

Analisis sentimen terhadap aplikasi Bukalapak sebelum dan sesudah IPO menggunakan algoritma Naïve Bayes

Bayu Yanuargi*¹

- 1 Program Magister Teknik Informatika Program Pascasarjana
Universitas Amikom Yogyakarta
Jl. Padjajaran, Ring Road Utara, Kel. Condongcatur, Kec. Depok, Kab.
Sleman, Prop. Daerah Istimewa Yogyakarta
bayu.yanuargi@students.amikom.ac.id

Abstrak

Bukalapak adalah salah satu startup eCommerce paling awal yang berdiri di Indonesia. Bukalapak sudah menjembatani antara penjual dengan pembeli sejak 2010. Pada tahun 2021 Bukalapak memberanikan diri untuk melakukan *initial public offers* (IPO) di Bursa Efek Indonesia. Banyak ragam tanggapan dari pengguna Bukalapak terhadap langkah Bukalapak tersebut, baik positif maupun negatif. Sentimen negatif atau positif tersebut dapat menjadi masukan dan evaluasi bagi Bukalapak sendiri untuk menjaga loyalitas penggunanya. Proses penelitian ini dimulai dari pengumpulan data yang diperoleh dari *scrapping* data *review* produk Bukalapak di Google *Playstore* sebelum dan sesudah IPO. Kemudian dilakukan preprosesing data mulai dari *casefolding*, penghilangan *stop words*, tokenisasi, *stemming* hingga TF-IDF. Hasil dari preprosesing tersebut kemudian dijadikan data untuk melakukan klasifikasi menggunakan Naïve Bayes. Klasifikasi tersebut kemudian diuji dan mendapatkan nilai akurasi untuk data sebelum IPO sebesar 77% dan data setelah IPO sebesar 76%.

Kata Kunci analisis sentimen, bukalapak, naive Bayes, tf-idf, confusion matrix

Digital Object Identifier 10.36802/jnanaloka.v3-no1-17-25

1 Pendahuluan

Bukalapak merupakan salah satu startup yang bergerak di bidang e-Commerce sejak 2010, dan pada tahun 2021 Bukalapak sebagai salah satu Unicorn e-Commerce melakukan *initial public offers* (IPO). Dengan diluncurkannya BUKA ke ranah publik pada akhir Juli 2021, maka Bukalapak haruslah mampu meningkatkan pelayanan dan peka terhadap suara konsumen. Dinamika harga saham (naik / turun) selain ditentukan oleh kinerja perusahaan, maka suara konsumen juga akan menjadi salah satu faktor penentu dalam naik / turunnya nilai saham. Sebagai salah satu unicorn dengan jumlah pengguna sebesar 104 juta per Desember 2020 [1] maka menjaga kesetiaan pelanggan adalah hal yang sangat penting. Salah satu proses untuk menjaga kesetiaan pelanggan adalah dengan mengidentifikasi permasalahan yang dihadapi oleh pengguna. Salah satu kanal dalam mengidentifikasi permasalahan pengguna adalah berdasarkan umpan balik yang diberikan pengguna di *Playstore*.

Umpan balik atau komentar yang diberikan pengguna di kanal *Playstore* sangat penting untuk dikelola. Hal tersebut disebabkan karena komentar di *Playstore* dapat mempengaruhi citra Bukalapak bagi pengguna baru maupun lama. Untuk itu segala permasalahan yang

* Corresponding author.



ditulis di kolom komentar aplikasi Bukalapak harus bisa diserap untuk ditindak lanjuti. Tetapi dengan jumlah pengguna yang sedemikian besar dan jumlah komentar yang besar pula, maka proses pengelolaan tentu saja tidak bisa dilakukan dengan cara manual, karena hal tersebut tidak saja tidak produktif, tetapi informasi yang dihasilkan juga tidak akurat. Salah satu metode yang dapat digunakan untuk melakukan analisis tersebut di atas adalah dengan menggunakan analisis sentimen atau biasa disebut opinion mining. Analisis sentimen ialah proses mengekstraksi, mengolah dan memahami data berupa teks yang tidak terstruktur secara otomatis guna mengambil informasi sentimen yang terdapat pada sebuah kalimat pendapat atau opini [2]. Sedangkan menurut [3], analisis sentimen pada *review* adalah proses menyelidiki *review* produk di internet untuk menentukan opini atau perasaan terhadap suatu produk secara keseluruhan.

Untuk melakukan analisis sentimen analisis, salah satu metode dalam machine learning yang sering digunakan adalah Naïve Bayes, dimana menurut [4], pengklasifikasi Naïve Bayes biasanya digunakan dalam kategorisasi teks. Ide dasarnya adalah menggunakan probabilitas gabungan kata dan kategori untuk memperkirakan probabilitas kategori yang diberikan dokumen. Dalam penelitiannya, terlepas dari asumsi independensi yang tidak realistis, pengklasifikasi Naïve Bayes secara mengejutkan mencapai kinerja yang sebanding, atau lebih baik daripada SVM. Interaksi antara metode klasifikasi dan opsi penyajian fitur diamati, dan frekuensi bigram terbukti merupakan fitur yang efektif dalam menangkap sentimen dalam teks Kanton. Selain itu, kami melihat efek ukuran set fitur pada kinerja klasifikasi. Saat ukuran fitur meningkat, akurasi kedua pengklasifikasi mencapai puncaknya dan kemudian menurun karena overfitting, dengan pengklasifikasi SVM (*support vector machine*) kurang sensitif terhadap jumlah fitur [4]. Beberapa penelitian menunjukkan bahwa Naïve Bayes memiliki hasil yang cukup bagus, seperti pada penelitian oleh [5] dengan score akurasi yang cukup tinggi sebesar 0.865. Pada penelitian lainnya yang dilakukan oleh [6] membuktikan bahwa Naïve Bayes nilai akurasi sebesar 98.67%. Kedua penelitian tersebut akan diterangkan lebih lanjut pada bagian metodologi. Berdasarkan pada nilai akurasi yang tinggi tersebut, maka pada penelitian ini metode klasifikasi yang digunakan adalah Naïve Bayes dengan harapan agar bisa menghasilkan hasil yang optimal.

Penggunaan metode pengklasifikasian teks Naïves Bayes, diharapkan dalam penelitian ini akan menghasilkan kata-kata yang diklasifikasikan sebagai sentiment negative sebagai masukan untuk pengembangan produk Bukalapak menjadi lebih baik. Karena dengan berhasilnya identifikasi permasalahan dalam *review* produk akan dapat meningkatkan kepuasan pengguna, dimana pengguna akan merasakan bahwa masukan mereka telah ditindaklanjuti oleh Bukalapak. Dalam tulisan ini akan dilakukan analisis sentimen sebelum dan sesudah Bukalapak meluncurkan IPO, dengan tujuan untuk melihat bagaimana response konsumen setelah peluncuran IPO tersebut terhadap produk Bukalapak.

Banyak penelitian telah dilakukan terkait sentimen analisis dengan menggunakan Naïves Bayes, seperti [5] yang melakukan analisis sentimen untuk proses pemesanan tiket dari pegipegi.com dengan menggunakan Naïves Bayes. Data Latih yang digunakan adalah data yang berhasil disrape dari website pegipegi.com dengan jumlah kalimat positif 500 dan kalimat negatif sebanyak 500. Pada penelitian tersebut dilakukan preprosesing teks hingga TF-IDF hingga perolehan *information gain* sebagai target utama dalam preprocessing teks. Klasifikasi Naïves Bayes dilakukan dengan menghitung nilai probabilitas dari data testing yang kemudian dikategorikan ke dalam dua kelas yaitu positif dan negatif, maskapai dengan jumlah kalimat positif terbanyak disimpulkan merupakan Maskapai terbaik. Hasil observasi metode Naïves Bayes setelah dan sebelum ditambahkan *information gain*, diperoleh bahwa nilai akurasi terbaik adalah setelah ditambahkan *information gain* dengan nilai akurasi

0.865, sedangkan *F1-score* memiliki nilai lebih tinggi setelah penambahan *information gain* yaitu sebesar 0,864. Sedangkan secara *performance* kecepatan, *information gain* memberikan efisiensi yang cukup tinggi dengan pengolahan 1000 data hanya membutuhkan waktu 93.63 detik.

Penelitian [7] melakukan analisis sentimen menggunakan algoritma *support vector machine* (SVM), dimana dalam penelitian tersebut dilakukan analisis sentimen pada data tweet terkait rencana pemindahan ibu kota negara. Dengan menggunakan SVM klasifikasi dilakukan pada data tweet menjadi positif dan negatif, dimana tujuan akhirnya adalah untuk mendapatkan bagaimana kecenderungan sentimen publik pada wacana pemindahan ibukota. Metode pengumpulan data pada penelitian tersebut adalah dengan melakukan *crawling* data twitter dengan menggunakan metode *crawling* URL. Dengan menggunakan dataset sebanyak 1236 tweet dilakukan pelabelan dengan menggunakan dua orang anotator, dimana anotator pertama melakukan klasifikasi positif dan negatif, sedangkan anotator kedua melakukan pengujian terhadap hasil klasifikasi anotator pertama. Hasil pelabelan diketahui bahwa sentimen positif sebanyak 404 dan negatif sebanyak 832. Hasil dari pelabelan tersebut kemudian dilakukan pengujian dengan menggunakan SVM dengan menghasilkan *confusion matrix* dengan akurasi=96,68%, *precision*=95.82%, *recall*=94.04% dan AUC = 0,979.

Penelitian sentimen analisis multi aspek menggunakan data ikon emosi dan *review* yang diberikan oleh pengguna Tripadvisor dengan data dikumpulkan menggunakan teknik scraping rapidminer melalui interaksi dengan API [6]. Penelitian tersebut menggunakan algoritma Naïve Bayes dikombinasikan dengan metode pelabelan multi aspek yang disertai konversi ikon emosi menjadi teks ulasan. Setelah data melalui tahap text preprocessing, maka didapatkan ribuan kata yang sudah siap untuk diproses ke tahap selanjutnya yaitu *multi-aspect sentence labelling*. Proses ini akan melalui beberapa tahap yaitu klustering (pengelompokan), dan mengidentifikasi dokumen awal untuk setiap kelas. Dalam penelitian ini dihasilkan tiga jenis probabilitas, yaitu pertama probabilitas bersyarat (*likelihood*) dokumen dan menghasilkan positif 19.029206, negatif 0, netral 449.121081. Kedua adalah Probabilitas Prior dengan hasil kelas positif 0.33, kelas negatif 0.33 dan kelas netral 0.33. Probabilitas ketiga adalah probabilitas Posterior dengan hasil kelas positif 0.697737, negatif 0 dan netral 16.467772. Berdasarkan probabilitas ketiga, diketahui bahwa kelas yang tertinggi adalah kelas Netral, karena memiliki nilai yang terbesar. Pengujian yang dilakukan menggunakan model 10-fold cross validation dimana hasil analisis sentimen yang dilakukan menggunakan metode Naïve Bayes classifier dengan penambahan fitur konversi *emoticon* dan *multiaspect sentence labelling* mendapatkan nilai akurasi sebesar 98.67% dibanding dengan sebelum digunakan *multiaspect sentence labelling* yang hanya sebesar 88.78%.

2 Metodologi

Dalam penelitian ini dilakukan pengambilan data secara scrapping *review* dari Google Playstore untuk aplikasi Bukalapak yang pada alamat com.bukalapak.android. Pemfilteran manual dilakukan untuk membatasi pengambilan data hanya pada tiga bulan sebelum dan tiga bulan sesudah IPO. Seperti kita ketahui bahwa IPO Bukalapak dilakukan pada bulan Juli 2021, maka data sebelum IPO digunakan data pada bulan April sampai Juni, dan setelah IPO digunakan data pada bulan Agustus hingga Oktober. Dengan jumlah dataset disetiap periode adalah 500 baris, dengan data uji sebanyak 20% dan data latih 80%. Untuk data yang berhasil diambil, dilakukan pelabelan negatif dan positif dengan menggunakan rating yang diberikan oleh pengguna, untuk bintang 4 dan 5 diberikan label positif dan bintang 1,2 dan 3 diberikan label negatif. Walaupun sebenarnya rating 3 masih cenderung netral, tetapi

secara produk rating 3 merupakan rating yang tidak bagus. Berdasarkan pada pelabelan tersebut diperoleh jumlah data seperti dalam Tabel 1.

■ **Tabel 1** Sebaran dataset

sentimen	pra IPO	pasca IPO
negatif	245	264
positif	255	236
total	500	500

2.1 Tahapan penelitian

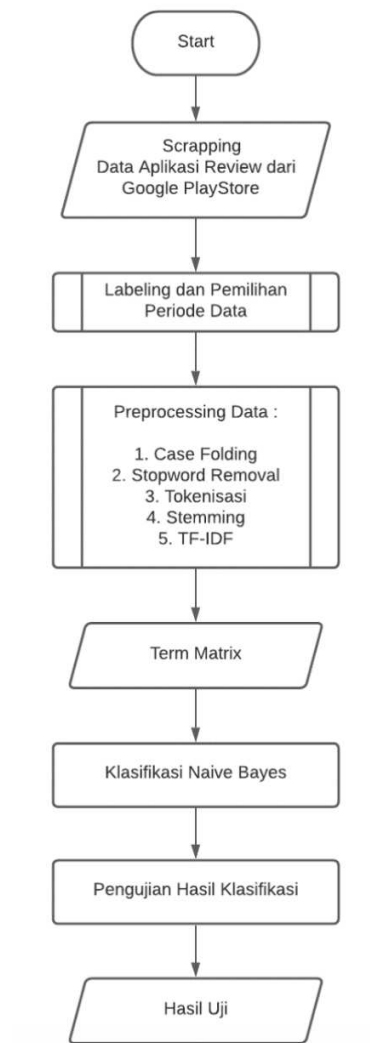
Tahapan penelitian yang dilakukan dari dua jenis data tersebut di atas adalah sama, mulai dari pengambilan data, preprocessing teks, klasifikasi teks dan pengujian algoritma. Untuk lebih jelasnya terkait dengan tahapan penelitian dapat dilihat pada Gambar 1.

Sebelum analisis sentimen dapat dilakukan dengan efektif dan efisien dan juga untuk meningkatkan akurasi dari hasil analisis atau klasifikasi, maka data yang kita miliki haruslah dipersiapkan dan diolah dengan baik dilakukan *preprocessing*. Langkah *preprocessing* berfungsi untuk mengolah data sedemikian rupa sehingga bisa dianalisis dengan cepat dan akurat, penghilangan simbol, kata-kata tidak penting dan juga imbuhan akan bisa menghasilkan hasil yang akurat dan cepat. Dalam penelitian ini dilakukan langkah *preprocessing* yakni, *casefolding*, penghilangan *stopwords*, tokenisasi, *stemming* dan pembobotan dengan *Term Frequency Invers Document* (TF-IDF).

Casefolding secara definisi adalah melakukan standardisasi atau konsistensi dalam penggunaan huruf kapital. Hal ini perlu dilakukan karena huruf kapital tidak dibutuhkan dalam *preprocessing* dan klasifikasi, sehingga dengan standardisasi teks, proses data menjadi lebih efisien dan rapi. Penghilangan *stopwords* dilakukan dengan menggunakan *filtering* karena kata-kata ini tidak deskriptif dan tidak membantu dalam proses analisis. Sebagai contoh *stopwords* dalam bahasa Indonesia adalah “yang”, “di”, “atau”, “dan”, sedangkan dalam bahasa Inggris contohnya “about”, “again”, “and”, “or”, “any”, dst. Untuk melakukan *filtering stopwords* digunakan *script* menggunakan library NLTK.

Tokenisasi merupakan proses penting dalam NLP agar data dapat dimulai untuk dilakukan sebuah analisis, klasifikasi, *clustering* berdasarkan kata, daripada kalimat. Kamus yang digunakan untuk menilai sebuah kalimat pastilah berdasarkan kata perkata, bukan kalimat, sehingga tokenisasi sangat penting untuk dilakukan dalam *preprocessing*. Proses tokenisasi dalam penelitian ini menggunakan library NLTK. *Stemming* merupakan proses untuk mengembalikan kata kedalam bentuk bakunya, sehingga setelah dilakukan pengembalian kedalam bentuk bakunya diharapkan proses pengelompokan kata akan menjadi lebih akurat. Sebagai contoh misalnya kepunyaannya, punyaku, punyamu, bisa dikelompokkan kedalam satu kata yakni punya. Pada penelitian ini digunakan *stemmer* Sastrawi, dengan algoritma dasarnya adalah Nazief Adriani. Metode *stemming* ini menggunakan teknik *lemmatization* yang mengembalikan kata yang berimbuhan menjadi kata dasar. Selain itu, pemilihan metode ini berdasarkan pada penelitian oleh [8] yang telah membandingkan akurasi dan kecepatan algoritma Nazief Adriani dan Porter untuk mengukur plagiarisme. Hasil dari penelitian tersebut diperoleh bahwa walaupun kecepatan Porter dua kali lebih cepat dari Nazief Adriani, akurasi Stemmer Nazief Adriani lebih tinggi yaitu 70.5%.

Data yang telah melalui tahap *preprocessing* harus berbentuk numerik. Untuk mengubah data tersebut menjadi numerik yaitu menggunakan metode pembobotan TF-IDF. Metode



■ **Gambar 1** Tahapan analisis sentimen

TF-IDF merupakan metode yang digunakan menentukan seberapa jauh keterhubungan kata (*term*) terhadap dokumen dengan memberikan bobot setiap kata. Proses perhitungan TF-IDF pada penelitian ini dengan cara menghitung bobot setiap kata yang ada dalam data pelatihan. Proses perhitungan dilakukan dengan menggunakan *library* Sklearn yang terdiri dari dua model yaitu *count vectorizer* (*word count*) dan *Tfidfvectorizer* (*word frequencies*). Hasil dari preprocessing TF-IDF ini bisa dikatakan menjadi acuan penting dalam penelitian kali ini karena berdasarkan hasil TF-IDF kita akan dapat memperoleh informasi kata apa yang sering muncul dalam *review* konsumen untuk dua periode yang berbeda tersebut.

Matrix matematika yang mendeskripsikan seberapa besar frekuensi sebuah kata muncul dalam suatu dokumen atau korpus (*document term matrix*) digunakan untuk mempermudah pembacaan dan analisis, karena jumlah kata individual setiap dokumen akan ditampilkan dengan lebih jelas. Berdasarkan kecenderungan kata yang muncul dan dapat dibaca dengan jelas urutan dokumen kemunculannya, maka akan dapat dengan jelas mengetahui informasi kata-kata yang penting dalam *review* produk.

Klasifikasi Naïve Bayes adalah sebuah metode klasifikasi dengan menggunakan metode statistik dan probabilitas yang dikemukakan oleh ilmuwan Thomas Bayes. Menurut penelitian [4] yang membandingkan penggunaan Naïve Bayes dengan SVM dalam klasifikasi sentiment restoran online, Naïve Bayes memiliki performa yang lebih baik dari SVM, dimana akurasi Naïve Bayes mencapai 95.33% sedangkan SVM memiliki akurasi 90%. Metode Naïve Bayes memiliki dua tahapan, yaitu tahapan pelatihan dan tahapan klasifikasi. Dalam penelitian ini untuk tahapan pelatihan digunakan nilai 20% dari total data dan tahap klasifikasi menggunakan nilai 80%.

Untuk mengetahui performa klasifikasi yang dilakukan, maka perlu pengujian atau evaluasi hasil klasifikasi. Dalam penelitian ini digunakan metode *confusion matrix* untuk mendapatkan tingkat akurasi, presisi, F1-score dan *recall* dari hasil klasifikasi. *Confusion matrix* membagi dua kelas dalam masalah klasifikasi seperti pada gambar di bawah ini. Hal itu bisa membentuk empat hasil yang berbeda dalam perkiraan. Hasil klasifikasi yang benar dapat dibagi menjadi dua benar-benar positif (*true positive*) dan benar-benar negatif (*true negative*), sedangkan nilai positif yang salah (*false positive*) dan nilai negatif yang salah (*false negative*) merupakan dua tipe kesalahan. Contoh dari *false positive* adalah dimana data yang salah diklasifikasikan menjadi data yang benar, sedangkan contoh *false negative* adalah dimana data yang benar diklasifikasikan sebagai data yang salah.

3 Hasil dan pembahasan

Tabel 2 merupakan contoh hasil pemrosesan data yang berhasil dilakukan menggunakan Google Colab. Dimana tahapan preprosesing yang dilakukan dimulai dari *case folding*, *stop words removal*, tokenisasi, *stemming*.

■ **Tabel 2** Contoh hasil pemrosesan data

	sebelum IPO	sesudah IPO
teks asli	User friendly market place, a lot of beneficial promotional programs.	Beberapa kali menggunakan aplikasi ini untuk jual beli barang, sangat membantu.
	Aplikasi nya bagus. Sedikit masukan aplikasi nya kurang responsif, user experience nya jadi agak lemot	Terima kasih karna buka lapak saya bisa membeli kelas prakerja dengan mudah dan cepat
casefolding	user friendly market place, a lot of beneficial promotional programs.	beberapa kali menggunakan aplikasi ini untuk jual beli barang, sangat membantu.
	aplikasi nya bagus. sedikit masukan aplikasi nya kurang responsif, user experience nya jadi agak lemot	terima kasih karna buka lapak saya bisa membeli kelas prakerja dengan mudah dan cepat
stop words, tokenisasi, stemming	['user', 'friendly', 'market', 'place', 'a', 'lot', 'of', 'beneficial', 'promotional', 'programs']	['beberapa', 'kali', 'menggunakan', 'aplikasi', 'jual', 'beli', 'barang', 'bantu']
	['aplikasi', 'bagus', 'butuh', 'tambah', 'fitur', 'lengkap', 'goodjob', 'bukalapak']	['terima', 'kasih', 'karna', 'buka', 'lapak', 'beli', 'kelas', 'prakerja', 'mudah', 'cepat']

Data yang berhasil diambil kemudian dilakukan pelabelan negatif dan positif dengan menggunakan rating yang diberikan oleh pengguna, untuk bintang 4 dan 5 diberikan label positif dan untuk data dengan bintang satu hingga tiga diberikan label negatif. Tahapan

teks *preprocessing* terakhir yang dilakukan adalah melakukan pembobotan TF-IDF dan menghasilkan *document term matrix*.

Dalam proses ini dihasilkan kata-kata yang sering muncul dalam *review* aplikasi Bukalapak di Google *Playstore*. Hasil *document term matrix* dari TF-IDF dapat dilihat pada Tabel 3 yang menampilkan urutan 20 besar. Dari *document term matrix* di Tabel 3, dapat dilihat bahwa urutan kata yang sering keluar dalam *review* produk di Google *Playstore* sama pada dua periode sampai urutan ke lima, kemudia mulai masuk beberapa kata baru termasuk kata “Saham” yang terkait dengan IPO sebanyak 49 kali.

■ **Tabel 3** Hasil term matrix 20 kata teratas

no	sebelum IPO	jumlah	sesudah IPO	jumlah
1	bukalapak	163	bukalapak	197
2	aplikasi	105	aplikasi	123
3	barang	103	barang	96
4	beli	90	beli	90
5	jual	80	jual	70
6	dana	79	app	104
7	lapak	73	belanja	57
8	buka	69	lapak	51
9	seller	64	saham	49
10	belanja	63	dana	48
11	transaksi	52	seller	43
12	bayar	47	bayar	43
13	mudah	46	bintang	42
14	ongkir	46	cs	38
15	bagus	43	kirim	38
16	app	43	update	37
17	kirim	43	marketplace	37
18	online	39	mudah	37
19	akun	38	transaksi	34
20	masuk	38	kecewa	33

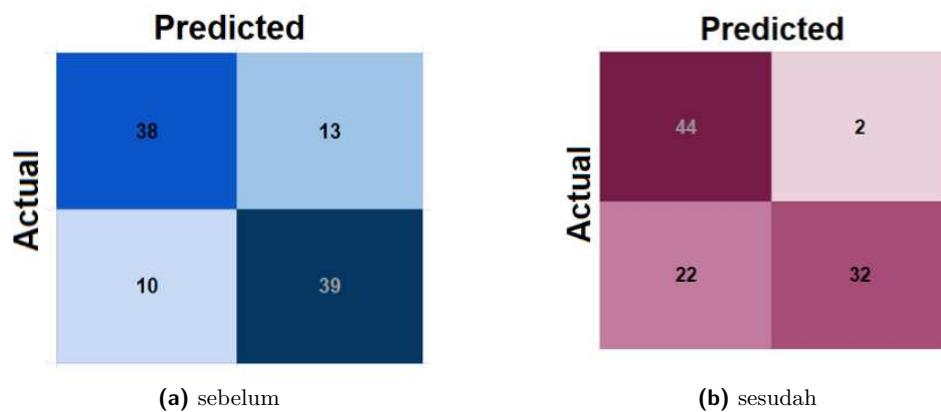
Pada penelitian ini dilakukan pemodelan labeling positif dan negatif menggunakan algoritma Naïve bayes. Dimana akan dibuat labeling secara otomatis berdasarkan dataset yang sudah ada untuk memprediksi labeling baru. Hasil dari pelabelan yang baru dapat dilihat pada Tabel 4.

■ **Tabel 4** Jumlah label setiap periode

sentimen	pra IPO	pasca IPO
negatif	51	46
positif	49	54

Berdasarkan pada Table 4 di atas, bisa diketahui bahwa komposisi negatif dan positif adalah 52.58% dibanding 47.57%. Komposisi ini berbeda dengan komposisi data training dimana perbandingan negatif dan positif adalah 48.13% dibanding 47.57%. Selain itu juga ditunjukkan bahwa sentimen negatif mengalami penurunan setelah peluncuran saham ke publik.

Untuk mengetahui tingkat akurasi dari hasil klasifikasi Naïve bayes, maka perlu dilakukan pengujian. Pada penelitian kali ini dilakukan pengujian dengan menggunakan metode *confusion matrix*. Pengujian algoritma menggunakan confusion matrix ini berguna untuk mengkalkulasi berapa banyak prediksi yang berhasil (TP) dengan prediksi sentimen positif dan aktualnya positif, serta TN dengan prediksi sentimen negatif aktualnya pun negatif. Hasil dari pengujian algoritma Naïve Bayes yang digunakan adalah seperti pada Gambar 2.



■ **Gambar 2** Data *confusion matrix* sebelum & sesudah IPO

Berdasarkan pada *confusion matrix* tersebut dapat dilihat bahwa nilai TP (*true positive*) mengalami kenaikan sebelum dan sesudah IPO dari 38 ke 44, dan nilai TN (*true negative*) mengalami penurunan dari 39 ke 32. Adapun nilai akurasi, presisi, recall dan *f1-score* dari hasil pengujian dalam penelitian ini yang didapatkan dari perhitungan *confusion matrix* adalah seperti tertampil di Tabel 5.

■ **Tabel 5** Hasil pengujian klasifikasi

	sebelum IPO	sesudah IPO
akurasi	0.7700	0.7600
presisi	0.7713	0.6667
recall	0.7700	0.9565
F1-score	0.7699	0.7857

Dalam tabel Tabel 5 tampak terjadi penurunan nilai akurasi dan presisi dari data sebelum IPO dan data setelah IPO walaupun tidak signifikan akan menarik untuk diteliti lebih lanjut. Hal ini bisa dikarenakan periode waktu pengambilan data yang berbeda sehingga menghasilkan data set yang berbeda. Perbedaan dataset tentunya akan berpengaruh pada pre processing data sehingga berpengaruh pada tingkat akurasi dan presisi.

4 Kesimpulan

Dari keseluruhan kegiatan analisis sentimen pada penelitian ini dengan menggunakan klasifikasi Naïves bayes, nilai akurasi klasifikasi yang didapatkan adalah 77% untuk dataset sebelum IPO berlangsung dan 76% untuk dataset yang diambil setelah IPO. Walaupun akurasi tidak terlalu tinggi, tapi sudah cukup bagus untuk menghasilkan nilai prediksi yang akurat berdasarkan dataset yang digunakan. Berdasarkan prediksi yang dilakukan, sentimen

negatif pada aplikasi Bukalapak mengalami penurunan setelah IPO dan kenaikan sentimen positif. Hal ini cukup bagus dimana konsumen cukup antusias dengan berlangsungnya IPO.

Pengembangan penelitian lebih lanjut adalah bagaimana melakukan prediksi dan klasifikasi sentimen untuk mengetahui jenis kata apa saja yang ada dalam sentimen negatif atau positif. Hal ini tentunya sangat bermanfaat untuk mengetahui secara pasti apa yang menjadi keluhan konsumen dengan mengetahui kata-kata terbanyak (*term frequency*) yang ada dalam sentimen negatif.

Pustaka

- 1 “Rekor bukalapak di bursa saham,” 2021, diakses 29 Oktober 2021. [Online]. Available: <https://katadata.co.id/ariayudhistira/infografik/610d4b596bf63/rekor-bukalapak-di-bursa-saham>
- 2 B. Brahimi, M. Touahria, and A. Tari, “Improving sentiment analysis in arabic: A combined approach,” *Journal of King Saud University-Computer and Information Sciences*, vol. 33, no. 10, pp. 1242–1250, 2021.
- 3 B. Liu, “Sentiment analysis and opinion mining,” *Synthesis lectures on human language technologies*, vol. 5, no. 1, pp. 1–167, 2012.
- 4 Z. Zhang, Q. Ye, Z. Zhang, and Y. Li, “Sentiment classification of internet restaurant reviews written in cantonese,” *Expert Systems with Applications*, vol. 38, no. 6, pp. 7674–7682, 2011.
- 5 A. B. P. Negara, H. Muhandi, and I. M. Putri, “Analisis sentimen maskapai penerbangan menggunakan metode naive bayes dan seleksi fitur information gain,” *J. Teknol. Inf. dan Ilmu Komput*, vol. 7, no. 3, 2020.
- 6 S. A. Azzahra and A. Wibowo, “Analisis sentimen multi-aspek berbasis konversi ikon emosi dengan algoritme naïve bayes untuk ulasan wisata kuliner pada web tripadvisor,” *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 7, no. 4, pp. 737–744, 2020.
- 7 P. Arsi and R. Waluyo, “Analisis sentimen wacana pemindahan ibu kota indonesia menggunakan algoritma support vector machine (svm),” *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 8, no. 1, pp. 147–156, 2021.
- 8 A. Rahmatulloh, N. I. Kurniati, I. Darmawan, A. Z. Asyikin, and D. Witarsyah, “Comparison between the stemmer porter effect and nazief-adriani on the performance of winnowing algorithms for measuring plagiarism,” *International Journal on Advanced Science, Engineering and Information Technology*, vol. 9, no. 4, pp. 1124–1128, 2019.