

ASPECT-BASED SENTIMENT ANALYSIS IPHONE 14 PRO MENGUNAKAN ALGORITMA XGBOOST

Galih Arum Prabowo, Basuki Rahmat, Henni Endah Wahanani

Program Studi Informatika S1, Fakultas Ilmu Komputer
Universitas Pembangunan Nasional Veteran Jawa Timur,
Jl Raya Rungkut Madya Gunung Anyar Surabaya, Indonesia
galiharum181100@gmail.com

ABSTRAK

Penelitian ini akan mengimplementasikan algoritma XGBoost untuk melakukan analisis sentimen berbasis aspek pada kolom komentar video youtube review iPhone 14 Pro, yang bertujuan untuk mengukur tingkat performansi algoritma XGBoost dalam mengklasifikasi ulasan ke dalam dua kelas, yaitu kelas sentimen positif-negatif, sekaligus mengetahui kelebihan dan kelemahan Iphone iPhone 14 pro. Dari hasil pengujian didapat tingkat akurasi yang cukup tinggi, yaitu berkisar antara 89% sampai 97%. Selain itu, dari penelitian ini juga dapat diketahui bahwa kelebihan yang dimiliki terdapat pada aspek citra merek dan performa. Sedangkan untuk kelemahannya, iPhone 14 pro masih banyak mendapat sentimen negatif terutama pada aspek desain dan spesifikasi

Kata kunci: Analisis Sentimen, Aspect-Based, iPhone 14 pro, XGBoost

1. PENDAHULUAN

Di masa sekarang teknologi membawa banyak kemudahan, salah satu kemudahan dalam mengakses informasi melalui internet. Dilansir dari hasil survei yang dilakukan oleh hootsuite (We Are Social) mengenai Indonesia Digital Report, terdapat sebanyak 80,1% menggunakan internet untuk menemukan informasi. Hal tersebut didukung dengan banyak platform penyedia layanan informasi seperti Youtube, Twitter, Instagram dan banyak lagi. Dengan kepopuleran Youtube terbukti dari hasil survei yang dilakukan oleh hootsuite (We Are Social) mengenai Indonesia Digital Report, Youtube menempati posisi kedua setelah google dalam kategori website yang banyak dikunjungi orang Indonesia pada tahun 2022 dengan total 241 miliar orang[1]. Hal ini dimanfaatkan banyak pihak sebagai wadah untuk branding dan penyebaran informasi, termasuk penyebaran informasi mengenai barang elektronik, salah satu channel yang cukup menarik pengguna Youtube khususnya para penggemar gadget yaitu GadgetIn

GadgetIn membahas seputar barang elektronik. Salah satu videonya yang membahas iPhone 14 Pro. Di dalam video menjelaskan tentang Iphone 14 Pro Indonesia yang sedang marak di kalangan masyarakat indonesia karena inovasi baru pada iphone 14. Banyak masyarakat indonesia yang bertanya mengenai performa teknologi baru dari brand Iphone ini. Dibutuhkan informasi untuk melihat apakah masyarakat tertarik pada gadget baru ini atau tidak. Maka dari itu untuk mengatasi masalah ini, dilakukan penelitian analisis sentimen melalui komentar pengguna youtube pada channel GadgetIn untuk melihat minat beli masyarakat sebagai konsumen terhadap produk Iphone 14 pro. Analisis sentimen secara singkat dapat diartikan sebagai studi komputasi yang memanfaatkan teks berisi opini,

pandangan, pendapat, penilaian, serta emosi seseorang terhadap suatu objek [2].

Sentimen pada ulasan-ulasan tersebut nantinya akan diklasifikasikan ke dalam dua kategori, yaitu sentimen positif dan sentimen negatif. Untuk algoritma klasifikasinya, peneliti akan fokus pada algoritma XGBoost sebagai algoritma klasifikasi dan metode TF-IDF (Term Frequency - Inverse Document Frequency) sebagai metode word embedding. Secara singkat word embedding adalah bentuk representasi numerik dari sebuah frasa [3]. Sedangkan algoritma Extreme Gradient Boosting (XGBoost) adalah salah satu algoritma berkinerja terbaik yang digunakan untuk supervised learning. Dapat digunakan untuk masalah regresi dan klasifikasi. Xgboost disukai oleh ilmuwan karena kecepatan eksekusinya yang tinggi[4]. Hal tersebut dibuktikan dengan metode XGBoost menjadi metode yang banyak digunakan pada kompetisi machine learning[5].

Penelitian ini bertujuan untuk menghasilkan sebuah sistem yang dapat membantu para calon user untuk mengetahui kelebihan dan kelemahan dari iPhone 14 Pro selain itu hasil dari penelitian ini diharapkan dapat menjadi sumber rujukan pihak iPhone memberikan model algoritma yang cukup akurat dalam menganalisis sentimen publik di Youtube mengenai produk iPhone 14 pro. Oleh karena itu, pada penelitian ini akan menggunakan TF-IDF sebagai metode word embedding. Lalu algoritma klasifikasi XGBoost

2. TINJAUAN PUSTAKA

Teori dasar yang akan dipakai untuk penelitian diperlukan untuk mendukung penelitian ini. Landasan teori yang akan digunakan adalah Analisis Sentimen, Preprocessing, TF-IDF, XGBoost, Confusion Matrix.

2.1. Analisis Sentimen

Analisis sentimen adalah teknik pengolahan bahasa alami yang digunakan untuk mengidentifikasi, mengekstrak, dan mengevaluasi opini dan sentimen yang terkandung dalam teks dalam rangka menentukan apakah opini tersebut bersifat positif, negatif, atau netral. Analisis sentimen juga biasa disebut sebagai opinion mining. Analisis sentimen atau opinion mining bertujuan untuk memahami kecenderungan pendapat publik terhadap suatu objek permasalahan, apakah cenderung ke arah opini positif atau ke arah opini negatif. Analisis sentimen ini mulai populer sebagai topik penelitian sejak tahun 2002[6].

2.2. Aspect-Based Sentiment Analysis (ABSA)

Aspect-based sentiment analysis (ABSA) adalah sebuah pendekatan yang akan mengambil aspek dan sentimen pada sebuah kalimat sehingga setelahnya dapat diketahui sentimen-sentimen dari masing-masing aspek yang ada. Contoh dari aspect-based sentiment analysis bisa diambil dari sebuah kalimat ulasan produk, “Ni hape jujur bagus, dengan layar yg memanjakan mata, performa yang tinggi, apalagi fitur dynamic island nya cukup dieksekusi dengan mantap oleh apple, yang biasanya tompel ditutup tutupi di hp android, di iphone malah dibuat fitur wkwkwkw, tetep aja inovasi apple sekarang engga se gahar dulu”. Dari kalimat tersebut, aspek ulasan yang dibahas ada tiga, yakni citra merek, spesifikasi, dan desain. Untuk aspek citra merek termasuk sentimen negatif karena dianggap tidak segahar dulu, sedangkan untuk aspek spesifikasi dan desain termasuk ke dalam sentimen positif karena dapat merubah kekurangan di merek lain menjadi kelebihan dan terobosan baru.

2.3. Preprocessing

Teks atau kalimat yang diambil dari sumber bisa disebut sebagai data mentah. Hal ini dikarenakan dataset tersebut merupakan data yang tidak terstruktur sehingga komputer akan sulit untuk mengolah data tersebut. Karena hal tersebut, diperlukanlah tahap text preprocessing sebelum nantinya dilakukan penentuan fitur, atau dalam penelitian ini dilakukan dengan pembobotan setiap kata dengan menggunakan metode TF-IDF[7]. Pada penelitian ini, terdapat 5 tahap text preprocessing yaitu case-folding, cleaning, tokenizing, stopword removal, dan stemming.

2.4. TF-IDF

Term Frequency-Inverse Document Frequency (TF-IDF) adalah sebuah metode word embedding atau pengubahan kata menjadi nilai numerik, yang mana nilai numerik tersebut dapat mewakili bobot dari setiap kata yang ada di dalam sebuah kumpulan dokumen [8].

2.5. XGBoost

Gradient Boosting adalah algoritma pembelajaran mesin yang pengembangan algoritma akselerasi adaptif. Gradient Boost mampu mencegah

over-tuning membangun pohon keputusan berdasarkan pertumbuhan struktur pohon Belajar itu lemah, begitu juga memperbaiki kesalahan pohon keputusan[10].

2.6. Confusion Matrix.

Confusion matrix merupakan matriks atau tabel yang mampu merepresentasikan performa algoritma dalam proses pengklasifikasian kelas-kelas [11]

3. METODE PENELITIAN

Pada bagian ini memuat metode yang digunakan pada pemuatan skripsi.

3.1. Pengumpulan Data

Tahap pengambilan data merupakan tahapan awal pada perancangan sistem ini, dimana sistem akan melakukan scrapping data sebanyak 1000 buah data yang nantinya akan diproses lebih lanjut untuk diklasifikasi. Dataset di-scrapping dengan menggunakan google developer. Google developer merupakan sebuah package berbasis python yang mampu mengambil berbagai data yang ada Youtube, salah satunya adalah kolom komentar Youtube. Selanjutnya data tersebut dibagi menjadi 4 aspek, yaitu:

1. Citra Merek (Seberapa Iphone mendapatkan perhatian terhadap inovasi dan harga yang ada)
2. Desain (seberapa suka dengan desain iphone 14 terbaru dan desain fitur baru yaitu dynamic island)
3. Spesifikasi (seberapa puas pengguna dengan spesifikasi kamera layar speaker)
4. Performa (seberapa puas user dengan iphone 14 terutama seperti baterai, chipset 16 bionic)

3.2. Preprocessing

Tahap selanjutnya ialah tahap preprocessing. Secara singkat, tahap ini dilakukan untuk memproses data agar siap digunakan pada tahapan-tahapan selanjutnya. Beberapa tahapan yang akan dilakukan pada preprocessing data ini ialah case-folding, cleaning, tokenizing, stopword removal, dan stemming.

3.2.1. Case Folding

Tahap text preprocessing yang pertama ialah case-folding. Case-folding ialah proses yang akan merubah karakter yang ada pada dataset agar menjadi huruf kecil semua.

3.2.2. Cleaning

Tahap text preprocessing yang selanjutnya adalah cleaning. Cleaning adalah proses pembersihan kalimat dari atribut-attribut yang mengganggu, seperti angka, simbol, tanda baca, dan karakter kosong.

3.2.3. Tokenizing

Lalu selanjutnya ialah tahap tokenizing. Tokenizing adalah proses yang akan memotong

kalimat pada dataset agar terpisah-pisah menjadi satuan kata.

3.2.4. Stopword Removal

Stopword Removal adalah suatu tahapan penghapusan kata-kata pada dataset yang tidak relevan, tidak berhubungan, dan tidak berpengaruh pada tahap pemrosesan data selanjutnya. Kata tanya, kata hubung, kata depan, dan kata-kata tidak bermakna lainnya juga akan dibuang pada tahap ini. Contoh kata-kata yang akan dibuang pada tahap stopwords removal ini adalah di, ke, adalah, merupakan, sangat, dan, jika, agak, saja, itu, begitu, banget, saja.

3.2.5. Stemming

Kata-kata yang ada di dataset biasanya bermacam-macam, ada yang hanya berbentuk kata dasar, ada pula kata yang memiliki imbuhan. Untuk memaksimalkan tahapan preprocessing, kata-kata yang memiliki imbuhan, perlu diubah ke bentuk kata dasar. Tahap inilah yang disebut tahap stemming.

3.3. TF-IDF

Setelah di lakukan preprocessing langkah selanjutnya ialah pembobotan setiap kata tadi menggunakan Metode TF-IDF memberikan nilai atau bobot pada setiap kata didasarkan pada frekuensi kemunculan sebuah kata. Namun sebelum metode TF-IDF diterapkan, diperlukan juga proses label encoding untuk mengubah kelas target menjadi bentuk angka. Tahapan-tahapan dalam melakukan perhitungan secara manual dengan menggunakan metode TF-IDF ialah sebagai berikut:

1. Ketahuilah nilai DF dan nilai D (jumlah kalimat) di dalam corpus
2. Hitung nilai TF
3. Hitung nilai IDF
4. Hitung nilai w (bobot) dengan mengalikan nilai TF dan IDF.

$$W = TF \times IDF \quad (1)$$

$$W = f_{t,d} \times \frac{D}{df(t)} \quad (2)$$

3.4. XGBoost

Setelah setiap kata mendapat bobot masing masing menggunakan metode TF-IDF lalu dilakukan klasifikasi menggunakan Algoritma XGBoost dalam melakukan perhitungan secara manual dengan menggunakan Algoritma XGBoost. Adapun contoh proses perhitungan Algoritma XGBoost :

1. Langkah 1

Menentukan Probabilitas

Menentukan nilai probabilitas awal (Pr) dan residual Nilai probabilitas ditentukan lebih awal, karena label terdiri dari 2 kategori maka probabilitas yang digunakan bernilai 0.5 dan contoh perhitungan nilai residual untuk instance 1, sebagai berikut:

$$residual = Y - F_0(x) \quad (5)$$

2. Langkah 2

Menentukan SW

Untuk menentukan simpul 1 atau cabang yang pertama pada pohon ditentukan dengan nilai gain. Setelah didapatkan hasil dari nilai gain maka dilakukan perhitungan similarity weight (SW) atau similarity score

$$SW = \frac{(\sum Residual)^2}{\sum (Pr(1-Pr) + \lambda)} \quad (6)$$

3. Langkah 3

Menentukan Gain

Setelah mendapatkan nilai similarity weight (SW) atau similarity score selanjutnya nilai SW ini akan digunakan untuk menentukan Gain. Dimana nilai gain digunakan untuk menentukan atribut mana yang akan digunakan untuk menjadi simpul sesuai dengan nilai gain tertinggi. Berikut contoh perhitungan untuk mendapatkan nilai Gain.

$$Gain = Sw_{kiri} + Sw_{kanan} - Sw_{awal} \quad (7)$$

3.5. Confusion Matrix

Setelah melakukan klasifikasi menggunakan algoritma XGBoost untuk mendapatkan nilai akurasi maka di lakukan menggunakan perhitungan confusion matrix. Terdapat empat elemen pada confusion matrix yaitu TP (True Positive) yaitu hasil dari klasifikasi yang tepat, TN (True Negative) yaitu hasil dari klasifikasi yang tidak tepat, FP (False Positive) yaitu hasil dari klasifikasi yang tepat tetapi faktanya kurang tepat, FN (False Negative) yaitu hasil dari klasifikasi yang kurang tepat tetapi faktanya tepat. Empat elemen tersebut selanjutnya dapat diolah untuk memperoleh data data lain yang mampu mengukur tingkat performansi sebuah algoritma, yaitu accuracy, precision, recall, dan F1 score.

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

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

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

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

4. METODE PENELITIAN

Adapun hasil dan pembahasan yang akan dijelaskan menjadi beberapa sub bab.

4.1. Pengumpulan Data

Dalam proses pengambilan dataset, memanfaatkan Youtube Data API v3. Sebelum nya penulis sudah mengaktifkan layanan google developer console dimana untuk mendapatkan API key. Setelah mendapatkan API key penulis meng-import beberapa library dasar yang akan membantu proses pengambilan dataset, seperti library pandas untuk mengolah data dan library google api client. Setelah import berhasil penulis membuat fungsi untuk crawling. Dimana

penulis memanggil fungsi dari library google api client untuk mengakses/mengambil data komentar.

	publishedAt	authorDisplayName	textDisplay	likeCount
0	2023-03-01T06:05:49Z	Ramzi Channel	Sampai saat ini masih bertahan di 11 Pro Max u...	0
1	2023-02-28T08:43:16Z	Reno Official	Kasih, iPhone tidak ada kamera depan. Selalu...	0
2	2023-02-27T13:29:06Z	Zeid	Selalu pengen ipin, mau secinta apapun sma and...	0
3	2023-02-26T05:04:44Z	Perdjem Huwae	5 tahun kedepan baru beli lah... 🍷	0
4	2023-02-26T02:13:34Z	Kang Mujo	Biasa aja sih menurutku gak terlalu berguna bgt	0
...	...	...	...	...
4023	2022-09-29T07:15:48Z	Alex	first kah?	0
4024	2022-09-29T07:15:48Z	Josep willis		3
4025	2022-09-29T07:15:48Z	Mahesa_Alif	Hallo bg	0
4026	2022-09-29T07:15:48Z	ArmyZhuA	Keduaaa	1
4027	2022-09-29T07:15:48Z	Kenin...		5

Gambar 1. Tampilan Keseluruhan Data

Data seluruh komentar ada 4027 komentar akan tetapi yang di gunakan pada penelitian kali ini hanya 1000 data karena banyak komentar yang tidak sesuai dengan batasan seperti berbahasa Indonesia lalu sesuai dengan 4 aspek yang di tentukan.

4.2. Preprocessing

Preprocessing data, penulis terlebih dahulu melakukan install atau import beberapa library yang tersedia di python, di antaranya yaitu Sastrawi, nltk, dan beberapa basic library seperti pandas, numpy, string, dan lainnya. Berikut hasil preprocessing menggunakan 5 tahapan (case folding, cleaning, tokenizing, stopword removal, stemming).

Tabel 1. Hasil Preprocessing

Proses	Hasil
Case Folding	wah yang Dynamic Island inovasinya keren sih, kelemahan dibuat jadi keunggulan yang mempermudah dan bermanfaat, pasti bentar lagi banyak ditiru ðŸ™...
Cleaning	wah yang dynamic island inovasinya keren sih, kelemahan dibuat jadi keunggulan yang mempermudah dan bermanfaat, pasti bentar lagi banyak ditiru ðŸ™...
Tokenizing	['wah', 'yang', 'dynamic', 'island', 'inovasinya', 'keren', 'sih', 'kelemahan', 'dibuat', 'jadi', 'keunggulan', 'yang', 'mempermudah', 'dan', 'bermanfaat', 'pasti', 'bentar', 'lagi', 'banyak', 'ditiru']
Stopword Removal	['dynamic', 'island', 'inovasinya', 'keren', 'kelemahan', 'dibuat', 'keunggulan', 'mempermudah', 'bermanfaat', 'bentar', 'banyak', 'ditiru']
Stemming	['dynamic', 'island', 'inovasi', 'keren', 'lemah', 'buat', 'unggul', 'mudah', 'manfaat', 'bentar', 'banyak', 'tiru']

Hasil dari preprocessing ini berupa kata dasar dari setiap sentimen yang ada yang nantinya akan di proses kembali menggunakan metode TF-IDF untuk mendapatkan bobot untuk setiap kata yang ada.

4.3. TF-IDF

TF-IDF ini hanya dilakukan pada kolom 'Ulasan' saja, dengan TF-IDF ini, tiap kata yang ada pada dataset telah memiliki bobot masing-masing untuk

selanjutnya dapat digunakan pada proses klasifikasi data.

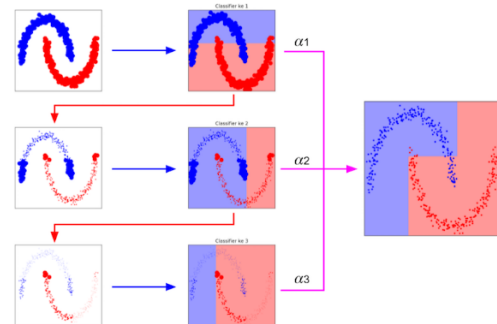
	aaa	aaain	abia	ad	adap	adick	adil	ah	ahabhaa	...	yaasala	ye	yes	yg	youtube	youtube	zaman	zamanbr
0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.000000	0.0	0.0	0.0	0.0	0.0
1	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.266234	0.0	0.0	0.266234	0.0	0.0	0.129485	0.0	0.0	0.0	0.0
2	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.171157	0.0	0.0	0.0	0.0
3	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.0	0.0
4	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.0	0.0
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
245	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.0	0.0
246	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.172202	0.0	0.0	0.0	0.0
247	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.0	0.0
248	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.0	0.0
249	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.000000	0.0	0.0	0.0	0.0

Gambar 2. Hasil TF-IDF

Data yang sebelumnya berupa kata tiap sentimennya dirubah menjadi angka dengan menggunakan metode pembobotan TF-IDF.

4.4. XGBoost

Setelah mendapatkan bobot pada setiap kata pada bagian ini proses pembuatan sebuah model yang digunakan untuk melakukan klasifikasi. Dimana yang sebelumnya sudah melakukan proses preprocessing data, dan labelling pada tahap ini data tersebut akan digunakan sebagai masukkan algoritma klasifikasi, yaitu Extreme Gradient Boosting (XGboost).



Gambar 3. Cara kerja Algoritma XGBoost

Berdasarkan gambar 3 ilustrasi dimana xgboost ini algoritma yang belajar dari kesalahan seperti contoh pada iterasi pertama biru sebagai positif dan merah sebagai negatif setelah di lakukan iterasi pertama hasilnya sudah benar akan memudar dan yang salah akan lebih tebal dan begitu juga pada iterasi kedua hasil kesalahan dari iterasi pertama menjadi bahan untuk dipelajari dan hasilnya yang tadinya sudah benar jadi semakin memudar lagi dan begitu seterusnya lalu hasil dari banyak weak learner tadi di gabungkan. Untuk menghitung pohon yang ada menggunakan persamaan (6) dan untuk menentukan evaluasi model pada pohon selanjutnya menggunakan nilai gain menggunakan persamaan (7). Nantinya hasil prediksi XGBoost akan di hitung pada pengujian yang dilakukan menggunakan metode pengujian Confusion Matrix.

4.5. Confusion Matrix

Bagian ini proses pembuatan sebuah model yang digunakan untuk melakukan klasifikasi. Dimana yang sebelumnya sudah melakukan proses preprocessing

data, labelling dan metode optimasi pada tahap ini data tersebut akan digunakan sebagai masukan algoritma klasifikasi, yaitu Extreme Gradient Boosting (XGboost).

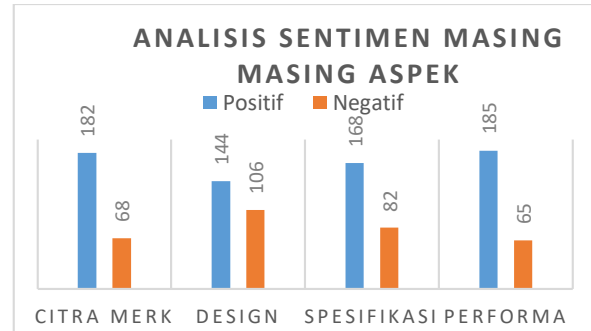
Tabel 2. Hasil Confusion Matrix

Aspek	Training	Testing	Accuracy
Citra Merek	90%	10%	96%
	80%	20%	94%
	70%	30%	96%
Desain	90%	10%	92%
	80%	20%	92%
	70%	30%	89%
Spesifikasi	90%	10%	96%
	80%	20%	92%
	70%	30%	95%
Performa	90%	10%	96%
	80%	20%	96%
	70%	30%	97%

Berdasarkan diagram pada tabel 4.1, dapat dilihat bahwa tingkat akurasi paling baik untuk hasil pengujian ialah yang menggunakan perbandingan data training dan data testing sebesar 90% : 10%. Namun untuk aspek performa akurasi lebih tinggi didapatkan menggunakan perbandingan 70% : 30% tapi tidak terpaut jauh hanya berbeda 1%. Lalu untuk perbandingan 80% : 20% dan 70% : 30% tidak dapat ditentukan yang mana yang lebih baik. Hal ini dikarenakan pada aspek desain pengujian yang menggunakan perbandingan 80% : 20% lebih tinggi nilai akurasinya dibanding dengan pengujian yang menggunakan perbandingan 70% : 30%. Hal ini berbanding terbalik dengan tiga aspek lainnya. Pada aspek citra merek, spesifikasi dan performa perbandingan 80% : 20% nilai akurasinya lebih rendah dibanding dengan pengujian yang menggunakan perbandingan 70% : 30%. Meskipun begitu, pada penelitian tugas akhir ini perbandingan pembagian data yang memberikan nilai akurasi paling tinggi ialah skenario pengujian yang menggunakan perbandingan 80% : 20% dengan hasil akurasi paling tinggi berada di pada skenario pengujian 12 yang mengukur keakuratan proses klasifikasi aspek performa dengan tingkat akurasi mencapai 97%, disusul oleh aspek citra merek, desain, dan spesifikasi, dengan tingkat akurasi masing-masing sebesar 96%, 92%, dan 96%.

4.6. Analisis Sentimen Aspek

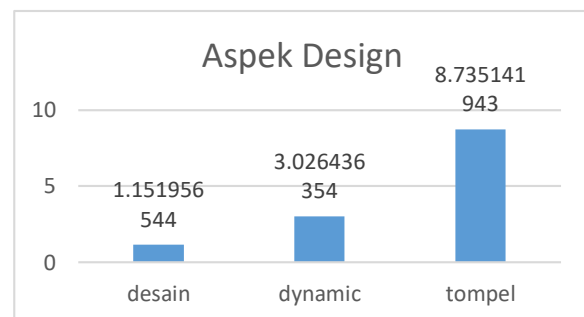
Untuk mengetahui kelebihan dan kelemahan Iphone 14 pro, penulis akan menganalisis kecenderungan ulasan pada masing-masing aspek, apakah cenderung mengandung sentimen positif atau sentimen negatif. berikut akan ditampilkan diagram jumlah sentimen masing-masing aspek.



Gambar 4. Analisis Sentimen Masing Masing Aspek

Berdasarkan gambar 4 dari 1000 data ulasan paling relevan yang di scraping dari kolom komentar youtube pada tanggal 1 Maret 2023, terlihat bahwa ulasan pada aspek desain mengandung sentimen negatif paling banyak di antara keempat aspek lainnya.

Jumlah ulasan dan persentase sentimen pada masing-masing aspek secara tidak langsung mampu memberikan gambaran terkait kelebihan dan kelemahan iphone 14 pro ini. Berdasarkan hasil analisis di atas, dapat diketahui bahwa kelemahan iphone 14 pro ini terletak pada aspek desain dan spesifikasi. Para pengguna iphone merasa desain pada 14 pro ini terlalu monoton sama dengan seri sebelumnya dan inovasi baru dynamic island juga mengganggu di beberapa aplikasi yang ada seperti tombol search pada aplikasi tokopedia



Gambar 5. Analisis Sentimen Aspek Design

Berdasarkan gambar 5 dari pembobotan kata terdapat kata pada sentimen yang di prediksi negatif pada aspek desain, terlihat bahwa kata dynamic dan tombol (kata lain dari dynamic island karena bentuknya menyerupai tombol) menjadi kata teratas dalam aspek desain hal ini membuktikan beberapa masyarakat Indonesia memberikan sentimen negatif terhadap kata dynamic

Selain itu masalah pada kamera iphone 14 pro dimana kekurangan pada seri sebelumnya belum terselesaikan yaitu kamera yang ngeflare pada malam hari. Sedangkan untuk kelebihannya, iphone 14 pro ini cukup unggul untuk chipset yang digunakan yaitu A16 bionic dan inovasi yang selalu dihadirkan di setiap serinya dari hasil ulasan dan persentase sentimen iphone 14 pro berhasil membuat banyak masyarakat indonesia puas terhadap keseluruhan inovasi dan spesifikasi yang diberikan di iphone 14 pro ini.

5. KESIMPULAN DAN SARAN

Berdasarkan penelitian yang telah dilakukan, kesimpulan yang didapat dari hasil penelitian tugas akhir ini ialah sebagai berikut: Algoritma Extreme Gradient Boost berhasil diimplementasikan untuk analisis sentimen pada kolom komentar video youtube iPhone 14 Pro pada channel GadgetIn. Tingkat akurasi penggunaan algoritma Extreme Gradient Boost dalam analisis sentimen pada kolom komentar video youtube iPhone 14 Pro pada channel GadgetIn berkisar antara 89% hingga 97%. Akurasi tertinggi pada masing-masing aspek didapatkan dari skenario pengujian yang menggunakan perbandingan 80% : 20% pada pembagian data training dan data testing. iPhone 14 pro memiliki kelemahan iPhone 14 pro ini terletak pada aspek desain dan spesifikasi. Para pengguna iPhone merasa desain pada 14 pro ini terlalu monoton sama dengan seri sebelumnya dan inovasi baru dynamic island juga mengganggu di beberapa aplikasi yang ada seperti tombol search pada aplikasi tokopedia selain itu masalah pada kamera iPhone 14 pro dimana kekurangan pada seri sebelumnya belum terselesaikan yaitu kamera yang ngeflare pada malam hari. Sedangkan untuk kelebihan, iPhone 14 pro ini cukup unggul untuk chipset yang digunakan yaitu A16 bionic dan inovasi yang selalu dihadirkan di setiap serinya dari hasil ulasan dan persentase sentimen iPhone 14 pro berhasil membuat banyak masyarakat Indonesia puas terhadap keseluruhan inovasi dan spesifikasi yang diberikan di iPhone 14 pro ini.

Berikut ini adalah beberapa saran berdasarkan penelitian yang telah dilakukan, yaitu : Pada pengembangan selanjutnya dapat ditambahkan jumlah dataset yang digunakan agar meningkatkan akurasi dan prediksi dari proses analisis sentimen. Diharapkan penelitian selanjutnya dapat menggunakan aspek-aspek lain agar kelebihan dan kelemahan aplikasi dapat makin teridentifikasi.

DAFTAR PUSTAKA

- [1] Andi Dwi Riyanto. (2022, February 15). *Hootsuite (We are Social): Indonesian Digital Report 2022*. Hootsuite.
- [2] Pasek, P., Mahawardana, O., Sasmita, G. A., Agus, P., & Pratama, E. (2022). Analisis Sentimen Berdasarkan Opini dari Media Sosial Twitter terhadap "Figure Pemimpin" Menggunakan Python. In *JITTER-Jurnal Ilmiah Teknologi dan Komputer* (Vol. 3, Issue 1).
- [3] Muhammad Arief Rahman, Herman Budianto, & Ester Irawati Setiawan. (2019). Aspect Based Sentimen Analysis Opini Publik Pada Instagram dengan Convolutional Neural Network. *JOURNAL OF INTELLIGENT SYSTEMS AND COMPUTATION*, 1(2), 50–57. <https://doi.org/https://doi.org/10.52985/insyst.v1i2.83>
- [4] Ibrahim Ahmed Osman, A., Najah Ahmed, A., Chow, M. F., Feng Huang, Y., & El-Shafie, A. (2021). Extreme gradient boosting (Xgboost) model to predict the groundwater levels in Selangor Malaysia. *Ain Shams Engineering Journal*, 12(2), 1545–1556. <https://doi.org/10.1016/j.asej.2020.11.011>
- [5] Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 13-17-August-2016, 785–794. <https://doi.org/10.1145/2939672.2939785>.
- [6] Andi Nurul Hidayat. (2015). ANALISIS SENTIMEN TERHADAP WACANA POLITIK PADA MEDIA MASA ONLINE MENGGUNAKAN ALGORITMA SUPPORT VECTOR MACHINE DAN NAIVE BAYES (Vol. 1, Issue 1).
- [7] Andi Nurul Hidayat. (2015). ANALISIS SENTIMEN TERHADAP WACANA POLITIK PADA MEDIA MASA ONLINE MENGGUNAKAN ALGORITMA SUPPORT VECTOR MACHINE DAN NAIVE BAYES (Vol. 1, Issue 1).
- [8] Diki Hendriyanto, M., Ridha, A. A., & Enri, U. (2022). ANALISIS SENTIMEN ULASAN APLIKASI MOLA PADA GOOGLE PLAY STORE MENGGUNAKAN ALGORITMA SUPPORT VECTOR MACHINE SENTIMENT ANALYSIS OF MOLA APPLICATION REVIEWS ON GOOGLE PLAY STORE USING SUPPORT VECTOR MACHINE ALGORITHM. *Journal of Information Technology and Computer Science (INTECOMS)*, 5(1).
- [9] Syulistyo, A. R., Jati Purnomo, D. M., Rachmadi, M. F., & Wibowo, A. (2016). PARTICLE SWARM OPTIMIZATION (PSO) FOR TRAINING OPTIMIZATION ON CONVOLUTIONAL NEURAL NETWORK (CNN). *Jurnal Ilmu Komputer Dan Informasi*, 9(1), 52. <https://doi.org/10.21609/jiki.v9i1.366>
- [10] Rifky Hendrawan, I. (2022). "Jurnal TRANSFORMASI (Informasi & Pengembangan Iptek)" (STMIK BINA PATRIA) PERBANDINGAN ALGORITMA NAIVE BAYES, SVM DAN XGBOOST DALAM KLASIFIKASI TEKS SENTIMEN MASYARAKAT TERHADAP PRODUK LOKAL DI INDONESIA. In *Jurnal TRANSFORMASI* (Vol. 18, Issue 1).
- [11] Farah Zhafira, D., Rahayudi, B., & Korespondensi, P. (2021). ANALISIS SENTIMEN KEBIJAKAN KAMPUS MERDEKA MENGGUNAKAN NAIVE BAYES DAN PEMBOBOTAN TF-IDF BERDASARKAN KOMENTAR PADA YOUTUBE (Vol. 2, Issue 1).