

Analisis Sentimen Program Coding Anak SD Menggunakan Metode Naive Bayes

Intan Gustina¹, Aditia Yudhistira^{*2}

^{1,2}Sistem Informasi, Fakultas Teknik dan Ilmu Komputer, Universitas Teknokrat Indonesia, Indonesia
Email: 1intan_gustina@teknokrat.ac.id, 2aditiayudhistira@teknokrat.ac.id

Abstrak

Pemanfaatan teknologi dalam pendidikan anak usia dini semakin berkembang, termasuk pembelajaran coding yang diklaim mampu meningkatkan keterampilan berpikir analitis dan sistematis. Penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap pembelajaran coding anak dengan menggunakan data dari media sosial Twitter. Untuk analisis ini, menggunakan metode Naive Bayes untuk mencapai temuan yang lebih komprehensif, penelitian ini mengevaluasi akurasi klasifikasi sebanyak 4.357 tweet terdiri 1.869 tweet negatif, dan 750 tweet positif. Data diproses melalui tahapan *preprocessing* seperti pembersihan teks, tokenisasi, dan normalisasi. Klasifikasi dilakukan menggunakan tiga varian *Naive Bayes* yaitu *MultinomialNB*, *GaussianNB*, dan *BernoulliNB*. yang berfokus pada analisis sentimen positif dan negative. Dari ketiga metode *Naive Bayes*, *MultinomialNB* memiliki akurasi terbaik dalam mengenali sentimen negatif, dengan precision 72% dan recall 99%, yang berarti hampir semua tweet negatif terdeteksi dengan benar. Namun, metode ini kurang efektif dalam mengenali sentimen positif, dengan recall hanya 3%, sehingga banyak opini positif yang terabaikan. *BernoulliNB* memiliki pola serupa dengan recall 96% untuk sentimen negatif tetapi hanya 21% untuk sentimen positif, menunjukkan bias terhadap opini negatif. Sementara itu, *GaussianNB* memiliki precision lebih rendah untuk sentimen positif (34%) dibandingkan sentimen negatif (76%), dengan recall yang kurang stabil, yaitu 60% untuk positif dan 53% untuk negatif. Dari hasil ini, *MultinomialNB* tetap menjadi pilihan terbaik untuk analisis sentimen, terutama dalam mendeteksi opini negatif secara real-time.

Kata kunci: analisis sentimen, coding anak, naive bayes classifier, pendidikan digital

Sentiment Analysis of Elementary School Children's Coding Program Using the Naive Bayes Method

Abstract

The use of technology in early childhood education is increasingly developing, including coding learning which is claimed to be able to improve analytical and systematic thinking skills. This study aims to analyze public sentiment towards children's coding learning using data from Twitter social media. For this analysis, using the Naive Bayes method to achieve more comprehensive findings, this study evaluated the classification accuracy of 4,357 tweets consisting of 1,869 negative tweets and 750 positive tweets. Data was processed through preprocessing stages such as text cleaning, tokenization, and normalization. Classification was carried out using three Naive Bayes variants, namely *MultinomialNB*, *GaussianNB*, and *BernoulliNB*. which focuses on analyzing positive and negative sentiments. Of the three Naive Bayes methods, *MultinomialNB* has the best accuracy in recognizing negative sentiment, with a precision of 72% and a recall of 99%, which means that almost all negative tweets are detected correctly. However, this method is less effective in recognizing positive sentiment, with a recall of only 3%, so that many positive opinions are ignored. *BernoulliNB* has a similar pattern with a recall of 96% for negative sentiment but only 21% for positive sentiment, indicating a bias towards negative opinions. Meanwhile, *GaussianNB* has lower precision for positive sentiment (34%) than negative sentiment (76%), with less stable recall, namely 60% for positive and 53% for negative. From these results, *MultinomialNB* remains the best choice for sentiment analysis, especially in detecting negative opinions in real-time.

Keywords: digital education, kids coding, naive bayes classifier, sentiment analysis

1. PENDAHULUAN

Salah satu cara yang efektif untuk mencegah kecanduan game pada anak-anak adalah dengan mengikutsertakan anak ke dalam kursus pembelajaran coding, dengan belajar coding anak tidak hanya bermain game, tetapi juga melatih anak dalam berpikir secara sistematis. Coding merupakan salah satu kegiatan yang dapat

digunakan untuk mengembangkan kemampuan berpikir secara komputasi seperti berpikir analitik, pemecahan masalah, kebiasaan dan pendekatan yang digunakan dalam ilmu komputer. Pembelajaran coding meliputi logika, analisis, berpikir kritis, dan pemecahan masalah. Dalam lingkungan pemrograman terbuka, anak-anak mempelajari dasar ilmu komputer, seperti algoritma, debugging, dan modularitas dengan merakit blok-blok pemrograman.[1]

Dalam penelitian melakukan analisis sentimen untuk melihat persepsi masyarakat terkait informasi seputar dampak yang terjadi dari pemanfaatan teknologi dalam pendidikan anak usia dini melalui media sosial di Twitter menggunakan berbagai algoritma machine learning dengan mengelompokkan opini dengan nilai positif atau negatif, dengan nilai positif atau negatif. Data hasil penelitian ini nantinya dapat digunakan untuk perbaikan atau evaluasi di masa mendatang, seperti pemahaman yang lebih baik mengenai dampak pembelajaran dalam pendidikan anak usia dini dan memberikan panduan praktis bagi Pendidikan anak dan orang tua anak. Memilih dan menggunakan pembelajaran coding dengan bijak, serta dengan memahami sentimen dan pendapat pengguna terhadap teknologi pendidikan, penelitian ini dapat membantu para pengembang teknologi dalam meningkatkan desain, fungsionalitas, dan efektivitas produk mereka untuk anak-anak pada usia dini.[2]

Analisis sentimen adalah sebuah metode yang digunakan untuk mengekstrak data opini, memahami serta mengolah tekstual data secara otomatis untuk melihat sentimen yang terkandung dalam sebuah opini. banyaknya opini masyarakat pada twitter yang membahas mengenai coding anak di twitter. Opini tersebut bersifat subyektif. Pendapat tersebut merupakan data teks yang dapat dianalisis dan digunakan untuk memperoleh informasi guna memantau sentimen masyarakat terhadap media sosial di Twitter dengan menggunakan teknik text mining, yaitu analisis sentimen. Sentimen tersebut dikategorikan menjadi sentimen positif, netral, dan negatif. Dalam penelitian ini, proses klasifikasi sentimen pengguna Twitter dimodelkan untuk mengetahui pengkodean mana yang bermanfaat bagi anak sekolah dengan menggunakan metode Naive Bayes Classification (NBC).[3]

Keterbatasan data empiris mengenai respons siswa terhadap metode belajar aktif dalam program coding anak menjadi salah satu masalah utama yang perlu diatasi. Tanpa pemahaman yang jelas tentang persepsi siswa, sulit untuk mengevaluasi efektivitas metode ini dalam meningkatkan kemampuan pemecahan masalah mereka. Oleh karena itu, diperlukan penelitian berbasis analisis sentimen yang dapat mengungkap opini siswa secara lebih mendalam. Metode Naive Bayes, yang telah terbukti akurat dalam klasifikasi sentimen, dapat digunakan untuk menganalisis data ini, memberikan wawasan yang mendukung pengembangan kurikulum yang lebih efektif..

Dalam penelitian yang dilakukan oleh Fajar Ratnawati, mengenai penerapan algoritma Naive Bayes untuk analisis sentimen terhadap opini film di Twitter, tujuan penelitian ini tercapai dengan baik. Penerapan algoritma Naive Bayes Classifier dalam menganalisis sentimen dari data opini film berbahasa Indonesia di Twitter berhasil dilakukan. Hasil pengujian dan analisis menunjukkan bahwa klasifikasi data opini film berbahasa Indonesia berdasarkan sentimen menggunakan algoritma Naive Bayes Classifier dapat dilakukan dengan efektif, terutama dengan pemisahan dataset menggunakan metode *5-fold cross-validation*. Ditemukan bahwa semakin banyak data latih yang digunakan, kinerja sistem juga semakin meningkat. Hal ini berpengaruh pada tingginya akurasi hasil yang menunjukkan bahwa sistem berhasil melakukan klasifikasi. Akurasi tertinggi dicapai pada fold kedua dengan nilai 90%, ditunjang oleh precision sebesar 92%, recall 90%, dan f-measure 90%.[4]

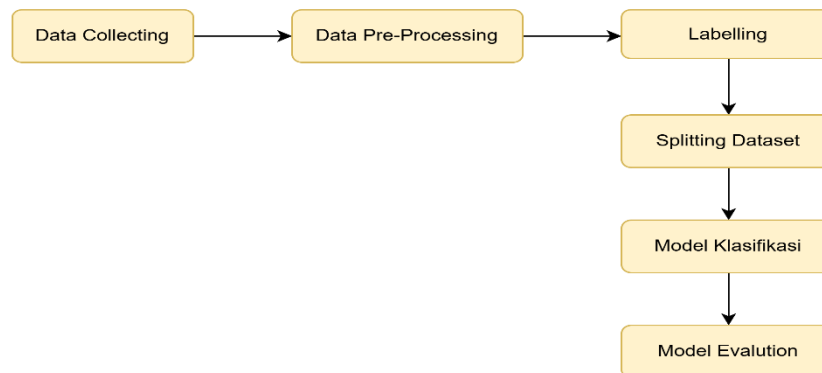
Penelitian yang dilakukan oleh Kaswili Sriwenda Putri, Iwan Rizal Setiawan, Agung Pambudi mengenai analisis sentiment terhadap brand skincare local menggunakan naive bayes classifier dalam proses klasifikasi yang menggunakan naive bayes classifier ini menghasilkan nilai akurasi sebesar 79% untuk dataset Avoskin, 78% untuk dataset Azarine dan 75% untuk dataset Somethinc. Sedangkan pengujian dengan k-fold cross validation menghasilkan nilai sebesar 79% untuk dataset Avoskin serta Somethinc dan 78% untuk dataset Azarine.[5]

Penelitian yang dilakukan oleh Ernianti Hasibuan dan Elmo Allistair Heriyanto menganalisis sentimen pada paper “Sentimen analisis Pada Media Sosial Twitter Menggunakan Metode Naive Bayes Classifier (NBC)” membahas tentang analisis sentimen ulasan pengguna aplikasi Amazon Shopping menggunakan algoritma Naive Bayes. Penelitian ini menunjukkan efektivitas metode Naive Bayes dalam menganalisis sentimen, dengan akurasi mencapai 82,15%, dan hasil Multinomial NB terbaik mencapai 86,74%[6]

Tujuan dari penelitian ini adalah mengetahui hasil analisis sentimen terhadap masyarakat mengenai coding anak dan mengetahui hasil akurasi dari klasifikasi analisis sentimen. Selain itu, penelitian ini juga akan menganalisis sentimen yang diungkapkan di platform Twitter untuk mendapatkan perspektif yang lebih luas mengenai persepsi masyarakat terhadap coding anak.

2. METODE PENELITIAN

Penulis dalam penelitian ini memberikan penjelasan mengenai tahapan-tahapan yang sistematis dalam membangun model klasifikasi, dimulai dari pengumpulan data hingga evaluasi kinerja model. Tahapan penelitian ini dapat dilihat sebagai berikut:



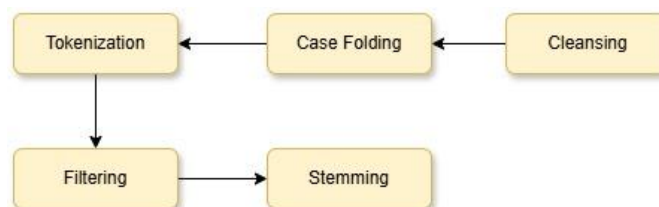
Gambar 1. Tahapan Penelitian

2.1. Data Collecting

Data *Collecting* adalah tahap pertama pada penelitian ini adalah pengumpulan data tweet dengan topik coding anak melalui twitter menggunakan bahasa pemrograman python dengan hastag "#codinganakSD" sebanyak 4357 tweet yang berhasil dikumpulkan. Setelah pengumpulan, data disimpan dalam format CSV untuk keperluan analisis lebih lanjut.

2.2. Data Pre-Processing

Data yang didapat dari hasil crawling belum bisa langsung diklasifikasikan karena data tersebut masih terdapat banyak simbol dan kata-kata yang tidak diperlukan, karena itu kita memerlukan pre-processing data agar data lebih terstruktur dan bersih sehingga bisa diklasifikasikan. Ada beberapa tahapan yang dilakukan preprocessing. Tahapan dapat dilihat pada Gambar 2.



Gambar 2. Tahapan Pada Data Pre-Processing

Berbagai teknik preprocessing yang diterapkan dalam penelitian ini meliputi:

1. Data Cleaning

Data *Cleaning* adalah proses yang dilakukan untuk menghilangkan gangguan (noise) pada data, sehingga kualitas data dapat ditingkatkan [7]. Dengan membersihkan data dari elemen yang tidak relevan, hasil penelitian menjadi lebih akurat dan valid. Proses ini penting untuk memastikan bahwa analisis yang dilakukan mencerminkan kondisi sebenarnya, sehingga pengambilan keputusan yang didasarkan pada hasil penelitian menjadi lebih tepat.

Tabel 1. Hasil *Cleaning*

Tweet	<i>Cleaning</i>
@toe_giman Dia bisa berbicara tentang Coding didepan anak SD	Dia bisa berbicara tentang Coding didepan anak SD

2. Case Folding

Tahap dalam pengolahan teks yang bertujuan untuk mengkonversi seluruh teks dalam dokumen menjadi bentuk standar yang konsisten, biasanya dengan mengubah semua huruf menjadi huruf kecil. Proses ini penting untuk mengurangi variasi yang tidak perlu dalam data, sehingga memudahkan analisis dan pencocokan kata.[8]

Tabel 2. Hasil *Case Folding*

Tweet	Case Folding
@toe_giman Dia bisa berbicara tentang Coding didepan anak SD	Dia bisa berbicara tentang Coding didepan anak SD

3. Tokenization

Tokenization adalah proses memisahkan teks menjadi unit-unit kecil, seperti kata atau frasa, untuk memudahkan analisis. Proses ini memungkinkan penghitungan frekuensi kemunculan setiap kata dan mempermudah pemrosesan data dalam analisis bahasa.

Tabel 3. Hasil *Tokenization*

Tweet	Tokenization
@toe_giman Dia bisa berbicara tentang Coding didepan anak SD	[Dia, bisa, berbicara, tentang, Coding, didepan, anak SD]

4. Filtering (Stopword Removal)

Filtering adalah proses menghapus kata-kata yang dinilai kurang berpengaruh untuk hasil klasifikasi [9]. Proses ini penting untuk meningkatkan kualitas analisis dengan mengurangi kebisingan dalam data, sehingga model klasifikasi dapat fokus pada kata-kata yang lebih signifikan dan relevan. Dengan menghilangkan stopwords, hasil klasifikasi menjadi lebih akurat dan efisien.

Tabel 4. Hasil *Filtering*

Tweet	Filtering
@toe_giman Dia bisa berbicara tentang Coding didepan anak SD	[berbicara, coding, didepan, anak, sd]

5. Stemming

Stemming merupakan proses yang dilakukan untuk menghilangkan kata imbuhan, kata depan, kata ganti, dan kata belakang yang sesuai dengan KBBI [10]. Proses ini bertujuan untuk menghilangkan imbuhan dan variasi infleksi dari kata, sehingga kata-kata yang memiliki makna serupa dapat dikenali sebagai bentuk yang sama.

Tabel 5. Hasil *Stemming*

Tweet	Stemming
@toe_giman Dia bisa berbicara tentang Coding didepan anak SD	[bicara, coding, depan, anak, sd]

2.3. Labeling

Proses selanjutnya adalah memberi label atau data menetapkan polaritas ke dataset yang telah dibersihkan. Labelling atau pelabelan adalah tahap yang penting dalam analisis sentimen yang mencakup pemberian label atau kategori tertentu pada data ulasan untuk mengindikasikan sentimen yang terkandung dalam ulasan tersebut. *Labelling* atau pelabelan dilakukan untuk mengelompokkan setiap ulasan pengguna ke dalam beberapa kategori sentimen positif dan sentimen negatif [11].

Dari dataset yang sudah dilakukan penghapusan duplikat data twitter untuk klasifikasi class sentimen Positif data twitter untuk klasifikasi class sentimen Negatif. Persentase dari kedua kategori tersebut dapat dilihat pada tabel 6.

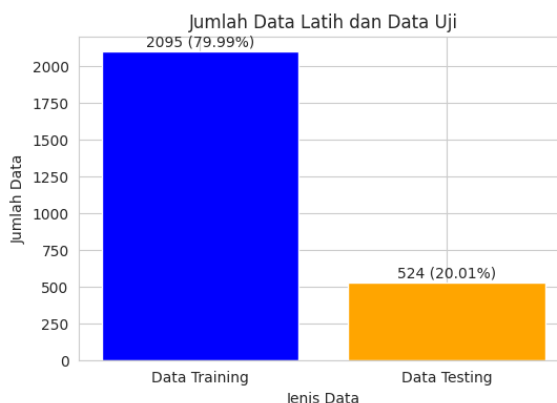
Tabel 6. Hasil *Labelling*

Tweet	Labelling
siswa sd ajar coding anak matematika	Positif
salah sebab sih anak ikut les coding makan larut malam akibat tidur novita kurang tidur sebab turun halusinasi	Negatif

2.4. Splitting Dataset

Setelah data melalui proses pelabelan, langkah selanjutnya adalah membagi data menjadi dua subset utama: data latih (training data) dan data uji (testing data) yang digunakan untuk melatih dan mengevaluasi model. Rasio

yang paling umum digunakan adalah 80% data latih dan 20% data uji [12]. Pertama, dalam skenario di mana jumlah data latih lebih kecil daripada data uji, model dilatih dengan data yang terbatas untuk menguji kemampuannya dalam mengenali pola dari data yang lebih besar. Kedua, dalam skenario di mana jumlah data latih lebih banyak daripada data uji, model dilatih dengan lebih banyak informasi, yang dapat meningkatkan akurasi dan keandalan prediksi.



Gambar 3. Grafik Dataset

Gambar 3 menunjukkan grafik menggambarkan pembagian dataset menjadi data latih (training) dan data uji (testing), dengan 2095 data (79,99%) digunakan untuk melatih model dan 524 data (20,01%) digunakan untuk menguji kinerja model. Warna biru mewakili data latih, sedangkan warna oranye mewakili data uji. Proporsi ini memastikan model memiliki data yang cukup untuk pembelajaran sekaligus data independen untuk evaluasi, sehingga hasil evaluasi dapat merepresentasikan kemampuan model pada data baru.

2.5. Model Klasifikasi

Dalam penelitian ini metode klasifikasi yang digunakan oleh penulis yaitu metode klasifikasi Naive Bayes.

1. Naive Bayes Klasifikasi (NBC)

Naive Bayes Classifier adalah metode untuk mengklasifikasikan data menerapkan probabilitas sederhana dengan menggunakan teorema Bayes dan mengasumsikan tingkat independensi yang tinggi. Metode ini sangat sesuai dalam menangani data dengan kinerja yang cepat untuk klasifikasi data dengan akurasi yang tinggi [13]. Metode ini sudah biasa diterapkan untuk pengklasifikasian data, seperti analisis sentimen, pengenalan teks, dan spam filtering. Ada beberapa varian dari algoritma Naive Bayes, antara lain:

2. Naive Bayes Gaussian

Gaussian Naïve Bayes (GNB) merupakan algoritma klasifikasi yang menggunakan teori Bayes dengan cara kerja klasifikasi pada label kelas yang sebelumnya telah ditentukan berdasarkan aturan yang telah ditetapkan dari variable prediktor yang bersifat independent. [14] Dalam Naive Bayes Gaussian, setiap kelas diasumsikan memiliki distribusi normal yang berbeda, dan model ini menghitung probabilitas setiap kelas berdasarkan nilai rata-rata dan deviasi standar dari fitur-fitur tersebut.

3. Naive Bayes Multinomial

Klasifikasi multinomial Naïve Bayes adalah pengembangan dari algoritma Bayes yang sering diterapkan dalam klasifikasi teks. Model ini mempertimbangkan frekuensi kemunculan setiap kata dalam dokumen, dan analisis frekuensi tersebut dapat mempermudah proses klasifikasi [15].

4. Naive Bayes Bernoulli

Algoritma Naive Bayes Bernoulli menerapkan proses klasifikasi pada data yang mengikuti distribusi Bernoulli multivariat; artinya, fitur-fitur yang ada mungkin memiliki beberapa variasi, tetapi setiap fitur dianggap sebagai variabel biner (Bernoulli, boolean) [16]. Oleh karena itu, kelas ini memerlukan sampel yang diwakili dalam bentuk vektor fitur dengan nilai-nilai biner.

2.6. Evaluasi Model

Evaluasi model dilakukan setelah proses pengujian model selesai. Tujuan dari evaluasi dan perbandingan algoritma adalah membuktikan bahwa model benar-benar mempresentasikan sesuai pemodelan yang dilakukan dan mengukur performa dari algoritma dengan menentukan tingkat keakuratan dari dua

algoritma yang digunakan. Penulis menggunakan hasil pengujian *confusion matrix*, yang mencakup metrik seperti Akurasi, *Precision*, *Recall*, dan *F1-Score*. Gambar *confusion Matrix* pada tabel 7. sebagai berikut:

$$Acurasy = \frac{TP+TN}{TP+TN+FP+PN} \times 100 \quad (1)$$

$$Precision = \frac{TP}{TP+FP} \quad (2)$$

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

$$F1 - Score = 2 \times \frac{Presisi \times Recall}{Presisi + Recall} \quad (4)$$

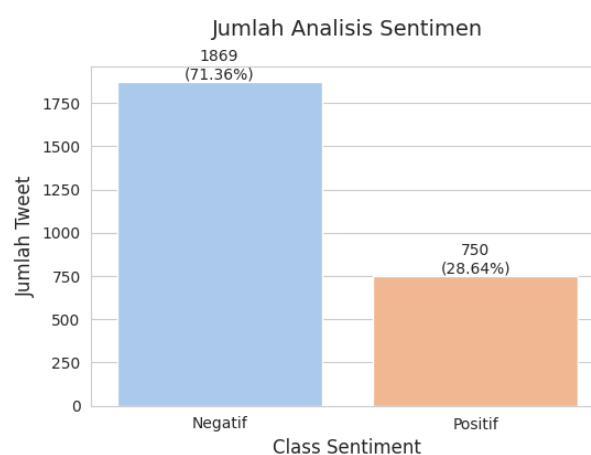
Tabel 7. Confusion Matrix

Class	Classified as positive	Classified as negative
Positif	True Positive (TP)	False Negative (FN)
Negatif	False Positive (FP)	True Negative (TN)

Dalam mengevaluasi model klasifikasi, terdapat empat metrik utama yang digunakan untuk mengukur kinerjanya, yaitu *True Positive* (TP), *False Positive* (FP), *True Negative* (TN), dan *False Negative* (FN). Metrik-metrik ini sangat penting dalam menilai akurasi model. *True Positive* merujuk pada jumlah data positif dalam set yang berhasil diklasifikasikan sebagai positif. Sementara itu, *True Negative* adalah jumlah data negatif yang tepat diklasifikasikan sebagai negatif. Sebaliknya, *False Positive* adalah jumlah data negatif yang salah diklasifikasikan sebagai positif, dan *False Negative* adalah data positif yang salah diklasifikasikan sebagai negatif. Dari matriks ini, kita dapat menghitung berbagai metrik lainnya seperti *acurasy*, *precision*, *recall*, dan *F1-score* untuk mendapatkan gambaran yang lebih komprehensif mengenai kinerja model tersebut..

3. HASIL DAN PEMBAHASAN

3.1. Dataset



Gambar 4. Jumlah Klasifikasi

Data yang digunakan dalam penelitian ini terdiri dari tweet yang telah dianalisis, dengan total mencapai 4357 tweet. Dari jumlah tersebut, 1869 tweet dikategorikan sebagai negatif, sementara 750 tweet dikategorikan sebagai positif. Pengelompokan ini sangat penting untuk memahami bagaimana masyarakat merespons topik yang diteliti. Dengan menganalisis sentimen ini, penelitian dapat memberikan wawasan yang lebih jelas mengenai pandangan dan persepsi publik terhadap isu yang sedang dibahas. Dapat dilihat Gambar 4 jumlah klasifikasinya:

Jumlah tweet negatif mencapai 1.869, jauh lebih tinggi dibandingkan tweet positif yang hanya 750. Hal ini dapat disebabkan oleh ekspresi sentimen negatif yang lebih dominan di media sosial, topik yang memicu pandangan negatif, serta ketidakseimbangan data dalam analisis. Sentimen negatif terhadap program coding anak mungkin muncul karena metode pembelajaran yang dianggap terlalu kompleks, kurangnya fasilitas pendukung, atau ketidaksesuaian metode dengan kebutuhan anak. Untuk mengatasinya, diperlukan pendekatan pembelajaran yang interaktif dan sesuai usia, pelatihan bagi guru dan orang tua, serta penyediaan alat bantu yang mendukung proses belajar.

3.2. Tahap Visualisasi Wordcloud

Word cloud digunakan untuk mengidentifikasi kata-kata yang paling sering muncul dalam sebuah teks. Dengan menampilkan kata-kata dalam berbagai ukuran sesuai dengan frekuensi kemunculannya, word cloud memudahkan kita untuk memvisualisasikan tema dan sentimen yang terkandung dalam data teks. Hasil visualisasi Word Cloud ditunjukkan pada Gambar 5. Kata-kata yang berukuran lebih besar menunjukkan frekuensi kemunculan yang lebih tinggi, seperti "Coding Anak".



Gambar 5. Visualisasi Word Cloud

3.3. Tahap Pengujian

Dalam pengujian ini data tweet dibagi menjadi data latih dan data uji dengan perbandingan 80% untuk data latih dan 20% untuk data uji, yang telah diberi label sebelumnya. Evaluasi dilakukan terhadap tiga model *Naïve Bayes* yaitu *GaussianNB*, *MultinomialNB*, dan *BernoulliNB*. Metrik yang digunakan meliputi akurasi, presisi, *recall* dan *F1-score*, yang masing-masing memberikan gambaran tentang kemampuan model dalam mengklasifikasi data secara keseluruhan. Berikut rumus akurasi. Metrik-metrik ini memberikan analisis menyeluruh mengenai efektivitas masing-masing model dalam mengklasifikasikan data secara akurat dan efisien.

1. Nilai accuracy, precision, recall dan F1-Score

Tabel 8. menampilkan tingkat akurasi Naive Bayes dapat dilihat dibawah ini:

Tabel 8. Hasil Nilai Accuracy

Metode	Accuracy Score
Naive Bayes Gaussian	0.58
Naive Bayes Multinomial	0.72
Naive Bayes Bernoulli	0.72

Tabel 8 menunjukkan hasil akurasi dari tiga varian metode Naive Bayes, yaitu Gaussian, Multinomial, dan Bernoulli, yang digunakan untuk melakukan klasifikasi. Metode Naive Bayes Gaussian memiliki akurasi sebesar 58%, menunjukkan performa yang paling rendah dibandingkan dua metode lainnya. Sementara itu, metode Naive Bayes Multinomial dan Naive Bayes Bernoulli masing-masing mencapai akurasi 72%, yang menunjukkan bahwa keduanya memiliki kinerja yang lebih baik dalam memprediksi data dibandingkan Gaussian. Perbedaan ini menunjukkan bahwa pilihan varian Naive Bayes dapat memengaruhi tingkat akurasi berdasarkan jenis data yang digunakan. Dalam studi ini, selain mengukur akurasi, juga digunakan F1-score, presisi, dan recall. Hasil klasifikasi dilihat pada Tabel 9.

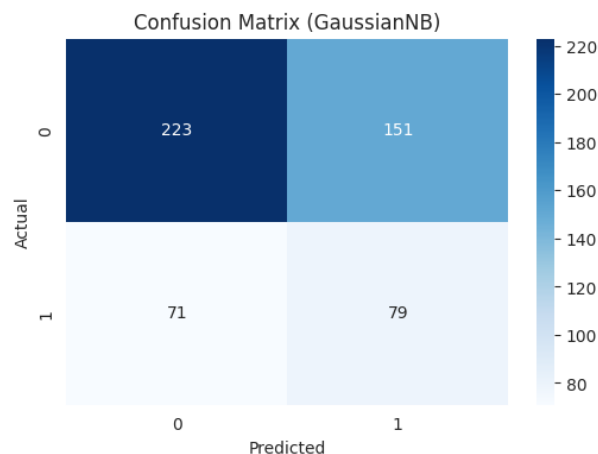
Tabel 9. Hasil Klasifikasi

Modelling	Sentimen	Precision	Recall	F1-score
Naive Bayes Gaussian	Positif	0.34	0.60	0.67
	Negatif	0.76	0.53	0.42
Naive Bayes Multinomial	Positif	0.71	0.03	0.06
	Negatif	0.72	0.99	0.84
Naive Bayes Bernoulli	Positif	0.68	0.21	0.32
	Negatif	0.75	0.96	0.84

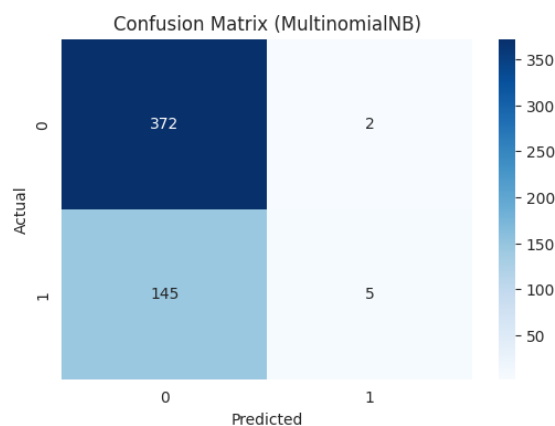
Tabel 9 menyajikan hasil evaluasi kinerja tiga jenis metode Naive Bayes, yaitu Gaussian, Multinomial, dan Bernoulli, yang berfokus pada analisis sentimen positif dan negative. Dari ketiga metode Naive Bayes, MultinomialNB memiliki akurasi terbaik dalam mengenali sentimen negatif, dengan precision 72% dan recall 99%, yang berarti hampir semua tweet negatif terdeteksi dengan benar. Namun, metode ini kurang efektif dalam mengenali sentimen positif, dengan recall hanya 3%, sehingga banyak opini positif yang terabaikan. BernoulliNB memiliki pola serupa dengan recall 96% untuk sentimen negatif tetapi hanya 21% untuk sentimen positif, menunjukkan bias terhadap opini negatif. Sementara itu, GaussianNB memiliki precision lebih rendah untuk sentimen positif (34%) dibandingkan sentimen negatif (76%), dengan recall yang kurang stabil, yaitu 60% untuk positif dan 53% untuk negatif. Dari hasil ini, MultinomialNB tetap menjadi pilihan terbaik untuk analisis sentimen, terutama dalam mendeteksi opini negatif secara real-time

3.4. Hasil Confusion Matrix

Selain itu, penelitian ini juga membandingkan matriks konfusi ketiga algoritma. Tujuan dari perbandingan ini adalah untuk mempertimbangkan indikator-indikator seperti true positive (TP), false positive, true negative (TN), dan false negative (FN) untuk mengidentifikasi emosi positif dan negatif setiap algoritma.



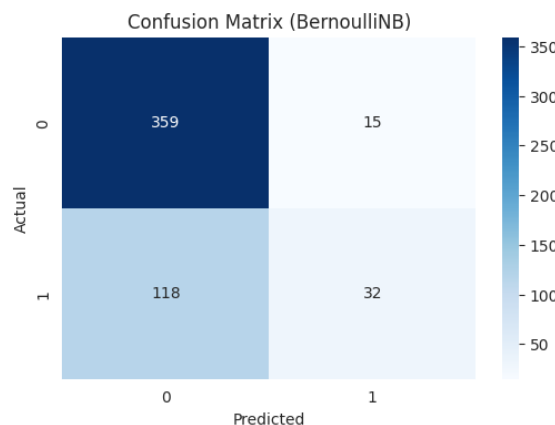
Gambar 6. Confusion Matrix GaussianNB



Gambar 7. Confusion Matrix MultinomialNB

Dari gambar 6 confusion matrix Gaussian Naive Bayes, dapat dilihat bahwa model menunjukkan performa yang lebih baik dalam mengidentifikasi data sentiment negatif (223 TN). Namun, model kurang optimal dalam mengidentifikasi data sentiment positif, (79 TP). Kesalahan juga terjadi pada prediksi, di mana sejumlah senimen negatif salah diidentifikasi sebagai positif (151 FP), dan beberapa data positif salah diidentifikasi sebagai negative (71 FN). Hasil ini menunjukkan bahwa meskipun model cukup baik dalam mengenali kelas negatif, masih diperlukan perbaikan untuk meningkatkan akurasi dalam mengenali kelas positif.

Gambar 7 confusion matrix Multinomial Naive Bayes menunjukkan bahwa model memiliki kinerja yang sangat baik dalam mengenali data negatif, dengan 372 data negatif diprediksi dengan benar (True Negative) dan hanya 2 data negatif salah diprediksi sebagai positif (False Positive). Namun, model kurang mampu mengenali data positif, di mana hanya 5 data positif diprediksi dengan benar (True Positive), sementara 145 data positif salah diprediksi sebagai negatif (False Negative). Hal ini menunjukkan bahwa model lebih cenderung fokus pada kelas negatif dan memerlukan perbaikan untuk meningkatkan akurasi dalam mengenali kelas positif.



Gambar 8. Confusion Matrix BernoulliNB

Gambar 8, matriks kebingungan ini menggambarkan bahwa model Bernoulli Naive Bayes lebih unggul dalam mengidentifikasi kelas negatif (359 True Negatif) dibandingkan kelas positif (32 True Positif), terdapat 15 prediksi positif yang salah (False Positive). Namun, model sering salah memprediksi kelas positif sebagai negatif (118 kasus False Negative), yang menunjukkan potensi bias terhadap kelas negatif. Untuk meningkatkan akurasi model pada kelas positif, langkah-langkah seperti menyeimbangkan data atau menggunakan algoritma lain dapat dipertimbangkan.

4. KESIMPULAN

Penelitian ini membahas isu penting mengenai pemanfaatan coding anak di Indonesia dan dampaknya terhadap perkembangan mereka dengan menganalisis sentimen dari 4.357 tweet. Metode Naive Bayes, termasuk Naive Bayes Gaussian 72%, Naive Bayes Multinomial 72%, dan Naive Bayes Bernoulli menunjukkan akurasi 58%. Dari total tweet yang dianalisis, 1.869 tweet dikategorikan sebagai negatif, sementara 750 tweet dikategorikan sebagai positif serta memanfaatkan bahasa pemrograman Python di Google Collab. menyajikan hasil evaluasi kinerja tiga jenis metode Naive Bayes, yaitu Gaussian, Multinomial, dan Bernoulli, yang berfokus pada analisis sentimen positif dan negative. Dari ketiga metode Naive Bayes, MultinomialNB memiliki akurasi terbaik dalam mengenali sentimen negatif, dengan precision 72% dan recall 99%, yang berarti hampir semua tweet negatif terdeteksi dengan benar. Namun, metode ini kurang efektif dalam mengenali sentimen positif, dengan recall hanya 3%, sehingga banyak opini positif yang terabaikan. BernoulliNB memiliki pola serupa dengan recall 96% untuk sentimen negatif tetapi hanya 21% untuk sentimen positif, menunjukkan bias terhadap opini negatif. Sementara itu, GaussianNB memiliki precision lebih rendah untuk sentimen positif (34%) dibandingkan sentimen negatif (76%), dengan recall yang kurang stabil, yaitu 60% untuk positif dan 53% untuk negatif. Hasil penelitian ini dapat menjadi dasar untuk mengembangkan kebijakan yang lebih baik terkait penggunaan teknologi di kalangan anak-anak. Selain itu, penelitian ini juga membuka peluang untuk studi lebih lanjut yang dapat menyelidiki dampak jangka panjang dari penggunaan metode analisis dan teknologi digital lainnya.

DAFTAR PUSTAKA

- [1] P. Silvia, "Analisis Kemampuan Computational Thinking Melalui Pembelajaran Coding Pada Anak Usia Dini 0-8 Tahun," *J. Islam. Early Child. Educ. PIAUD-Ku*, vol. 1, no. 2, pp. 50–59, 2022, doi: 10.54801/piaudku.v1i2.140.
- [2] N. T. Romadloni and W. Supriyanti, "Analisis Sentimen Penggunaan Teknologi Pada Pendidikan Anak Usia Dini," *J. Ilm. SINUS*, vol. 21, no. 2, p. 101, 2023, doi: 10.30646/sinus.v21i2.759.
- [3] F. V. Sari and A. Wibowo, "Analisis Sentimen Pelanggan Toko Online Jd.Id Menggunakan Metode Naïve Bayes Classifier Berbasis Konversi Ikon Emosi," *J. SIMETRIS*, vol. 10, no. 2, pp. 681–686, 2019.
- [4] F. Ratnawati, "Implementasi Algoritma Naive Bayes Terhadap Analisis Sentimen Opini Film Pada Twitter," *INOVTEK Polbeng - Seri Inform.*, vol. 3, no. 1, p. 50, 2018, doi: 10.35314/isi.v3i1.335.
- [5] K. S. Putri, I. R. Setiawan, and A. Pambudi, "Analisis Sentimen Terhadap Brand Skincare Lokal Menggunakan Naïve Bayes Classifier," *Technol. J. Ilm.*, vol. 14, no. 3, p. 227, 2023, doi: 10.31602/tji.v14i3.11259.
- [6] Ernianti Hasibuan and Elmo Allistair Heriyanto, "Analisis Sentimen Pada Ulasan Aplikasi Amazon Shopping Di Google Play Store Menggunakan Naive Bayes Classifier," *J. Tek. dan Sci.*, vol. 1, no. 3, pp. 13–24, 2022, doi: 10.56127/jts.v1i3.434.
- [7] R. Nooraeni, A. B. Safiruddin, A. F. Afifah, K. D. Agung, and N. N. Rosyad, "Analisis Sentimen Publik terhadap Sistem Zonasi Sekolah Menggunakan Data Twitter dengan Metode Naïve Bayes Classification," *Fakt. Exacta*, vol. 12, no. 4, p. 315, 2020, doi: 10.30998/faktorexacta.v12i4.5205.
- [8] S. Berliani and S. Lestari, "Analisis Sentimen Masyarakat Terhadap Isu Pecat Sri Mulyani Pada Twitter Menggunakan Metode Naive Bayes Dan Support Vector Machine," *J. Sains dan Teknol.*, vol. 5, no. 3, pp. 951–960, 2024, doi: 10.55338/saintek.v5i3.2746.
- [9] S. N. Salsabila, B. N. Sari, and R. Mayasari, "Klasifikasi Ulasan Pengguna Aplikasi Discord Menggunakan Metode Information Gain Dan Naïve Bayes Classifier," *INFOTECH J.*, vol. 9, no. 2, pp. 383–392, 2023, doi: 10.31949/infotech.v9i2.6277.
- [10] N. Agustina, D. H. Citra, W. Purnama, C. Nisa, and A. R. Kurnia, "Implementasi Algoritma Naive Bayes untuk Analisis Sentimen Ulasan Shopee pada Google Play Store," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 2, no. 1, pp. 47–54, 2022, doi: 10.57152/malcom.v2i1.195.
- [11] N. Cahyono and Anggista Oktavia Praneswara, "Analisis Sentimen Ulasan Aplikasi TikTok Shop Seller Center di Google Playstore Menggunakan Algoritma Naive Bayes," *Indones. J. Comput. Sci.*, vol. 12, no. 6, pp. 3925–3940, 2023, doi: 10.33022/ijcs.v12i6.3473.
- [12] R. Apriani and D. Gustian, "Analisis Sentimen Dengan Naïve Bayes Terhadap Komentar Aplikasi Tokopedia," *J. Rekayasa Teknol. Nusa Putra*, vol. 6, no. 1, pp. 54–62, 2019, doi: 10.52005/rekayasa.v6i1.86.
- [13] F. Meila, A. Sofyan, N. Sulistiyowati, and A. Voutama, "ANALISIS SENTIMEN TERHADAP RESPON PERUBAHAN NAMA TWITTER MENJADI ' X ' MENGGUNAKAN METODE NAÏVE BAYES CLASSIFIER," vol. 8, no. 5, pp. 10987–10994, 2024.
- [14] S. Pokhrel, "No TitleEAENH," *Ayan*, vol. 15, no. 1, pp. 37–48, 2024.
- [15] L. M. Siniwi, A. Prahutama, and A. R. Hakim, "QUERY EXPANSION RANKING PADA ANALISIS SENTIMEN MENGGUNAKAN KLASIFIKASI MULTINOMIAL NAÏVE BAYES (Studi Kasus : Ulasan Aplikasi Shopee pada Hari Belanja Online Nasional 2020)," *J. Gaussian*, vol. 10, no. 3, pp. 377–387, 2021, doi: 10.14710/j.gauss.v10i3.32795.
- [16] V. Oktaviana Yamin, A. Tenriawaru, L. Ode Saidi, and G. Arviana Rahman, "Penerapan Naïve Bayes Classifier dengan Algoritma Nazief dan Adriani Untuk Deteksi Hoaks," *Pros. Semin. Nas. Pemanfaat. Sains dan Teknol. Inf.*, vol. 1, no. 1, pp. 335–344, 2023.