

## PREDIKSI KEKERASAN TERHADAP PEREMPUAN DAN ANAK DI YABIKU NUSA TENGGARA TIMUR MENGGUNAKAN METODE NAIVE BAYES

Yohana Palestrina Kobesi, Silvester Tena, Wenefrida Tulit Ina, Jani Fredie Mandala

Teknik Elektro, Universitas Nusa Cendana

Jalan Adisucipto, Penfui, Kupang, Nusa Tenggara Timur, 85228, Indonesia

yunikobesi09@gmail.com

### ABSTRAK

Kekerasan terhadap perempuan dan anak merupakan bentuk pelanggaran hak asasi manusia yang serius dan dapat terjadi di berbagai ruang baik dalam lingkungan keluarga, masyarakat, maupun institusi. Di Nusa Tenggara Timur kasus kekerasan terhadap perempuan dan anak tergolong tinggi. Lembaga YABIKU-NTT fokus menangani kasus perlindungan perempuan dan anak, berupaya untuk mengimplementasikan berbagai program agar mencegah dan mengurangi kekerasan. Namun mengalami kesulitan dalam memprediksi dan mengklasifikasikan jenis kekerasan karena dilakukan secara manual. Teknologi kecerdasan buatan khususnya metode *machine learning* dapat membantu pengolahan data untuk memprediksi pola kejadian kekerasan. Salah satu metode *machine learning* yakni *Naive Bayes* dapat mengklasifikasikan pola kekerasan berbasis probabilitas dan statistik. Metode ini dapat memberikan gambaran mengenai performa *naive bayes* dalam memprediksi potensi kekerasan. Pada penelitian ini menggunakan 100 data kasus di YABIKU-NTT dengan pembagian 80% dan 20% berturut-turut untuk data pelatihan dan data uji. Hasil penelitian menunjukkan bahwa pada kelas Psikis nilai *precision* 83%, *recall* 100%, *F1-score* 91% dan pada kelas Seksual nilai *precision* 79%, *recall* 100%, *F1-score* tinggi 88% sedangkan pada kelas fisik tidak terprediksi. Metode Naive Bayes mencapai akurasi 80% dan mampu memprediksi pola kekerasan dengan baik pada kategori psikis dan seksual, meskipun pada kekerasan fisik masih kurang akurat karena keterbatasan data.

**Kata kunci :** kekerasan perempuan dan anak, machine learning, naive bayes, prediksi

### 1. PENDAHULUAN

Kekerasan terhadap perempuan dan anak merupakan bentuk pelanggaran hak asasi manusia yang serius dan dapat terjadi di berbagai ruang baik dalam lingkungan keluarga, masyarakat, maupun institusi. Kekerasan ini bisa bersifat fisik, psikologis, seksual, ekonomi, maupun sosial, dan sering kali terjadi secara terus-menerus. Berdasarkan data statistik jumlah kasus kekerasan terhadap perempuan dan anak di provinsi NTT masih cukup tinggi. Kekerasan Dalam Rumah Tangga (KDRT) masih menjadi bentuk kekerasan yang paling banyak terjadi, baik terhadap perempuan maupun anak. Selain itu, kekerasan seksual juga merupakan bentuk kekerasan yang cukup tinggi, khususnya terhadap anak. Eksploitasi seksual juga menjadi bentuk kekerasan yang perlu mendapat perhatian, khususnya terhadap anak. Data ini menunjukkan bahwa meskipun berbagai upaya telah dilakukan, tantangan dalam mengatasi kekerasan terhadap perempuan dan anak di NTT masih sangat besar [1]. YABIKU-NTT sebagai lembaga yang berfokus pada perlindungan perempuan dan anak, berkomitmen kuat untuk mengatasi isu kekerasan yang sering kali mengancam keselamatan dan kesejahteraan kelompok rentan ini. Sejak berdirinya, lembaga YABIKU-NTT bekerja bersama masyarakat akar rumput, berupaya untuk memberi pemahaman dan membangun kesadaran mengenai perlindungan terhadap hak-hak mereka. Dalam upaya untuk mengatasi hal ini, lembaga YABIKU-NTT berupaya untuk mengimplementasikan berbagai program intervensi yang bertujuan untuk mencegah

dan mengurangi kekerasan, dan memberikan dukungan kepada korban. Program-program ini mencakup pelatihan keterampilan, penyuluhan tentang hak-hak perempuan, dan pendampingan bagi korban kekerasan. Berdasarkan survei awal, terdapat beberapa faktor yang berpotensi memengaruhi terjadinya kekerasan, seperti faktor ekonomi, kondisi keluarga, karakter yang dimiliki pelaku, serta suasana lingkungan. Namun, meskipun berbagai upaya telah dilakukan, YABIKU-NTT masih menghadapi kesulitan dalam menentukan data yang paling dominan untuk pengambilan keputusan yang efektif.

Perkembangan ilmu pengetahuan dan teknologi dapat membantu manusia untuk mengolah data sehingga mempercepat pengambilan keputusan. Teknologi kecerdasan buatan khususnya metode *machine learning* dapat membantu pengolahan data sehingga dapat memprediksi pola kejadian kekerasan. Salah satu metode *machine learning* yakni *Naive Bayes* [2] dapat digunakan untuk menganalisis data yang ada dan mengidentifikasi pola-pola yang relevan sehingga mendapatkan wawasan yang lebih mendalam tentang isu-isu yang dihadapi oleh perempuan dan anak di Nusa Tenggara Timur. Metode ini memungkinkan YABIKU-NTT untuk mengolah informasi dari berbagai sumber, seperti laporan kasus kekerasan sehingga dapat mengidentifikasi faktor-faktor risiko dan kebutuhan yang mendesak. Penelitian ini menggunakan lima variabel utama, yaitu karakteristik korban, pelaku, lokasi kejadian, faktor penyebab dan variabel jenis kekerasan. Pada penelitian ini, variabel jenis kekerasan dijadikan sebagai kelas

target yang akan diprediksi, sedangkan variabel lainnya digunakan sebagai atribut yang mempengaruhi proses klasifikasi.

Metode Naïve Bayes sangat sederhana, efisien, dan dapat menangani data yang memiliki banyak variabel. Metode ini juga dapat diandalkan untuk dataset yang terbatas. Dalam penelitian ini akan menggunakan Naïve Bayes untuk menganalisis data tentang kekerasan terhadap perempuan dan anak di YABIKU-NTT. Penelitian ini berkontribusi untuk memprediksi kemungkinan terjadinya kekerasan terhadap perempuan dan anak di masa depan. Dengan hasil penelitian ini, diharapkan dapat memberikan informasi yang berguna bagi pengambil keputusan dan pihak-pihak terkait dalam merancang program-program pencegahan yang lebih baik dan efektif.

## 2. TINJAUAN PUSTAKA

### 2.1. Penelitian Terdahulu

Beberapa penelitian terdahulu yang menunjukkan bahwa penggunaan algoritma Naïve Bayes untuk memprediksi kelulusan mahasiswa tepat waktu, dengan tingkat keberhasilan tinggi yang dipengaruhi oleh nilai IPK dan lama studi, menghasilkan precision, recall, dan akurasi berturut-turut 90%, 100% dan 90% [3]. Penelitian yang menggunakan Naïve Bayes untuk mengklasifikasikan kasus Anak Berhadapan Hukum (ABH) ke dalam kategori ringan, sedang, dan berat, dengan faktor paling berpengaruh yaitu status pernikahan (96%), hubungan dengan kepala keluarga (95%), usia (82%), status sekolah (72%), dan jenis kelamin (57%) [4]. Penelitian membahas tentang metode Naive Bayes dan penggunaan aplikasi weka dalam prediksi hasil produksi produk Shell pada PT. Terang Parts Indonesia dengan tingkat akurasi 92,7614% dan nilai error sebesar 7,2389 [5]. Penelitian ini telah berhasil membangun model untuk melakukan klasifikasi jenis kekerasan pada perempuan dan anak menggunakan algoritma Multinomial Naïve Bayes. Berdasarkan hasil pengujian model yang diperoleh hasil terbaik pada model skenario 2 dengan nilai akurasi, presisi, recall dan f-score berturut-turut 47%, 45%, 42%, 41% [6]. Penelitian lain membahas prediksi pola kasus kekerasan di Jawa Barat menggunakan teknik data mining dengan algoritma Regresi Linear. Hasil menunjukkan bahwa penerapan beberapa metode feature selection mampu menghasilkan nilai RMSE yang optimal sebagai dasar dalam perumusan strategi pencegahan dan penanganan kasus kekerasan [7]. Hasil-hasil tersebut memperkuat alasan penggunaan metode Naïve Bayes dalam penelitian ini.

### 2.2. Algoritma Naïve Bayes

Metode Naive Bayes adalah salah satu metode yang menggunakan metode perhitungan probabilitas dan statistik. Naive Bayes memiliki berbagai kegunaan dalam berbagai bidang. Algoritma ini sering digunakan untuk mengklasifikasikan dokumen teks, seperti artikel berita atau teks akademis. Dengan

menghitung probabilitas dari setiap fitur terhadap kelas tertentu, *Naive Bayes* dapat memprediksi kelas mana yang paling mungkin untuk data baru[8].

Dalam proses rumus probabilitas *Naive Bayes* adalah sebagai berikut:

$$P(C_i | X) = \frac{P(X|C_i) \cdot P(C_i)}{p(X)} \quad (1)$$

Keterangan

$X$  : Data dengan *Class* yang belum diketahui

$C_i$  : Suatu variabel yang harus dideskripsikan secara probabilistik

$P(C_i | X)$  : Probabilitas hipotesis  $C_i$  berdasarkan kondisi  $X$  (*Posterior probability*)

$P(C_i)$  : Probabilitas hipotesis  $C_i$  (*prior probability*)

$P(X | C_i)$  : Probabilitas  $X$  berdasar kondisi pada hipotesis  $C_i$

$P(X)$  : Probabilitas dari  $X$

### 2.3. Confusion Matrix

*Confusion matrix* adalah sebuah tabel yang digunakan untuk menampilkan jumlah data uji yang diklasifikasikan dengan benar dan salah, sehingga memudahkan dalam mengevaluasi akurasi suatu sistem klasifikasi [9].

Tabel 1. Confusion Matrix

| Kelas Aktual | Positif                                                                        | Negatif                                                                         |
|--------------|--------------------------------------------------------------------------------|---------------------------------------------------------------------------------|
| Positif      | <i>True Positif (TP)</i><br>Data yang bersifat positif dan terdeteksi positif  | <i>False Positive (FP)</i><br>Data yang bersifat positif dan terdeteksi negatif |
| Negatif      | <i>False Negatif (FN)</i><br>Data yang bersifat negatif dan terdeteksi positif | <i>True Negatif (TN)</i><br>Data yang bersifat negatif dan terdeteksi negatif   |

Dengan menggunakan Metrik *Confusion Matrix* pada Tabel 1, kita dapat menghitung berbagai metrik evaluasi yang penting untuk menilai efektivitas model jumlah data uji yang diklasifikasikan dengan benar dan salah, sehingga memudahkan dalam mengevaluasi akurasi suatu sistem klasifikasi [9].

Akurasi adalah proporsi jumlah prediksi yang benar dibandingkan dengan total prediksi yang dilakukan oleh model[10] dengan persamaan (2).

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

Presisi mengukur proporsi prediksi positif yang benar dari semua prediksi positif yang dibuat oleh model dengan persamaan (3).

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

*Recall* mengukur proporsi prediksi positif yang benar dari semua kasus positif yang sebenarnya dalam data dengan persamaan (4).

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

Setelah hasil klasifikasi dapat diukur kebenarannya, maka dilakukan perhitungan kombinasi nilai untuk dijadikan sebagai nilai pengukuran (*F1-score*). *F1-score* dapat dihitung dengan persamaan (5).

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

Kurva *Receiver Operating Characteristic* (ROC)-*Area Under the Curve* (AUC) digunakan untuk menilai hasil prediksi. Grafik ROC digunakan untuk mengevaluasi kemampuan model dalam membedakan antara setiap kelas jenis kekerasan, yaitu kekerasan fisik, psikis, dan seksual, berdasarkan nilai probabilitas hasil prediksi yang dihasilkan oleh model. Nilai AUC yang cukup tinggi pada sebagian besar kelas, menandakan bahwa model memiliki kemampuan klasifikasi yang baik dalam membedakan jenis kekerasan. Semakin besar nilai AUC (mendekati 1), maka semakin baik kinerja model dalam memprediksi kasus kekerasan dengan tepat. Penelitian uji diagnostik akan semakin baik bila nilai AUC mendekati 1. Hasil AUC dapat dibagi menjadi beberapa kelompok (0,90 - 1,00) = Klasifikasi Sangat Baik, (0,80 - 0,90) = Klasifikasi Baik, (0,70 - 0,80) = Klasifikasi Sedang, (0,60 - 0,70) = Klasifikasi Buruk, (0,50 - 0,60) = Kegagalan [11], [12].

## 2.4. Python

*Python* sering digunakan untuk scripting, analisis statistik, dan aplikasi dan pengembangan Web serba guna. Ini dapat digunakan secara interaktif atau terprogram. Banyaknya pustaka basis data, matematika, dan grafik yang cocok untuk proyek dalam eksplorasi data dan masalah data mining merupakan kekuatan pemrograman *Python*.

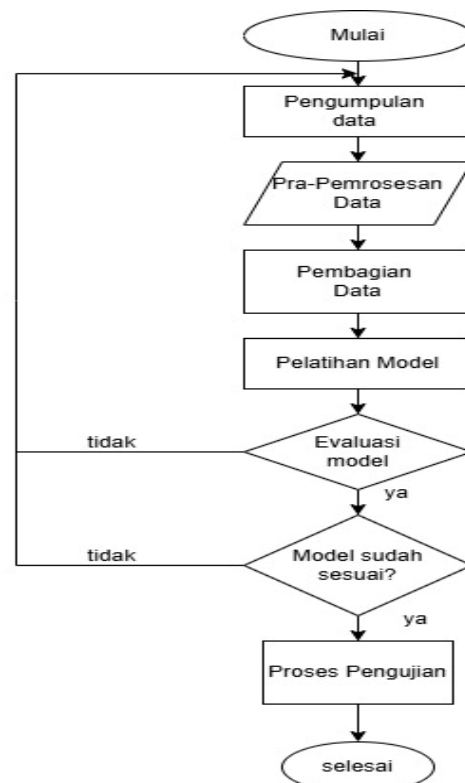
## 3. METODE PENELITIAN

Tahapan proses dalam penelitian ditunjukkan pada Gambar 1.

Penelitian ini terdiri dari beberapa tahapan utama yang mencakup pengumpulan data dengan jenis data yang diambil berupa data primer dan sekunder yang dikumpulkan dari Kantor Yabiku-NTT.

Tabel 2. Data Sekunder

| Variabel             | Label/isi Data                           |
|----------------------|------------------------------------------|
| Karakteristik Korban | Usia dan jenis kelamin                   |
| Karakteristik Pelaku | Usia dan jenis Kelamin                   |
| Jenis kekerasan      | Fisik, psikis, seksual                   |
| Lokasi kejadian      | Kota/Desa                                |
| Faktor               | Character, lingkungan, ekonomi, keluarga |



Gambar 1. Alur penelitian

Data sekunder pada tabel 2 diperoleh dari dokumentasi kasus kekerasan yang tersedia. Sementara itu, data primer diperoleh melalui Wawancara yang dilakukan untuk memperoleh informasi mengenai jenis kekerasan yang paling sering terjadi berdasarkan pengalaman mereka. Informasi dari wawancara ini turut memperkuat pemahaman terhadap pola-pola kekerasan yang terjadi. Dilanjutkan dengan, pemrosesan dan pembagian data, pelatihan dan pengujian model dengan metode *Naïve Bayes* [13], [14].

### 3.1. Metode *Naïve Bayes*

Penelitian ini menggunakan metode *Naïve Bayes* sebagai metode klasifikasi untuk memprediksi kasus kekerasan. Alur Proses *Naïve Bayes* secara detail dapat dilihat pada Gambar 2.

Tahapan proses *Naïve Bayes* dapat dijelaskan sebagai berikut:

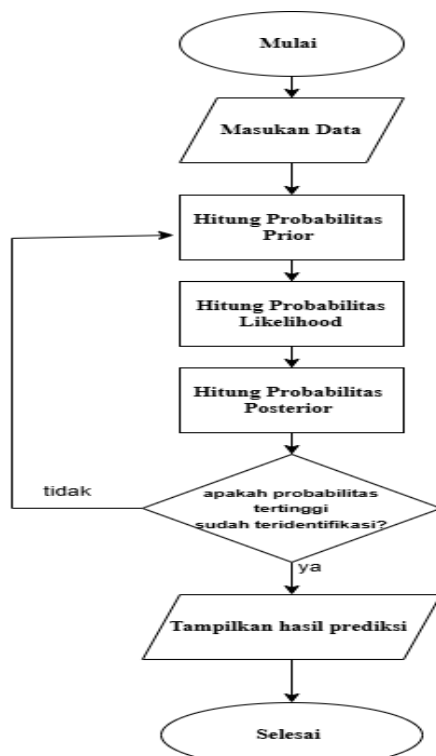
#### a. Masukan Data

Untuk menggunakan metode *Naïve Bayes* dalam memprediksi kasus kekerasan, data yang telah diproses terdiri dari beberapa fitur (variabel) yang telah diubah menjadi format numerik atau kategorikal.

#### b. Menghitung *Probabilitas Prior*

*Probabilitas Prior* (*Prior Probability*) yaitu *Probabilitas* sebelum kita melihat data. *Probabilitas* setiap kelas dihitung dari distribusi data pelatihan [15]. Klasifikasi kelas dibagi menjadi 3 yakni  $C_1$  untuk kekerasan fisik,  $C_2$  untuk kekerasan psikis, dan  $C_3$  untuk kekerasan seksual.

- $P(C_1)$ = probabilitas awal suatu kasus merupakan kekerasan fisik  
 $P(C_2)$ = probabilitas awal suatu kasus merupakan kekerasan psikis  
 $P(C_3)$ = probabilitas awal suatu kasus merupakan kekerasan seksual



Gambar 2. Alur Naïve Bayes

#### c. Menghitung Probabilitas *Likelihood*

Probabilitas *Likelihood* yakni Probabilitas data ( $X$ ) muncul dengan asumsi bahwa kelas/hipotesis  $C_i$  benar [15]. Hal ini berarti probabilitas setiap atribut muncul dalam kelas tertentu. Seberapa besar kemungkinan kita melihat data ini jika berasal dari kelas tersebut. Dalam konteks *Naïve Bayes*, simbol  $X$  merujuk pada fitur-fitur (variabel) yang digunakan sebagai dasar untuk melakukan prediksi kelas.

- $P(X | C_1) = P(X | \text{fisik})$ , yakni Probabilitas fitur-fitur muncul jika kasus adalah kekerasan fisik.  
 $P(X | C_2) = P(X | \text{psikis})$ , adalah Probabilitas fitur-fitur muncul jika kasus adalah kekerasan psikis.  
 $P(X | C_3) = P(X | \text{seksual})$ , adalah Probabilitas fitur-fitur muncul jika kasus adalah kekerasan seksual.

#### d. Menghitung Probabilitas *Posterior*

Probabilitas *Posterior* (*Posterior Probability*) yakni probabilitas yang mengukur seberapa besar kemungkinan suatu kelas terjadi setelah mempertimbangkan evidence yang telah diamati [16]. Probabilitas suatu hipotesis  $C_i$  setelah melihat data  $X$ . Tujuan akhir dari *Naïve Bayes* adalah memprediksi kelas berdasarkan fitur yang diberikan. Hasil akhir dari kombinasi *prior* dan *likelihood* dihitung dengan persamaan (1).

$$P(\text{kelas} | X) = \frac{P(X | \text{kelas}) \times P(\text{kelas})}{\sum \text{setiap kelas } P(X | \text{kelas}) \times P(\text{kelas})}$$

Simbol  $X$  merujuk pada sekumpulan fitur atau atribut dari suatu kasus kekerasan yang digunakan sebagai dasar klasifikasi. Sedangkan  $P$  menunjukkan probabilitas. Dalam konteks penelitian ini, fitur-fitur tersebut dapat mencakup variabel data sekunder pada Tabel 1. Fitur-fitur ini digunakan untuk memprediksi jenis kekerasan (fisik, psikis, atau seksual) pada data baru dengan pendekatan *Naïve Bayes*.

#### e. Pemilihan Probabilitas Tertinggi

Langkah terakhir dalam proses metode *Naïve Bayes* adalah memilih kelas dengan nilai probabilitas tertinggi. Setelah menghitung probabilitas *posterior* berdasarkan data fitur, nilai tersebut menunjukkan kemungkinan suatu kasus termasuk dalam kategori kekerasan, seperti fisik, psikis, atau seksual. Pemilihan kelas didasarkan pada nilai *posterior* terbesar, yang menunjukkan keyakinan tertinggi berdasarkan kombinasi probabilitas awal dan *likelihood*. Hasil klasifikasi ditentukan oleh kelas yang paling sesuai dengan pola data yang diamati.

### 3.2. Evaluasi Kinerja sistem

Tahapan evaluasi dilakukan setelah menghitung probabilitas yang ditampilkan dalam bentuk *confusion matrix*. *Confusion matrix* adalah sebuah tabel yang digunakan untuk menampilkan jumlah data uji yang diklasifikasikan dengan benar dan salah, sehingga memudahkan dalam mengevaluasi akurasi suatu sistem klasifikasi [9]. Berdasarkan *Confusion Matrix*, dapat dihitung beberapa metrik evaluasi yang penting untuk menilai efektivitas model klasifikasi, yakni akurasi, presisi, *recall*, dan *F1-Score*, *Receiver Operating Characteristic* (ROC)-Area Under the Curve (AUC) [17].

## 4. HASIL DAN PEMBAHASAN

Dataset yang digunakan dalam penelitian ini terdiri dari 100 kasus kekerasan terhadap perempuan dan anak, yang ditangani oleh lembaga YABIKU-NTT sejak Tahun 2020 hingga 2025. Dataset tabel 3 berisikan atribut data yang akan diolah pada proses klasifikasi dengan metode *Naïve Bayes*.

Tabel 3. Dataset

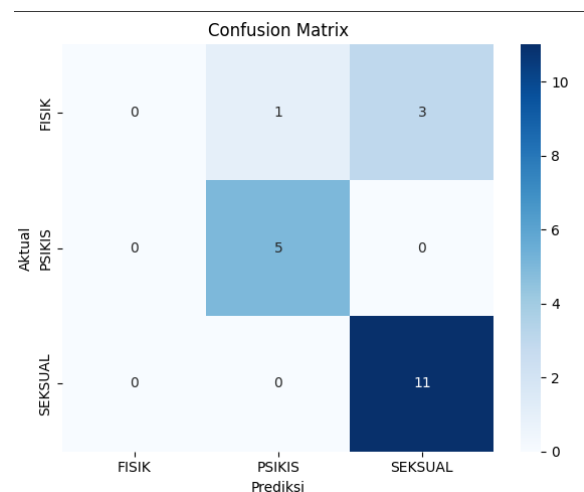
| U Korban | Jk Korban | U Pelaku | Jk Pelaku | Jenis Kekerasan | Lokasi | Faktor     |
|----------|-----------|----------|-----------|-----------------|--------|------------|
| 19       | P         | 49       | L         | Seksual         | Desa   | Karakter   |
| 17       | P         | 24       | L         | Psikis          | Kota   | Lingkungan |
| 20       | P         | 44       | L         | Fisik           | Kota   | Keluarga   |
| 20       | P         | 44       | L         | Seksual         | Kota   | Karakter   |
| 38       | P         | 40       | L         | Seksual         | Desa   | Karakter   |
| 23       | P         | 28       | L         | Psikis          | Kota   | Lingkungan |
| 23       | P         | 31       | L         | Fisik           | Kota   | Ekonomi    |
| 19       | P         | 31       | L         | Seksual         | Desa   | Karakter   |
| 23       | P         | 35       | L         | Seksual         | Kota   | Karakter   |
| 30       | P         | 28       | L         | Psikis          | Kota   | Lingkungan |
| 28       | P         | 34       | L         | Fisik           | Desa   | Karakter   |
|          |           |          |           |                 |        |            |
| 20       | P         | 25       | L         | Seksual         | Kota   | Karakter   |

Dalam penelitian ini data dibagi menggunakan metode *split* sederhana (80:20) [18] [19]. Sebanyak 80 data dialokasikan sebagai himpunan pelatihan (*training set*) untuk membangun dan melatih algoritma *machine learning*, sedangkan 20 data lainnya dijadikan data uji (*testing set*) untuk memeriksa keakuratan dan kemampuan model. Algoritma *Naïve Bayes* untuk klasifikasi diproses dengan program bantu yakni phyton melalui *google colab* untuk memperoleh tingkat akurasi model dan visualisasinya. Dalam implementasi dengan menggunakan program python, model dilatih dan dievaluasi menggunakan beberapa parameter utama. Untuk pembagian dataset menggunakan parameter *test\_size* = 0.2. Model yang digunakan adalah *Gaussian Naive Bayes* (*GaussianNB()*), yang memiliki parameter internal seperti *priors* dan *var\_smoothing* untuk perhitungan probabilitas serta menjaga stabilitas numerik dalam distribusi Gaussian. Pengaturan *random\_state* = 42 digunakan untuk menjaga konsistensi pembagian data dan pembagian fold dalam K-Fold Cross Validation.

Proses training data merupakan tahap penting dalam pembangunan model klasifikasi karena pada tahap ini algoritma mempelajari pola dan karakteristik data yang digunakan. Proses melatih data dimulai dengan mencari peluang awal (*prior probability*) dari setiap kategori data. Pada penelitian ini, prediksi kekerasan terhadap perempuan dan anak dibagi menjadi tiga kelas data yaitu kelas fisik, psikis, dan seksual. Setiap kelas memiliki jumlah data yang berbeda, dengan kelas jenis kekerasan fisik, psikis, dan seksual berturut-turut 4 data, 5 data, dan 11 data. Perbedaan jumlah data pada masing-masing kelas mencerminkan variasi kasus yang ditemukan selama pengumpulan data. Hasil perhitungannya memberikan nilai peluang awal bagi setiap kelas sebelum mempertimbangkan atribut lain, dan menjadi dasar untuk tahapan perhitungan *likelihood* dan *posterior probability*. Selanjutnya, algoritma menghitung *likelihood*, yaitu probabilitas kemunculan setiap atribut atau nilai fitur terhadap masing-masing kelas. Perhitungan *likelihood* dilakukan berdasarkan distribusi data latih. Untuk atribut numerik, *likelihood*

dihitung menggunakan parameter statistik seperti nilai rata-rata (*mean*) dan *varians* [20]. Sedangkan untuk atribut kategorikal dihitung berdasarkan frekuensi kemunculan dengan penerapan teknik *smoothing* untuk menghindari probabilitas nol [21].

Berdasarkan nilai *prior* dan *likelihood* yang telah diperoleh, algoritma kemudian menghitung probabilitas posterior (*posterior probability*) untuk setiap kelas. Probabilitas *posterior* merupakan peluang suatu data termasuk ke dalam kelas tertentu setelah mempertimbangkan seluruh atribut yang dimiliki. Nilai *posterior* diperoleh dari hasil perkalian antara *prior* dan seluruh *likelihood* dari atribut yang diamati. Kelas dengan nilai posterior tertinggi dipilih sebagai hasil prediksi. Hasil prediksi setiap kelas ditampilkan dalam bentuk *confusion matrix* pada Gambar 4.



Gambar 3. Hasil Confusion Matrix

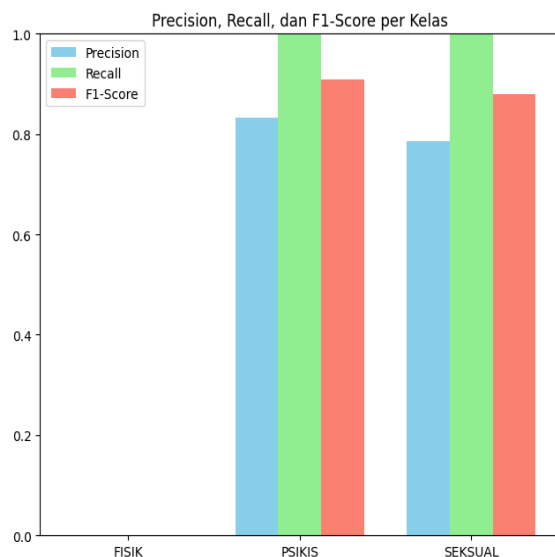
Berdasarkan *confusion matrix* pada Gambar 4 Total data yakni 20 sampel (4 Fisik, 5 Psikis, 11 Seksual). Pada kelas Fisik, Tidak ada yang benar (0/4) semua salah klasifikasi. Jumlah 1 salah prediksi Psikis, 3 salah prediksi Seksual. Artinya model gagal mengenali pola khas dari kasus Fisik. Hal ini diakibatkan karena jumlah data Fisik terlalu sedikit (*class imbalance*). Pada kelas Psikis, 5 benar dari 5

(100%) dengan performa sangat baik. Tidak ada *false positive* maupun *false negative*. Pada kelas Seksual, 11 benar dari 11 (100%) performa sempurna pada kelas ini karena tidak ada salah klasifikasi.

Model sangat kuat dalam mengenali Psikis dan Seksual, tapi masih gagal mendeteksi Fisik. Pola ini menandakan bahwa data *training* belum seimbang atau ciri khas (fitur) kasus Fisik terlalu mirip dengan dua kelas lain. Sehingga terdapat akurasi total sebagai berikut:

$$accuracy = \frac{TP+TN}{TP+FN+FP+TN} \times 100\% = \frac{16}{20} = 0,8 = 80\%$$

Model Naïve Bayes menunjukkan kemampuan yang cukup baik dalam melakukan klasifikasi jenis kekerasan terhadap perempuan dengan tingkat akurasi sebesar 80%. Model memberikan hasil terbaik pada kategori psikis dan seksual, sedangkan kategori fisik masih sering salah diklasifikasikan karena jumlah datanya lebih sedikit dibandingkan kelas lainnya.

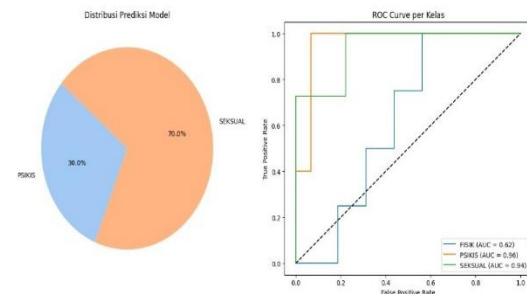


Gambar 4. Hasil pengukuran *Precision*, *Recall*, dan *F1-Score*

Berdasarkan grafik pada Gambar menunjukkan bahwa Fisik tidak terprediksi, jadi *precision* dan *recall* = 0. Psikis *precision* 83% dan semua prediksi Psikis benar, *recall* 100% semua Psikis terdeteksi, *F1-score* tinggi 91%. Seksual *precision* 79% ada beberapa salah prediksi dari kelas lain ke Seksual, *recall* 100% tidak ada Seksual yang terlewat dan *F1-score* tinggi (0,88). Model unggul dalam mendeteksi Psikis dan Seksual dengan tingkat kesalahan kecil. Nilai *F1-score* yang tinggi menunjukkan keseimbangan antara presisi. Namun Fisik tetap menjadi titik lemah utama model.

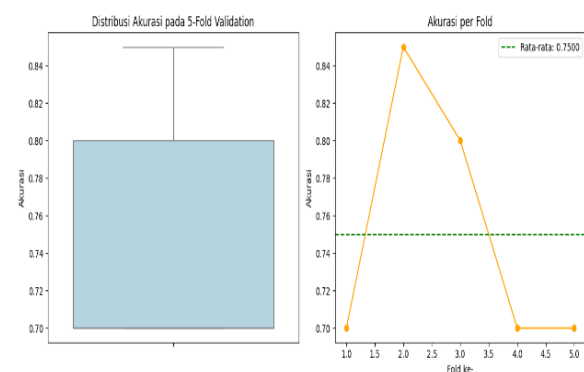
Berdasarkan Gambar 6 distribusi prediksi model tidak ada prediksi Fisik model cenderung hanya memilih dua kelas yakni Psikis dan Seksual. Ketimpangan ini menunjukkan bias model terhadap kelas dengan data lebih banyak (Seksual). Dan

berdasarkan kurva ROC di mana nilai AUC, Fisik (0,62) kemampuan membedakan antara Fisik dan non-Fisik sangat rendah, skala klasifikasi buruk. Psikis (0,96) sangat baik, hampir sempurna dalam membedakan Psikis dari kelas lain, skala klasifikasi sangat baik. Seksual (0,94) menunjukkan model sangat yakin untuk kasus ini, skala klasifikasi sangat baik. Model sangat baik untuk dua kelas (Psikis dan Seksual), tapi masih belum belajar baik pada Fisik.



Gambar 5. Grafik Distribusi Prediksi Model dan Grafik ROC-AUC

Pengujian kemampuan model yakni dengan penambahan *K-fold cross Validation*. *Cross Validation* merupakan teknik *resampling* untuk menyesuaikan parameter data saat membangun model [22]. Semua data dibagi menjadi beberapa bagian atau biasa disebut *fold*. Model kemudian dilatih menggunakan sebagian data, dan diuji menggunakan bagian data yang lain. Proses ini dilakukan berulang-ulang sampai semua bagian data pernah menjadi data uji. Dengan cara ini, hasil yang diperoleh akan lebih adil dan mendapatkan hasil akurasi yang maksimal [23].



Gambar 6. Distribusi Akurasi 5-fold.

Pada pengujian 5-Fold, rata-rata akurasi model adalah 75%. Dari grafik pada Gambar 4.10, terlihat bahwa nilai akurasi berkisar antara 70% hingga 85%. Grafik kiri menunjukkan penyebaran akurasi di tiap *fold*, dan grafik kanan memperlihatkan perubahan akurasi di setiap *fold*. Hasil ini menunjukkan bahwa model cukup stabil karena perbedaan nilai akurasi antar *fold* tidak terlalu besar.

## 5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian yang telah dilakukan menggunakan metode *Naïve Bayes* menunjukkan performa yang baik dalam memprediksi potensi kekerasan terhadap perempuan dan anak. Hasil penelitian menunjukkan bahwa pada kelas Psikis nilai *precision* 83%, *recall* 100%, *F1-score* 91%, dan pada kelas Seksual nilai *precision* 79%, *recall* 100%, *F1-score* 88%, sedangkan pada kelas fisik tidak terprediksi. Metode *Naïve Bayes* mencapai akurasi 80% dan mampu memprediksi pola kekerasan dengan baik, khususnya pada kategori psikis dan seksual, meskipun prediksi pada kekerasan fisik masih kurang akurat karena keterbatasan data. Jumlah data kategori fisik terbatas sehingga terjadi salah prediksi. Pada penelitian selanjutnya perlu melakukan penambahan jumlah data dan memperhatikan keseimbangan antar kelas sehingga model mampu mengenali seluruh kategori kekerasan dengan lebih akurat, terutama kategori fisik. Rentang waktu pengambilan data kasus lebih panjang misalnya 10 atau 20 tahun serta mengembangkan metode klasifikasi dengan metode lain sehingga dapat membandingkan tingkat akurasi.

## DAFTAR PUSTAKA

- [1] Dinas Pemberdayaan Perempuan dan Perlindungan Anak Provinsi Nusa Tenggara Timur, "Laporan Tahunan DP3A Tahun 2023," 2024.
- [2] S. Sinaga, R. W. Sembiring, and S. Sumarno, "Penerapan Algoritma Naive Bayes untuk Klasifikasi Prediksi Penerimaan Siswa Baru," *Journal of Machine Learning and Data Analytics (MALDA)*, vol. 1, no. 1, pp. 55–4, Jan. 2022.
- [3] Y. Apridiansyah, N. D. M. Veronika, and E. D. Putra, "Prediksi Kelulusan Mahasiswa Fakultas Teknik Informatika Universitas Muhammadiyah Bengkulu Menggunakan Metode Naive Bayes," *JSAI: Journal Scientific and Applied Informatics*, vol. 4, no. 2, pp. 236–247, 2021, doi: 10.36085.
- [4] Nanda Clariza Febriandini, "Pengklasifikasian Anak Berhadapan Hukum (ABH) Di Kota Balikpapan Dengan Penerapan Metode Naive Bayes," *Seminar Nasional Sains Data 2024 (SENADA 2024) UPN "Veteran" Jawa Timur*, pp. 755–766, 2024.
- [5] D. Sunandar and Nurhayati, "Penerapan Algoritma Naive Bayes Untuk Mengetahui Kualitas Produksi Helmet Honda Pada PT Terang Parts Indonesia," *JORAPI: Journal of Research and Publication Innovation*, vol. 1, no. 3, 2023.
- [6] G. Subroto, N. Sulistiyowati, and A. A. Ridha, "Classification of Types of Violence Against Women And Children Using The Multinomial Naive Bayes Algorithm," *Journal of Information Technology and Computer Science (INTECOMS)*, vol. 5, no. 1, Jun. 2022.
- [7] Y. Yulia, R. Kurniawan, and Y. A. Wijaya, "OPTIMALISASI PARAMETER FEATURE SELECTION PADA PREDIKSI KASUS KEKERASAN DI JAWA BARAT MENGGUNAKAN ALGORITMA REGRESI LINEAR," *Jurnal Mahasiswa Teknik Informatika*, vol. 7, no. 6, 2023.
- [8] A. Felicia Watratan, A. B. Puspita, and D. Moeis, "Implementasi Algoritma Naive Bayes Untuk Memprediksi Tingkat Penyebaran Covid-19 Di Indonesia," *JOURNAL OF APPLIED COMPUTER SCIENCE AND TECHNOLOGY (JACOST)*, vol. 1, no. 1, pp. 7–14, 2020, [Online]. Available: <http://journal.isas.or.id/index.php/JACOST>
- [9] R. Nurhidayat and K. E. Dewi, "Penerapan Algoritma K-Nearest Neighbor dan Fitur Ekstraksi N-Gram dalam Analisis Sentimen Berbasis Aspek," *KOMPUTA: Jurnal Ilmiah Komputer dan Informatika*, vol. 12, no. 1, pp. 91–100, Apr. 2023, [Online]. Available: <https://www.kaggle.com/datasets/hafidahmustha/anah/skincare-review?select=00.+Review.csv>.
- [10] A. A. Elrahman, "Predicting Adults' Income using Naive Bayes Classifier," Jan. 2024, doi: 10.21203/rs.3.rs-3890646/v1.
- [11] Achmad Ridwan, "Penerapan Algoritma Naive Bayes Untuk Klasifikasi Penyakit Diabetes Mellitus," *Jurnal Sistem Komputer dan Kecerdasan Buatan*, vol. 4, pp. 15–21, Sep. 2020.
- [12] A. Oktaviyani, A. Heryati, and M. F. Alie, "Penerapan Metode Naive Bayes Untuk Klasifikasi Kategori Olah Pangan (Studi Kasus Dinas Kesehatan Kota Palembang)," Jun. 2024.
- [13] N. J. Fauziah, F. Rahmania, M. Daniyal, and N. F. A. T. Sari, "Analisis dan Optimalisasi Performa Algoritma Gaussian Naive Bayes pada Prediksi Metabolic Syndrome Menggunakan SMOTE," *Jurnal Informatika Sunan Kalijaga*, vol. 9, no. 2, pp. 112–122, 2024.
- [14] Muqorobin and M. B. Pakarti, "Sistem Prediksi Lama Studi Kuliah Menggunakan Metode Naive Bayes," *Jurnal Informatika, Komputer dan Bisnis*, vol. 2, [Online]. Available: <https://jurnal.itbaas.ac.id/index.php/jikombis>
- [15] Yusril Haza Mahendra, Ririen Kusumawati, and Imamudin, "Analisis Data Mining Untuk Deteksi Diabetes Mellitus Menggunakan Naive Bayes," *Jurnal CoSciTech (Computer Science and Information Technology)*, vol. 6, no. 1, pp. 39–44, May 2025, doi: 10.37859/coscitech.v6i1.7954.
- [16] T. Salsabilla and D. Alita, "Analisis Sentimen Inses di Social Media menggunakan Algoritma Naive Bayes," *Jurnal Informatika: Jurnal Pengembangan IT*, vol. 9, no. 3, pp. 271–280, Dec. 2024, doi: 10.30591/jpit.v9i3.6611.
- [17] A. Cahyana, E. R. Susanto, and parjito, "Penerapan Algoritma XGBoost untuk Prediksi Diabetes: Analisis Confusion Matrix dan ROC

- Curve,” *Fountain of Informatics Journal*, vol. 10, no. 1, pp. 40–50, May 2025, doi: 10.21111/fij.v10i1.14311.
- [18] A. Byna and M. Basit, “Penerapan Metode Adaboost Untuk Mengoptimasi Prediksi Penyakit Stroke Dengan Algoritma Naïve Bayes,” *Jurnal Sisfokom (Sistem Informasi dan Komputer)*, vol. 9, no. 3, pp. 407–411, Nov. 2020, doi: 10.32736/sisfokom.v9i3.1023.
- [19] Nazla Chairina, “Penggunaan Metode Naives Bayes Untuk Memprediksi Kelulusan Mahasiswa Pendidikan Teknologi Informasi Menggunakan Rapidminer,” 2023.
- [20] T. Destiana, Y. Umaidah, and U. Enri, “Penerapan Algoritma Gaussian Naive Bayes Dalam Penentuan Prioritas Rehabilitasi Daerah Aliran Sungai Berdasarkan Parameter Lahan Kritis,” *INFOTECH journal*, vol. 9, no. 2, pp. 507–513, Aug. 2023, doi: 10.31949/infotech.v9i2.6501.
- [21] I. Listiowarni and E. R. Setyaningsih, “Analisis Kinerja Smoothing pada Naive Bayes untuk Pengkategorian Soal Ujian,” *Jurnal Teknologi&Manajemen Informatika*, vol. 4, p. 6, 2018.
- [22] R. N. Irawan, K. M. Hindrayani, and M. Idhom, “Penerapan Cross Validation sebagai Analisis Sentimen Pelayanan Publik Kereta Api Lokal Daop 8 Menggunakan Metode Multinomial Naïve Bayes,” *G-Tech: Jurnal Teknologi Terapan*, vol. 8, no. 2, pp. 954–963, Apr. 2024, doi: 10.33379/gtech.v8i2.4117.
- [23] E. Febriyani and H. Februariyanti, “Analisis Sentimen Terhadap Program Kampus Merdeka Menggunakan Algoritma Naive Bayes Classifier Di Twitter,” *Jurnal TEKNO KOMPAK*, vol. 17, no. 1, pp. 25–8.