setiap data ke- i dan fitur kategorikal ke- k , model hanya menggunakan data-data sebelumnya ($x_j \in D_k$), dengan fungsi indikator $\mathbf{1}(x_{ij} = x_{ik})$ bernilai 1 jika kategori pada data j sama dengan kategori milik data i . Notasi p menyatakan nilai prior berupa rata-rata global target, sedangkan $a > 0$ adalah parameter penyeimbang prior; pembilang menjumlahkan target dari kategori yang sama, sedangkan penyebut menghitung banyaknya kemunculan kategori tersebut.

Ordered boosting memastikan perhitungan gradien pada suatu sampel hanya menggunakan informasi dari data yang muncul sebelumnya dalam permutasi acak σ , sebagaimana ditunjukkan pada Persamaan (2).

$$r_i = y_i - M_{\sigma(i)-1}(x_i), \quad i = 1, \dots, N \quad (2)$$

r_i adalah residual, y_i adalah target sebenarnya, $M_{\sigma(i)-1}(x_i)$ adalah prediksi model sebelumnya yang hanya menggunakan data sebelum contoh i dalam permutasi σ , dan N adalah jumlah total sampel data. Pohon keputusan baru kemudian dilatih untuk memprediksi residual tersebut, sehingga model secara bertahap memperbaiki kesalahan sebelumnya.

2.5 Model Building

Pembentukan model diawali pembagian data secara *stratified split* berdasarkan variabel target dengan rasio 80:10:10 sebagai pelatihan, validasi, dan pengujian. Data latih digunakan untuk membangun model, data validasi digunakan untuk *hyperparameter optimization* dan optimasi *threshold*, sedangkan data uji digunakan untuk evaluasi generalisasi model.

Hyperparameter optimization dilakukan setelah data dibagi untuk memaksimalkan nilai *recall* mengingat *dataset* bersifat *imbalanced*. Pada penelitian ini, *Optuna* digunakan karena mampu melakukan optimasi *hyperparameter* secara efisien melalui mekanisme pencarian adaptif dan penghentian dini (*early pruning*) pada kombinasi parameter yang kurang menjanjikan (Akiba et al., 2019; Lai et al., 2023; Tikaningsih et al., 2024). Parameter yang dioptimasi meliputi *learning rate*, kedalaman pohon, koefisien regularisasi L2, dan bobot kelas minoritas.

Parameter terbaik digunakan untuk melatih model dengan mengimplementasi metode *Mutually Disjoint Datasets* (Koyyada and Singh, 2023). Sebagaimana ditunjukkan pada Gambar 3, data latih kelas mayoritas dibagi menjadi sembilan subset sesuai rasio *class imbalance* (1:9),

masing-masing digabung dengan seluruh data kelas minoritas, sehingga menghasilkan sembilan model *CatBoost* untuk mengurangi bias terhadap kelas mayoritas.

Data validasi yang telah dipisahkan sebelumnya, kemudian diproses oleh sembilan model *CatBoost*. Keluaran yang diambil pada penelitian hanya probabilitas kelas positif dari setiap baris data. Probabilitas keluaran kesembilan model digabung melalui *soft voting* (rata-rata probabilitas). *Threshold* bawaan cenderung menghasilkan prediksi yang bias terhadap kelas mayoritas pada *dataset imbalanced* (Miftahushudur et al., 2025). Oleh karena itu, pencarian *threshold* optimal dilakukan berdasarkan metrik yang ingin dioptimalkan (Gad, 2025; Leevy et al., 2023). Pada penelitian ini, metrik yang digunakan sebagai dasar optimasi adalah *recall*, sehingga *threshold* terbaik dipilih berdasarkan nilai *recall* tertinggi dan digunakan pada tahap evaluasi.

2.6 Evaluasi Model

Evaluasi dilakukan pada data uji yang sepenuhnya terpisah dari data latih dan data validasi untuk mengukur kemampuan generalisasi model. Metrik yang digunakan berupa *recall* kelas *True* dan akurasi, mengikuti Koyyada and Singh (2023). *Recall* kelas *True* dipilih sebagai metrik utama untuk mengukur proporsi kasus berisiko (*True*) yang berhasil terdeteksi, mengingat karakteristik *dataset* yang *imbalanced*. *Recall* kelas *True* dan akurasi dihitung menggunakan Persamaan (3) dan (4) (Sokolova and Lapalme, 2009).

$$\text{Recall} = \frac{TP}{TP + FN} \quad (3)$$

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (4)$$

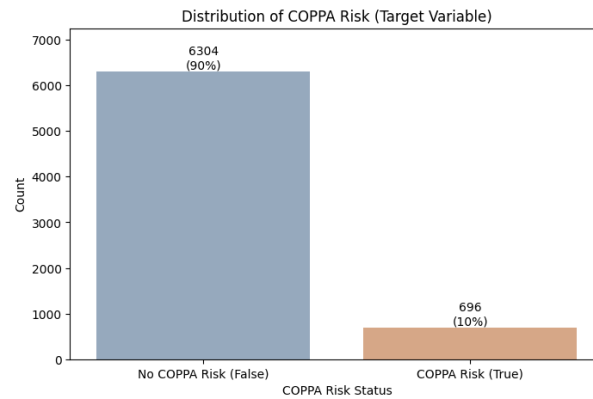
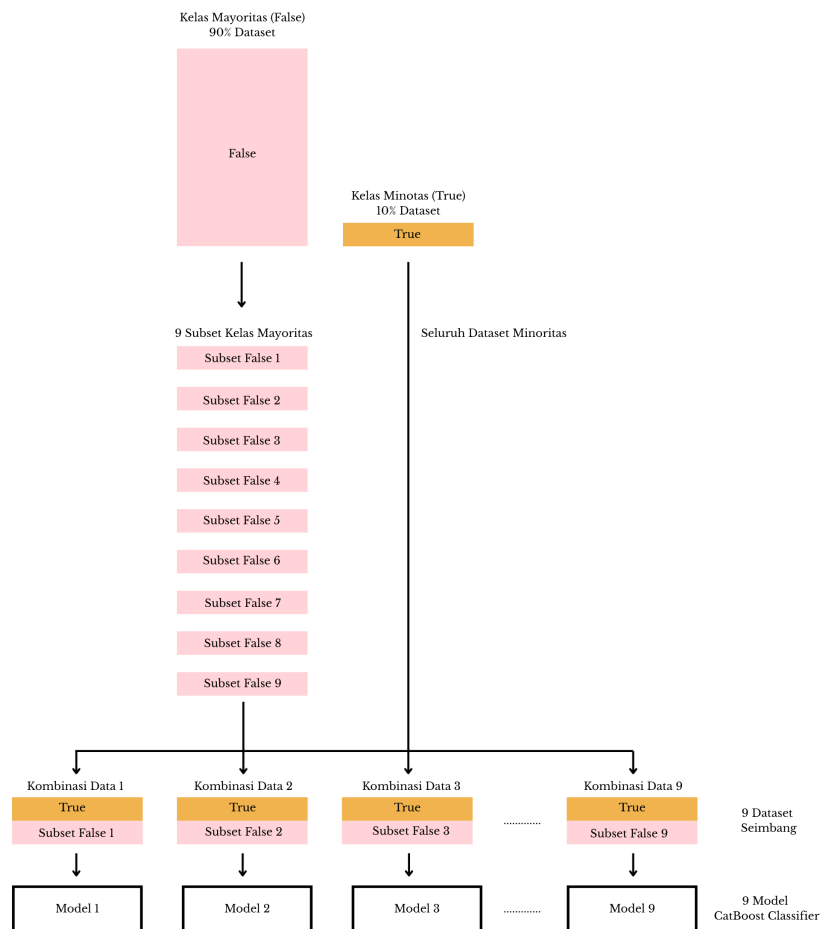
TP, *TN*, *FP*, dan *FN* masing-masing menyatakan jumlah sampel positif terprediksi benar, negatif terprediksi benar, negatif terprediksi positif, dan positif terprediksi negatif.

3 Hasil dan Pembahasan

Serangkaian perlakuan dilakukan sebelum pemodelan untuk memastikan kualitas data. Beberapa permasalahan ditemukan pada *dataset* mentah. Kolom negara pengembang mengandung nilai-nilai tidak konsisten yang semuanya bermakna identitas lokasi tidak tersedia, sehingga disragamkan menjadi “unknown”. Kolom kualitas tautan syarat ketentuan dihapus karena terbukti redundan secara sempurna dengan kolom tautan syarat ketentuan, keduanya membawa informasi identik dalam representasi berbeda. Kolom rentang jumlah unduhan mengandung beberapa nilai tidak valid tidak sesuai format, yang dikonversi menjadi *missing value* agar informasi baris lain tetap terjaga. Dari hasil cek duplikasi ditemukan 3 data dalam 3 baris yang sama atau disebut data duplikat. Sehingga sesuai proses prapengolahan data, data yang duplikat dihapus.

Setelah seluruh tahap persiapan, *dataset* bersih terdiri dari 6.981 entri dan 16 kolom diantaranya 10 kolom kategorikal, 5 kolom numerikal, dan kolom target. *Missing value* pada kolom kategorikal diimputasi dengan *string* kosong, sedangkan pada kolom numerikal dibiarkan karena *CatBoost* menanganinya secara otomatis melalui alokasi ke partisi minimum. Distribusi variabel target juga menunjukkan ketidakseimbangan kelas dengan rasio 1:9 antara kelas positif (696 sampel) dan kelas negatif (6.285 sampel).

Dominasi fitur kategorikal dalam *dataset* menjadi dasar pemilihan algoritma *CatBoost*, yang dirancang untuk menangani fitur kategorikal. Data dibagi menggunakan *stratified splitting* den-

Gambar 2: Distribusi Variabel Target pada *Dataset* MentahGambar 3: Implementasi Metode *Mutually Disjoint Datasets*

gan proporsi 80% data latih (5.584 sampel), 10% data validasi (698 sampel), dan 10% data uji (699 sampel).

Hyperparameter optimization dilakukan menggunakan *framework Optuna* selama 100 percobaan. Metrik yang dioptimalkan adalah *Macro Recall* guna memastikan model sensitif terhadap kelas minoritas. Hasil terbaik diperoleh pada percobaan ke-23 dengan *Macro Recall* sebesar 0,85 pada data validasi. Konfigurasi *hyperparameter* optimal disajikan pada Tabel 2.

Tabel 2: *Hyperparameter* Optimal Hasil *Optuna*

Parameter	Nilai Optimal
<i>Learning rate</i>	0.04
Kedalaman pohon (<i>depth</i>)	9
Koefisien regularisasi L2 (<i>L2 leaf regularization</i>)	5.75
Bobot kelas minoritas (<i>minority weight</i>)	4.2

Penanganan ketidakseimbangan kelas dilakukan menggunakan metode *Mutually Disjoint Datasets* (Koyyada and Singh, 2023). Data kelas mayoritas dibagi menjadi sembilan *subset* yang masing-masing berukuran sebanding dengan keseluruhan data kelas minoritas. Setiap *subset* digabungkan dengan seluruh data minoritas, menghasilkan sembilan *dataset* pelatihan yang relatif seimbang. Sembilan model *CatBoost* dilatih secara independen menggunakan konfigurasi *hyperparameter* optimal, sehingga seluruh informasi kelas mayoritas dimanfaatkan secara merata tanpa pembuangan data.

Nilai *threshold* default 0,5 berisiko menghasilkan bias terhadap kelas mayoritas pada data tidak seimbang. Oleh karena itu, *threshold* optimal dicari pada rentang 0,10–0,90 menggunakan probabilitas *soft voting* dari kesembilan model terhadap data validasi. *Threshold* 0,81 dipilih karena menghasilkan *Macro Recall* tertinggi sebesar 0,84, mencerminkan strategi model yang lebih selektif dalam menghasilkan prediksi positif guna meningkatkan sensitivitas terhadap kelas minoritas.

Evaluasi akhir dilakukan pada data uji yang sepenuhnya terpisah dari proses pelatihan dan validasi. Hasil prediksi diperoleh melalui *soft voting* dari kesembilan model, kemudian diklasifikasikan menggunakan *threshold* 0,81. Model mencapai akurasi 79% dengan *Macro Recall* sebesar 0,84. Hasil evaluasi per kelas disajikan pada Tabel 3.

Tabel 3: Hasil Evaluasi Model pada Data Uji

Kelas	<i>Recall</i>	Jumlah Sampel
<i>False</i> (Tidak Berisiko)	0.90	629
<i>True</i> (Berisiko Melanggar)	0.78	70
<i>Macro Recall</i>	0.84	699

Nilai *Recall* kelas minoritas sebesar 0,90 yang berarti 63 dari 70 aplikasi berisiko berhasil terdeteksi, hanya 7 yang tidak terdeteksi. Kondisi ini sesuai dengan prioritas penelitian dalam konteks perlindungan privasi anak. Kesalahan tidak mendeteksi aplikasi berbahaya memiliki dampak lebih serius dibandingkan kesalahan menandai aplikasi aman, karena kesalahan jenis kedua masih dapat dikoreksi melalui verifikasi manual.

Kontribusi optimasi diukur melalui perbandingan dengan model *baseline CatBoost* tanpa *hyperparameter optimization*, tanpa metode *Mutually Disjoint Datasets*, dan tanpa optimasi *threshold*. Hasil perbandingan disajikan pada Tabel 4.

Tabel 4: Perbandingan Model dengan dan Tanpa Optimasi

Metrik	Tanpa Optimasi	Dengan Optimasi
<i>Recall</i> kelas mayoritas	0.98	0.78
<i>Recall</i> kelas minoritas	0.10	0.90
<i>Macro Recall</i>	0.54	0.84
Akurasi	90%	79%

Model tanpa optimasi mencapai akurasi 90%, namun hanya mampu mendeteksi 10% kelas minoritas. Dari 139 aplikasi berisiko pada data uji, hanya 14 yang terdeteksi dan 125 lainnya terlewatkan.

Penerapan ketiga komponen optimasi secara bersamaan meningkatkan *recall* kelas minoritas dari 0,10 menjadi 0,90 dan *Macro Recall* dari 0,54 menjadi 0,84. Peningkatan ini dicapai dengan konsekuensi penurunan akurasi dari 90% menjadi 79% dan penurunan *recall* kelas mayoritas dari 0,98 menjadi 0,78. Penurunan tersebut merupakan kondisi yang tidak dapat dihindari dalam menangani data yang tidak seimbang, mengingat peningkatan sensitivitas terhadap kelas minoritas secara berdampak pada berkurangnya spesifisitas terhadap kelas mayoritas.

4 Kesimpulan

Model klasifikasi berbasis *CatBoost* berhasil dibangun melalui serangkaian tahapan yang sistematis. Tahap persiapan data mencakup pembersihan nilai tidak konsisten, penghapusan kolom redundan, penanganan nilai tidak valid, serta penyederhanaan kategori. Dilanjut dengan tahap *model building*, *dataset* dibagi secara *stratified* dengan rasio 80:10:10 untuk data latih, data validasi, dan data uji. Optimasi *hyperparameter* dilakukan menggunakan *Optuna*, hasil *hyperparameter* optimal yang telah diperoleh kemudian digunakan dalam proses pelatihan. Pelatihan model dilakukan implementasi metode *Mutually Disjoint Datasets* untuk mengatasi ketidakseimbangan kelas dengan rasio 1:9, yakni data kelas mayoritas dibagi menjadi sembilan subset yang masing-masing digabungkan dengan seluruh data kelas minoritas sehingga diperoleh sembilan sub-model. Prediksi akhir dihasilkan melalui mekanisme *soft voting* dengan *threshold* optimal hasil proses optimasi, yaitu sebesar 0,81.

Evaluasi model *CatBoost* optimasi menunjukkan bahwa model mampu mencapai akurasi sebesar 79% dengan *Macro Recall* sebesar 0,84. *Recall* pada kelas minoritas (aplikasi berisiko melanggar regulasi) mencapai 0,90, yang berarti 63 dari 70 aplikasi berisiko berhasil terdeteksi dengan benar. *Recall* pada kelas mayoritas (aplikasi tidak berisiko) mencapai 0,78. Dibandingkan dengan model tanpa optimasi mencapai akurasi 90%, namun *recall* kelas minoritas hanya sebesar 0,10. Kelas minoritas cenderung gagal terdeteksi oleh model, hanya 14 dari 139 aplikasi berisiko yang berhasil terdeteksi. Perbandingan dua kondisi tersebut menunjukkan bahwa model *CatBoost* yang dioptimasi mampu mendeteksi aplikasi digital berisiko melanggar regulasi, meskipun *dataset* memiliki karakteristik ketidakseimbangan kelas. Kemampuan ini khususnya terlihat dalam meminimalkan kasus yang berdampak pada terlewatnya aplikasi berbahaya bagi privasi anak.

Penelitian ini menggunakan *dataset* dari satu sumber kompetisi Kaggle yang memiliki keterbatasan dalam hal cakupan dan kemutakhiran data. Penelitian selanjutnya disarankan untuk menggunakan *dataset* yang lebih besar, lebih beragam, dan lebih mutakhir, termasuk kemungkinan pengumpulan data secara langsung dari platform distribusi aplikasi. Penelitian selanjutnya

juga dapat mengeksplorasi metode lain yang berpotensi memberikan kinerja lebih baik.

Ucapan Terima Kasih

Penulis mengucapkan terima kasih kepada Program Studi Sarjana Sains Data Universitas Negeri Surabaya serta seluruh pihak yang telah memberikan dukungan dan masukan selama pelaksanaan penelitian ini.

Daftar Pustaka

- Afina, S. S. and Prastyanti, R. A. (2025). Personal data protection analysis: Comparison of indonesia and the united states as federal countries. *International Journal of Law, Crime and Justice*, 2(2):245–259.
- Akiba, T., Sano, S., Yanase, T., Ohta, T., and Koyama, M. (2019). Optuna: A next-generation hyperparameter optimization framework. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery Data Mining*, pages 2623–2631. ACM.
- Andiani, F. N. and Wiraguna, S. A. (2025). Pelindungan data pribadi anak di tiktok: Kajian hukum terhadap penggunaan media sosial oleh pengguna di bawah umur.
- Badan Pusat Statistik (2025). *Profil Anak Usia Dini 2025*. Badan Pusat Statistik Indonesia.
- Caushaj, E. and Sugumaran, V. (2023). Classification and security assessment of android apps. *Discover Internet of Things*, 3(1):1–17.
- Federal Trade Commission (2025). Complying with coppa: Frequently asked questions. Diakses Juni 2026.
- Gad, I. (2025). Toca-iot: Threshold optimization and causal analysis for iot network anomaly detection based on explainable random forest. *Algorithms*, 18(2):117.
- Hakami, H. (2025). Automatic classification of mobile apps to ensure safe usage for adolescents. *PLOS ONE*, 20(1).
- Koyyada, S. P. and Singh, P. T. (2023). A multi stage approach to handle class imbalance: An ensemble method. *Procedia Computer Science*, 218:2666–2674.
- Lai, J.-P., Lin, Y.-L., Lin, H.-C., Shih, C.-Y., Wang, Y.-P., and Pai, P.-F. (2023). Tree-based machine learning models with optuna in predicting impedance values for circuit analysis. *Micromachines*, 14(2):265.
- Leevy, J. L., Johnson, J. M., Hancock, J., and Khoshgoftaar, T. M. (2023). Threshold optimization and random undersampling for imbalanced credit card data. *Journal of Big Data*, 10(1):58.
- Miftahushudur, T., Sahin, H. M., Grieve, B., and Yin, H. (2025). A survey of methods for addressing imbalance data problems in agriculture applications. *Remote Sensing*, 17(3):454.
- Nayak, J., Naik, B., Dash, P. B., Vimal, S., and Kadry, S. (2022). Hybrid bayesian optimization hypertuned catboost approach for malicious access and anomaly detection in iot framework. *Sustainable Computing: Informatics and Systems*, 36.
- Prokhorenkova, L., Gusev, G., Vorobev, A., Dorogush, A. V., and Gulin, A. (2018). Catboost: Unbiased boosting with categorical features. In *Proceedings of the 32nd Conference on Neural Information Processing Systems (NeurIPS 2018)*, Montréal, Canada.
- Putri, C. C., Sihabudin, and Ganindha, R. (2024). Legal construction of the protection and processing of children’s personal data in indonesia. In *Proceedings of the International Conference on Law and Social Justice*, pages 147–155. Atlantis Press.

- Sokolova, M. and Lapalme, G. (2009). A systematic analysis of performance measures for classification tasks. *Information Processing & Management*, 45(4):427–437.
- Stoilova, M., Livingstone, S., and Khazbak, R. (2021). Investigating risks and opportunities for children in a digital world: A rapid review of the evidence on children’s internet use and outcomes. Technical report, UNICEF Office of Research - Innocenti. Innocenti Discussion Paper 2020-03.
- Tikaningsih, A., Lestari, P., Nurhopipah, A., Tahyudin, I., Winarto, E., and Hassa, N. (2024). Optuna based hyperparameter tuning for improving the performance prediction mortality and hospital length of stay for stroke patients. *Telematika*, 17(1):1–16.
- Yandex (2025). Catboost documentation. <https://catboost.ai/docs/en/>. Diakses November 2025.
- Zakariya, R. A. I. and Ramli, K. (2023). Desain penilaian risiko privasi pada aplikasi seluler melalui model machine learning berbasis ensemble learning dan multiple application attributes. *Jurnal Teknologi Informasi dan Ilmu Komputer*, 10(4):831–842.