

Analisis Sentimen Kualitas Layanan Teknologi Pembayaran Elektronik pada Twitter (Studi Kasus Ovo dan Dana)

Enggarbela Ogi Intan Pratiwi¹, Wiyli Yustanti²

^{1,2} Jurusan Teknik Informatika/Program Studi S1 Sistem Informasi, Universitas Negeri Surabaya

¹enggarbela.17051214072@mhs.unesa.ac.id

²wiyliyustanti@unesa.ac.id

Abstrak— Teknologi memberikan kemudahan untuk pengguna dalam berbagai aspek. Teknologi Finansial atau *fintech* saat ini tengah berkembang di masyarakat. Namun, penyedia jasa *fintech* di Indonesia semakin banyak. OVO dan DANA merupakan salah satu penyedia jasa *fintech* terbesar dan memiliki jumlah pengguna terbanyak di Indonesia. Saat ini media sosial banyak dimanfaatkan oleh perusahaan untuk mendapatkan *feedback* terkait produk yang ditawarkan, salah satunya Twitter. Pelanggan yang puas dan kurang puas dengan layanan yang diberikan perusahaan dapat menuliskannya melalui *tweet* yang diunggah di Twitter. Kemudahan pengambilan data Twitter dimanfaatkan untuk melakukan penelitian ini. Penelitian ini bertujuan untuk menganalisa sentimen pengguna terhadap layanan OVO dan DANA di Twitter dalam cuitan pengguna melalui *tweets* dengan menggunakan tiga model klasifikasi yaitu Naive Bayes Classifier, Support Vector Machine (SVM) dan K-Nearest Neighbor dan mendapatkan hasil bahwa sentimen negatif pengguna masih lebih tinggi dibandingkan sentimen positif terhadap aplikasi OVO dan Dana, SVM merupakan model *machine learning* yang mendapatkan hasil akurasi paling tinggi pada data OVO dan DANA yaitu 81,33% dan 86,23%, NB hanya 80,63% dan 86,19% serta model K-NN sebesar 75,83% dan 86,19% dapat disimpulkan SVM menjadi model paling baik berdasarkan nilai metrics akurasi, presisi, *recall* dan *f-measure* atau *f1*. Selain itu hasil olah kata *wordcloud* ditemukan bahwa kendala terbesar yang dialami OVO dan Dana terletak pada keandalan aplikasi dalam melakukan transaksi dan kemampuan staf dalam merespon keluhan pengguna. Hal tersebut dibuktikan dengan banyaknya komentar negatif OVO dan Dana terkait kendala mereka dalam melakukan transfer uang dan bahkan kehilangan uang mereka karena lambatnya penanganan *customer service*.

Kata Kunci— Teknologi Finansial, Analisis Sentimen, Klasifikasi sentiment

I. PENDAHULUAN

Seiring perkembangan internet di Indonesia, penggunaan internet saat ini telah digunakan untuk membantu berbagai macam aktivitas. Data pengguna internet per-tahun 2020 [1] pengguna internet saat ini telah mencapai 4,9 miliar orang di seluruh dunia atau 60% dari seluruh populasi di dunia.

Teknologi pembayaran elektronik atau saat ini lebih dikenal *fintech* (*Financial technology*) merupakan sebuah inovasi berbasis teknologi yang dibentuk untuk bersaing dengan keuangan tradisional dalam melakukan pembawaan layanan keuangan [2] untuk memberikan kemudahan dan kenyamanan kepada masyarakat dalam melakukan transaksi. Di Indonesia telah terdapat banyak perusahaan yang menyediakan layanan teknologi finansial. Saat ini banyak masyarakat lebih memilih melakukan pembayaran dengan teknologi finansial atau *e-wallet*, daripada metode pembayaran lainnya dan diperkirakan akan terus meningkat hingga tahun 2023 [3].

OVO dan DANA merupakan salah satu penyedia jasa teknologi finansial yang berdiri dari sejak tahun 2017 dan 2018. Salah satu kanal komunikasi yang digunakan OVO dan DANA dengan pelanggan adalah melalui media sosial Twitter. Seiring meningkatnya pengguna media sosial yang ada di Indonesia media sosial dijadikan salah satu kanal komunikasi perusahaan dengan pelanggan. Twitter merupakan media sosial serta berita daring yang digunakan oleh penggunanya untuk berkomunikasi dalam pesan yang singkat yaitu maksimal sebanyak 280 karakter yang disebut "*tweets*" [4]. Limitasi dalam menyampaikan opini pengguna pada *tweets* atau *User Generated Content* ini yang membedakan Twitter dengan media sosial lainnya, sehingga membentuk karakteristik unik yang menyebabkan pengguna diminta untuk menjelaskan sesuatu dengan cepat, singkat, dan padat sehingga lebih cepat dipahami [5]. Penyajian data di media sosial Twitter lebih informatif dibandingkan dengan Facebook dan Instagram [6]. Saat ini pengguna Twitter membicarakan terkait OVO dan DANA melalui *tweets* untuk menyampaikan pendapat mereka.

Salah satu proses analisis yang dapat digunakan adalah text data mining yaitu dengan cara mencari sebuah informasi secara otomatis dari sumber teks yang berbeda [7] sumber data di Twitter adalah *tweets* dari pengguna. Analisis teks untuk mendapat informasi terkait penggambaran positif atau negatif dari suatu opini adalah sentimen analisis [8].

Dengan memanfaatkan data Twitter yaitu *tweets* yang disampaikan untuk OVO dan DANA akhirnya menimbulkan fenomena big data, yaitu terkumpulnya *dataset* yang besar atau kompleks sehingga sulit untuk diproses secara tradisional [9]. *Machine learning* dan *data mining* sudah populer digunakan dalam suatu penelitian dalam berbagai sektor. Meningkatnya

pengguna OVO dan DANA, membuat OVO dan DANA menjadi salah satu penyelenggara *fintech payment* terbesar di Indonesia, serta bagaimana masyarakat Indonesia sering menggunakan Twitter untuk menyampaikan pendapat atau melakukan komplain terkait kualitas layanannya, ditambah dengan data Twitter yang dapat diperoleh secara mudah, sehingga membuat fenomena tersebut menarik untuk diteliti.

Semakin bertambahnya penyedia layanan teknologi finansial di Indonesia, pengguna juga semakin selektif dalam memilih layanan *fintech* yang ingin digunakan. Oleh sebab itu, maka kualitas layanan merupakan salah satu faktor penting yang harus diperhatikan. Dengan memanfaatkan data mining dengan melibatkan metode dari berbagai bidang seperti statistika, sistem basis data, dan *machine learning* [10]. Penelitian ini memanfaatkan UGC (*User Generated Content*) twitter berbentuk teks. Pengumpulan data dari Twitter dilakukan dengan memanfaatkan Rstudio yang menggunakan Bahasa R dan API (*Application Programming Interface*) Twitter. Dengan menggunakan metode *Naïve Bayes Classifier* dikombinasikan dengan *Support Vector Machine* (SVM) dan K-NN (*K-Nearest Neighbor*) memiliki tujuan untuk mengetahui perbedaan performa ketiga model klasifikasi dalam menangani *dataset*, mengetahui hasil dari sentimen publik untuk mengetahui kualitas layanan yang diberikan oleh OVO dan DANA dan mengetahui dimensi kualitas layanan yang perlu diperbaiki atau mendapat perhatian lebih dari data yang didapatkan dari cuitan pengguna melalui *tweets* di media sosial twitter, sehingga dapat membantu perusahaan dalam meningkatkan kualitas layanan yang diberikan kepada pengguna.

II. METODOLOGI

Metode penelitian merupakan alur penelitian yang memiliki tujuan untuk melakukan pengumpulan data yang diperlukan pada penelitian. Alur penelitian ditunjukkan pada Gbr. 1.

Tahapan penelitian yang digunakan dalam penelitian ini dimulai dari membuat latar belakang, melakukan rumusan masalah, menentukan tujuan penelitian, studi literatur, pengumpulan data, pengolahan data, analisis data dan Pembahasan yang terakhir adalah Kesimpulan dan Saran. Pada penelitian ini menggunakan metode *Naïve Bayes*, *Support Vector Machine* (SVM) dan K- *Nearest Neighbor* (K-NN). Proses pengumpulan data dengan memanfaatkan Rstudio merupakan tahap dalam melakukan pembuatan data *training*, *preprocessing* data, dan klasifikasi data [11] dengan menggunakan model *Naïve Bayes Classifier* (NBC), SVM serta dengan menggunakan model K-NN. Tahapan analisa data dapat dilihat pada Gbr 1.



Gbr 1. Tahapan Analisis Data

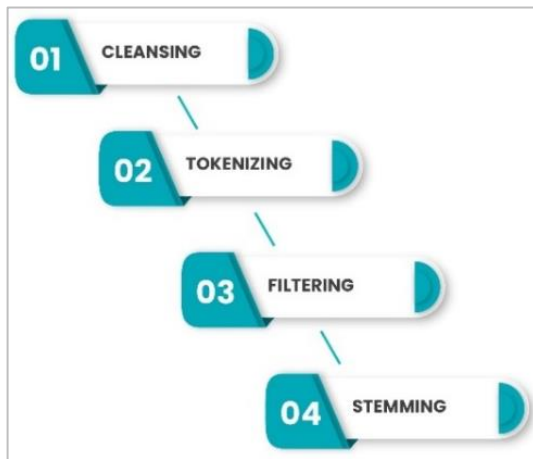
A. Pengumpulan Data

Data *crawling* merupakan tahapan pengumpulan data. Proses pengumpulan data didapatkan dari *tweets* pengguna aplikasi OVO dan DANA di Twitter. Data yang akan diuji dinamakan data latih. Data latih dibentuk berdasarkan *corpus* yang berisi *dataset* terkait opini atau tanggapan pengguna Twitter dengan kandungan sentimen yang dimiliki oleh pengguna [12]. Pengelompokan data akan dilakukan dengan memberikan label pada setiap tanggapan pada *tweet* pengguna dalam *corpus* agar dapat diterapkan pada *machine learning* sebelum melakukan klasifikasi terhadap data *test* [13].

B. Preprocessing Data

Proses *preprocessing* merupakan proses menyiapkan data yang bersih dan siap diolah. Pada tahap *preprocessing* juga proses menghilangkan kata-kata yang tidak relevan dengan topik penelitian ini. Pada Gbr 2. terdapat gambaran tahap yang dilakukan pada proses *preprocessing*.

Cleansing merupakan tahapan pertama dalam proses *preprocessing* yang merupakan proses membersihkan data yang tidak relevan dengan penelitian. *Tokenizing* merupakan tahap untuk memisahkan teks utuh menjadi sebuah token yaitu kata-kata, frasa, serta simbol. Pada tahap *filtering*, dilakukan pembersihan data dengan menghilangkan karakter atau kata yang tidak diperlukan, tahap *filtering* juga menjadi transformasi huruf menjadi lowercase. Tahap terakhir adalah tahap *stemming* yang merupakan proses mengubah kata dasar serta menghilangkan kalimat atau kata imbuhan.



Gbr 2. Tahap Preprocessing

Contoh proses *preprocessing data* dapat dilihat pada Tabel I.

TABEL I.
CONTOH PREPROCESSING DATA

<i>Tweet Mentah</i>	@ovo_id, hi. Saya top up lewat trf virtual acc. Di mutase BCA sdh berhasil, di history transaksi ovo jg sdh berhasil. Tp kok di display gak nambah ya saldonya :(
<i>Tokenizing</i>	@ovo_id hi Saya top up lewat trf virtual acc Di mutase BCA sdh berhasil di history transaksi ovo jg sdh berhasil Tp kok di display gak nambah ya saldonya
<i>Filtering</i>	hi saya top up lewat trf virtual acc Di mutase BCA sdh berhasil di history transaksi ovo jg sdh berhasil Tp kok di display gak nambah ya saldonya
<i>Steamming</i>	Saya top up menggunakan transfer virtual akun dan mutasi BCA sudah berhasil selain itu history transaksi ovo juga sudah berhasil namun display ovo tidak bertambah saldonya

C. Pembobotan

TF – IDF atau *Term Frequency - Inverse Document Frequency* merupakan statistik numerik yang sering digunakan sebagai faktor pembobotan dalam *text mining*, *information retrieval*, dan *user modelling* untuk menggambarkan seberapa penting satu kata bagi sebuah dokumen [14]. Kata-kata umum dalam sebuah kelompok kecil pada sebuah dokumen memiliki nilai TF-IDF yang lebih besar dibandingkan di artikel [15]. Penggunaan TF dapat menggunakan rumus di bawah ini yang disampaikan oleh Tokunaga dan Iwayama (1994) :

$$tf_{ij} = \frac{f_d(i)}{\max f_d(j)}$$

TF menunjukkan dokumen (d) seberapa banyak kata (t) yang muncul. Tokunaga dan Iwayama juga juga menyatakan terkait rumus IDF yaitu :

$$idf_t = \log \frac{N}{df_t}$$

N melambangkan jumlah kata dalam teks, df adalah jumlah teks yang memiliki kata t. Dengan menggabungkan TF dan IDF dalam pengerjaan dapat membantu meningkatkan performa, sehingga seperti halnya yang dipaparkan oleh Manning et al (2009) terkait rumus pembobotan TF-IDF yaitu :

$$W_{t,d} = tf_{d,t} \times idf_t$$

Keterangan :

t = kata kunci, *term*

d = dokumen

W d,t = Bobot d terhadap t

tf = Banyaknya t (kata) yang dicari dalam dokumen

idf = banyak t kebalikan dari kata yang dicari

Sehingga dapat disimpulkan, jika t berulang kali muncul pada sedikit dokumen maka bobot t akan lebih tinggi, kemudian jika t sedikit muncul pada banyak dokumen maka bobot t akan lebih rendah, dan jika t berulang kali muncul pada banyak dokumen maka bobot t akan memiliki nilai terendah.

D. Klasifikasi Teks

Klasifikasi teks pada penelitian ini menggunakan tiga model yaitu *Naïve Bayes Classifier*, *Support Vector Machine* dan *K-Nearest Neighbor* yang mana kedua model tersebut sering digunakan peneliti dalam penelitian analisa sentimen [16]. Tahap ini dilakukan setelah mendapat data bersih, data tersebut kemudian diberikan label secara manual menggunakan bantuan *Microsoft excel*. Pemberian label pada data diberikan batasan yaitu hanya untuk sentimen positif dan negatif. Pemberian label pada sentimen pengguna disesuaikan dengan kalimat yang diutarakan dengan kategori berikut:

1. Sentimen pengguna dianggap sebagai sentimen positif apabila mengandung kata-kata yang memiliki makna sentimen (emosi) positif dari pengguna. Salah satunya berupa kata sanjungan atau pujian. Begitu pula jika jumlah kata sentimen positif yang disampaikan pengguna lebih banyak dari jumlah kata sentimen negatif
2. Sentimen pengguna dianggap negatif jika kalimat pengguna mengandung kata-kata yang bermakna sentimen (emosi) negatif. Seperti halnya kata celaan, hinaan, dan keluhan pengguna. Sentimen juga dianggap negatif apabila jumlah kata dengan sentimen negatif lebih banyak dibandingkan dengan kata dengan sentimen positif

Setelah dilakukan proses *labelling* maka selanjutnya dapat dilakukan proses klasifikasi teks. Pada tahap ini proses klasifikasi dibantu dengan alat RapidMiner dan model yang ditentukan untuk melakukan klasifikasi.

E. Validasi dan Evaluasi Performa Klasifikasi Teks

Bentuk dasar metode *cross-validation* yaitu *k-fold* yang merupakan tahap pembagian data yang mendekati atau sama dengan nilai k. Pada *data training* validasi iterasi k dijalankan dari kumpulan data berbeda [17]. Pada tahap ini dilakukan pengukuran dengan menggunakan *confusion matrix*, pengukuran pada tahap ini dilakukan untuk mendapatkan nilai

akurasi, presisi, *recall* dan *f-measure* atau F1. Penelitian ini menggunakan *k-folds* ke 10.

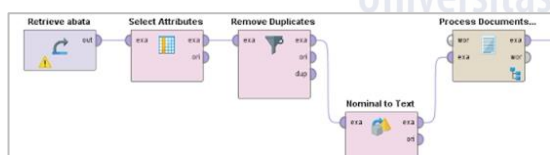
III. HASIL DAN PEMBAHASAN

A. Pengumpulan Data

Proses pengambilan data dilakukan dengan cara *crawling* menggunakan bantuan aplikasi Rstudio dan API yang didapatkan dari Twitter. Proses *crawling* data sentimen pengguna twitter diambil dari rentang waktu 1 Maret 2021 hingga 30 April 2021 yang dilakukan sebanyak tiga kali dan mendapat hasil sebanyak 12863 yaitu 6387 data yang berisi data ovo dan 6476 data dana. Struktur data awal berupa data kotor memiliki struktur yang terdiri dari 16 *field* yaitu *text*, *favorited*, *favoriteCount*, *replytoSN*, *created*, *truncated*, *replyToSID*, *id*, *replyToUID*, *statusSource*, *screenName*, *retweetCount*, *isRetweet*, *retweetes*, *longitude*, dan *latitude*. Namun dalam penelitian ini, hanya membutuhkan data *text* yaitu cuitan pengguna.

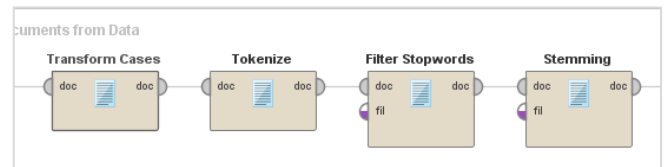
B. Preprocessing Data

Data yang awalnya berjumlah 16 *field* maka disederhanakan sesuai kebutuhan yaitu menjadi 2 yaitu sentimen pengguna dan klasifikasi pengguna. Pada tahap preprocessing ini juga dilakukan perubahan tipe data yang awalnya merupakan data *text* menjadi *numeric* atau *vector* dengan pendekatan TF/IDF. Pada tahap ini data yang didapatkan dan diproses dengan menggunakan alat bantu RapidMiner dan setelah dilakukan *preprocessing data* diperoleh data bersih sebanyak 1609 data dengan rincian 506 data tweets yang membahas ovo dan 1103 data yang membahas dana. Data yang tidak dipakai dan digunakan merupakan duplikasi data, serta cuitan pengguna yang tidak menggambarkan sentimen pengguna. Proses *preprocessing data* pada penelitian ini menggunakan dua bantuan mesin yaitu Rstudio dan RapidMiner. Proses cleansing data pada Rstudio untuk mendapatkan data bersih seperti halnya proses menghilangkan simbol, angka dan tanda baca.. Kemudian tahap *preprocessing data* dilanjutkan dengan menggunakan aplikasi RapidMiner. Gambaran proses *preprocessing data* pada RapidMiner dapat dilihat pada Gbr 3.



Gbr 3. Proses cleansing data RStudio

Pada Gbr 3. dapat dilihat bahwa data yang tadinya telah diolah terlebih dahulu pada Rstudio kemudian diolah kembali dengan bantuan RapidMiner. Pada rangkaian tersebut terdapat tahap *process document* rangkaian proses tersebut dapat dilihat pada Gbr 4.



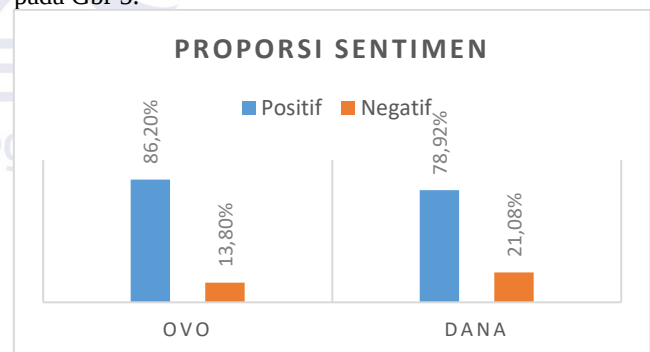
Gbr 4. Preprocessing data RapidMiner

Tahap pertama adalah melakukan *transform cases*, untuk mengubah kata menjadi *lowercase*. Selanjutnya tahap *Tokenize* atau tokenisasi yang digunakan untuk menghapus karakter yang bukan huruf. Tahap ketiga adalah *Stopwords filtering* fungsi ini digunakan untuk menghilangkan kata yang tidak terlalu diperlukan seperti kata hubung. Tahap terakhir adalah *Stemming* atau pembobotan, tahap ini dilakukan pengubahan kata imbuhan menjadi kata dasar. Proses ini dilakukan dengan menggunakan dua cara, yaitu menggunakan RapidMiner dan Rstudio. Namun karena keterbatasan kamus kata berbahasa Indonesia pada RapidMiner maka dilakukan kembali dengan Rstudio dengan memanfaatkan *package* katadasaR. Setelah mendapat data siap diolah selanjutnya adalah memberikan label. Proses *labelling* dibagi menjadi dua yaitu sentimen positif dan negatif yang nantinya diolah dengan RapidMiner yang berbasis rangkaian operator dan membandingkan performa model dalam melakukan *text mining* yang paling tepat. Hasil klasifikasi yang didapatkan dapat dilihat pada Tabel II.

TABEL II.
HASIL KLASIFIKASI SENTIMEN

Fintech	Sentimen Pengguna		Jumlah
	Positif	Negatif	
OVO	109	408	517
Dana	152	950	1102

Secara visualisasi grafik dari hasil sentiment dapat dilihat pada Gbr 5.



Gbr 5. Grafik proporsi Sentimen OVO dan Dana

Pada Gbr 5. menunjukkan dalam kurun waktu 2 bulan 1 Maret – 30 April didominasi oleh sentimen negatif yang disampaikan pengguna terhadap layanan yang diberikan oleh OVO dan Dana. Pada OVO terdapat 78,92% bersentimen negatif dan 21,08% bersentimen positif. Dari data tersebut dapat disimpulkan bahwa OVO tidak terlalu memberikan pelayanan yang baik. Sedangkan pada sentimen pengguna

Dana 82,22% pengguna melontarkan sentimen negatif dan hanya 17,78% bersentimen positif. Meskipun kedua penyedia teknologi finansial banyak mendapatkan sentimen negatif namun dari data tersebut OVO masih unggul melayani pengguna karena mendapat sentimen positif yang lebih banyak.

C. Klasifikasi Data

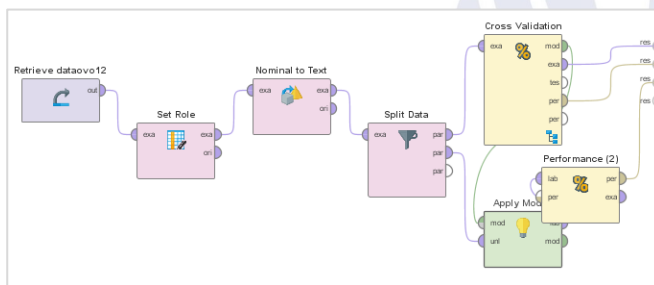
Penelitian dilakukan dengan melibatkan data teks, tahap klasifikasi data dilakukan dengan mengelompokkan data berdasarkan kelompok yang telah ditentukan. Pengelompokkan tersebut dapat digunakan untuk mengukur kualitas layanan yang diberikan OVO dan Dana. Data yang telah diberikan label nantinya dapat dilatih oleh mesin..

TABEL III.
KLASIFIKASI TWEETS BERDASARKAN SENTIMEN

Tweets	Kategori
kenapa saya tidak bisa isi ulang saldo akun ovo saya lewat mbanking bni dan gagal terus	Negatif
cashback ovo memang bikin untung	Positif

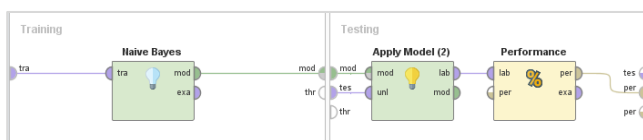
D. Evaluasi Performa Klasifikasi

Klasifikasi sentimen pengguna ini dilakukan dengan tiga model, NBC, SVM dan k-NN. Nilai performa klasifikasi didapatkan dari hasil perhitungan manual dengan menggunakan *metrics* dan juga menggunakan aplikasi RapidMiner yang dapat dilihat pada Gbr 6.



Gbr 6. Proses Performa Klasifikasi Model

Pada Gbr 6 dapat dilihat bahwa data yang ingin diolah kemudian ditentukan labelnya, kemudian dirubah tipe datanya yang awalnya berupa vector menjadi data *text*. Operator *Split Data* digunakan untuk membagi data menjadi dua jenis data yang dibutuhkan yaitu data *training* dan data *testing* kemudian dilakukan *cross validation* dari dua data tersebut, proses *cross validation* pada fold ke 10 dapat dilihat pada Gbr 7.



Gbr 7. Cross Validation

Klasifikasi pada penelitian ini dilakukan dengan *binary classification* atau dapat dikatakan proses untuk melakukan klasifikasi himpunan menjadi 2 kelompok klasifikasi, yaitu klasifikasi sentimen positif dan negatif. Adapun hasil evaluasi performa klasifikasi pada data *training* dan data data *testing* dapat dilihat pada Tabel IV.

TABEL IV
PERFORMA KLASIFIKASI SENTIMEN DATA TRAINING

Metrics	OVO			DANA		
	NB	SVM	K-NN	NB	SVM	K-NN
Akurasi	0,81	0,82	0,76	0,86	0,87	0,86
Presisi	0,90	0,91	0,60	0,43	0,43	0,43
Recall	0,54	0,56	0,57	0,50	0,50	0,50
F1	0,67	0,69	0,58	0,46	0,46	0,47

Berdasarkan pada Tabel IV. model SVM pada klasifikasi sentimen memiliki hasil yang lebih baik dibandingkan dengan dua model lainnya dengan nilai tingkat akurasi sebesar 0,82, dan pada data Dana 0,87 yang berarti prediksi pada model SVM mampu mengidentifikasi antara sentiment positif dan negatif dengan baik. *Recall* K-NN lebih besar pada OVO yang berarti model K-NN dapat melakukan prediksi yang sedikit lebih baik daripada dua model lain dalam memprediksi nilai *true value* dan *false value*. Namun dengan nilai *f-Measure* SVM lebih tinggi yaitu 0,69 dibandingkan dengan nilai NB 0,67 dan K-NN dengan nilai 0,58 dapat dikatakan secara keseluruhan model SVM memiliki tingkat akurasi lebih baik dalam melakukan klasifikasi sentimen. Hal sama didapatkan pada data *testing* model SVM unggul dalam nilai akurasi, presisi *recall* dan *f-measure* atau F1. Hasil Klasifikasi pada Data *Testing* dapat dilihat pada Tabel V.

TABEL V.
PERFORMA KLASIFIKASI SENTIMEN DATA TESTING

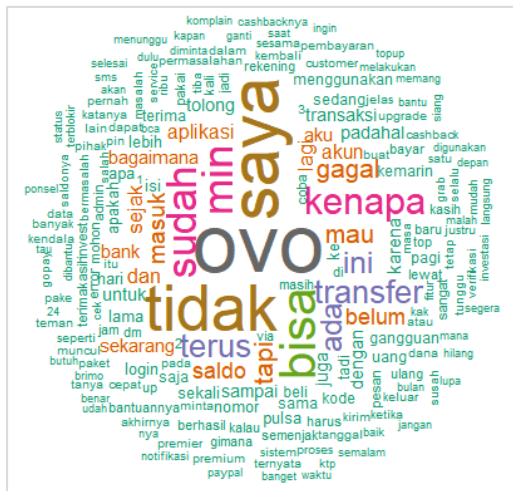
Metrics	OVO			DANA		
	NB	SVM	K-NN	NB	SVM	K-NN
Akurasi	0,79	0,80	0,79	0,86	0,86	0,85
Presisi	0,65	0,90	0,65	0,43	0,43	0,43
Recall	0,51	0,53	0,51	0,50	0,50	0,50
F1	0,57	0,66	0,57	0,46	0,46	0,46

Jika dilihat secara keseluruhan model SVM konsisten mendapatkan nilai performa klasifikasi yang tinggi dan dapat memprediksi lebih baik. Dapat disimpulkan hasil penelitian ini berbanding lurus dengan hasil penelitian yang dilakukan sebelumnya oleh Rachmadi dan Alamsyah (2017), yang juga menunjukkan bahwa model SVM dapat melakukan klasifikasi sentimen lebih baik daripada model NB dan K-NN.

E. Analisis Sentimen Teknologi Finansial

OVO dan Dana merupakan perusahaan yang bergerak pada bidang layanan teknologi finansial kepada publik dan kedua perusahaan tersebut saat ini menjadi salah satu penyedia jasa teknologi finansial terbesar di Indonesia. OVO dan Dana memanfaatkan media sosial sebagai media untuk berkomunikasi dengan pengguna mereka masing-masing. Komunikasi dan cuitan yang diunggah oleh pengguna dalam media sosial yang telah dilakukan penggalan data dan telah dilakukan pemrosesan data diperoleh hasil detail identifikasi sentimen dari opini pengguna OVO dan Dana di Twitter yang diperoleh dari hasil cuitan pengguna melalui *tweets*. Kata yang sering muncul yang disampaikan berupa keluhan terhadap layanan yang diberikan. Hasil dari kata yang sering muncul

pada data sentiment OVO yang didapatkan dari proses penggalan data dapat dilihat pada Gbr 8.



Gbr 8. Hasil Wordcloud data OVO

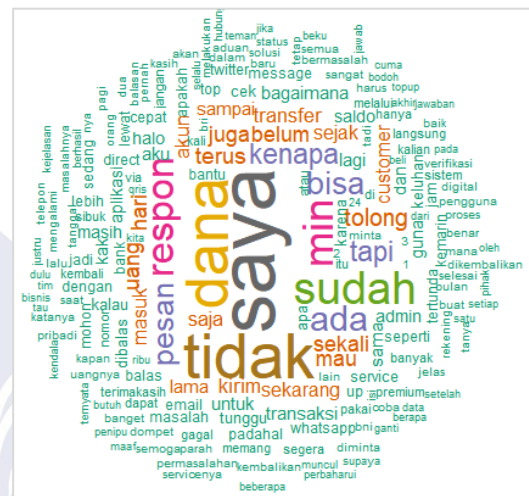
Pada Gbr. 8 merupakan kumpulan kata yang sering muncul pada data ovo yang didapatkan dari cuitan pengguna OVO di Twitter, 10 kata dengan bobot terbesar atau merupakan kata yang sering muncul adalah 'ovo', 'tidak', 'saya', 'bisa', 'sudah', 'min', 'kenapa', 'transfer', 'terus', dan 'gagal'. Kata tersebut merupakan kata yang sering sekali muncul dalam setiap kalimat yang dilontarkan oleh pengguna OVO di media sosial Twitter. Dari kata tersebut maka dapat disimpulkan bahwa sentimen pengguna terhadap layanan OVO dari hasil yang diperoleh yaitu seputar :

1. 'OVO' : merepresentasikan aplikasi ovo untuk menyampaikan kendala ataupun *feedback* yang dirasakan pengguna di Twitter;
2. 'Tidak' : merepresentasikan kegagalan ovo dalam memberikan layanan, yang dimaksud adalah dalam melakukan transaksi, transfer dan menangani keluhan pengguna;
3. 'Saya' : merepresentasikan pengguna OVO yang memberikan cuitan di Twitter;
4. 'Bisa' : merepresentasikan sebuah kemampuan dalam melakukan sesuatu pada aplikasi OVO;
5. 'Sudah' : merepresentasikan kesudahan pengguna dalam melakukan komplain kepada OVO;
6. 'Min' : merepresentasikan *customer service* yang diberikan OVO untuk menangani kendala yang dialami oleh pengguna;
7. 'Kenapa' : merepresentasikan keheranan yang diterima pengguna terhadap pelayanan OVO;
8. 'Transfer' : merepresentasikan kegiatan transfer yang dapat dilakukan atau tidak dapat dilakukan di aplikasi OVO;
9. 'Terus' : merepresentasikan kegiatan pengguna di aplikasi OVO yang dilakukan berulang kali;
10. 'Gagal' : merepresentasikan tentang gagal yang dimaksud oleh pengguna adalah gagal dalam melakukan transaksi

Dari hasil tersebut dapat menggambarkan sentimen pengguna terhadap aplikasi OVO yaitu :

1. Gambaran tentang layanan OVO yang banyak digunakan pengguna adalah transfer antar bank maupun sesama pengguna OVO;
2. Kesusahan pengguna dalam melakukan transfer pada aplikasi OVO yang sering terjadi kegagalan;
3. Kekecewaan pengguna dalam melakukan transaksi dengan OVO karena transaksi tidak berhasil namun saldo pengguna berkurang;
4. Lambatnya *customer service* OVO dalam menangani keluhan pengguna.

Selain itu data Dana juga diolah dan diketahui kata apa saja yang sering muncul dari aduan pengguna Dana di Twitter yang dapat dilihat pada Gbr 9.



Gbr 9. Hasil Wordcloud data Dana

Hasil yang didapatkan setelah diolah mendapatkan data yang hampir sama dengan data OVO, 10 kata yang paling sering muncul adalah kata 'saya', 'tidak', 'dana', 'sudah', 'min', 'respon', 'kenapa', 'pesan', 'bisa', dan 'ada'. Kata tersebut merupakan kata yang sering muncul dalam setiap cuitan pengguna yang diberikan di Twitter. Dari kata tersebut maka dapat disimpulkan bahwa sentimen pengguna yang diberikan terhadap layanan yang diberikan Dana dari hasil tersebut diperoleh seputar :

1. 'Saya' : merepresentasikan pengguna Dana yang memberikan cuitan di Twitter;
2. 'Tidak' : merepresentasikan kegagalan Dana dalam memberikan layanan, yang dimaksud adalah dalam menangani keluhan pengguna dan kemampuan aplikasi dalam melakukan transfer;
3. 'Dana' : merepresentasikan aplikasi Dana yang digunakan pengguna untuk menyampaikan kendala atau *feedback* yang dirasakan di Twitter;
4. 'Sudah' : merepresentasikan kesudahan pengguna dalam melakukan komplain pada Dana dan sudah melakukan aktivitas di aplikasi;
5. 'Min' : merepresentasikan *customer service* Dana yang menangani keluhan dan kendala yang dialami pengguna;

6. 'Respon' : merepresentasikan respon dari admin dana terhadap pengguna;
7. 'Kenapa' : merepresentasikan keheranan pengguna terhadap kegagalan atau kelambatan *customer service* dalam menangani keluhan pengguna;
8. 'Pesan' : merepresentasikan pesan yang diberikan pengguna di Twitter kepada *customer service* Dana;
9. 'Bisa' : merepresentasikan kemampuan atau ketidakmampuan aplikasi dan *customer service* dana dalam memberikan layanan kepada pengguna;
10. 'Ada' : merepresentasikan ketersediaan pesan atau respon *customer service* terhadap keluhan pengguna

Dari hasil tersebut dapat menggambarkan sentimen pengguna terhadap aplikasi Dana yaitu :

1. Gambaran tentang layanan Dana yang banyak digunakan pengguna adalah transfer antar bank maupun sesama pengguna Dana dan banyaknya keluhan pengguna yang disampaikan di Twitter;
2. Kegagalan pengguna dalam melakukan transfer ke bank atau sesama pengguna di Aplikasi Dana;
3. Lambatnya *customer service* Dana dalam menangani keluhan pengguna;
4. Kekecewaan pengguna terhadap respon dana setelah mengirim pesan berisi keluhan namun tidak pernah mendapatkan balasan;

Dapat disimpulkan dari kedua hasil pengolahan kata atau *wordcloud* yang dilakukan banyak sekali pengguna yang kecewa terhadap *Customer service* Dana dan OVO yang sangat lambat dalam menangani keluhan pengguna terhadap kendala yang diterima pengguna dalam melakukan transaksi maupun transfer pada aplikasi OVO

IV. KESIMPULAN

Kesimpulan yang dapat diambil setelah melakukan penelitian ini sebagai berikut:

1. Berdasarkan dengan nilai akurasi, presisi, *recall* dan *F-Measure* atau *F1*, model SVM dapat melakukan klasifikasi lebih baik pada klasifikasi sentimen pengguna di Twitter. Sehingga dapat dikatakan bahwa model SVM merupakan model terbaik untuk menangani teks yang bersumber dari platform media sosial yang mana hasil ini berbanding lurus dengan penelitian yang dilakukan oleh Rachmadi dan Alamsyah (2017).
2. Jika dilihat masih banyaknya keluhan dari pengguna maka didapatkan hasil proporsi sentimen dari data yang didapatkan OVO mendapatkan 78,92% sentimen negatif dan 21,08% sentimen positif, sedangkan Dana mendapatkan 86,20% sentimen negatif pengguna dan 13,80% sentimen positif. Namun, dapat disimpulkan pelayanan yang diberikan oleh OVO masih lebih baik dibandingkan dengan Dana. Meskipun hanya mendapatkan nilai 21,08% tapi nilai tersebut tidak lebih buruk dari nilai yang didapatkan oleh Dana.
3. Jika dilihat dari kata yang sering muncul banyak sekali pengguna yang mengalami kendala dalam melakukan transaksi dan juga transfer. Selain itu pengguna OVO dan Dana juga sangat kecewa terhadap pelayanan yang

diberikan oleh *customer service* OVO dan Dana karena sangat lambat dalam menangani keluhan yang diberikan kepada pengguna sehingga banyak pengguna kecewa bahkan sampai kehilangan uang pengguna di Aplikasi OVO dan Dana karena lambatnya penanganan permasalahan yang dialami dalam aplikasi.

Ketepatan sebuah model dalam memprediksi sentimen sangat perlu diperhitungkan, karena dapat membantu dalam menentukan keputusan untuk kemajuan perusahaan di masa depan. Oleh sebab itu pemilihan model yang baik perlu sekali diperhatikan. Permasalahan terbesar yang dialami OVO dan Dana terletak pada *reliability* atau keandalan dalam aplikasi yang disediakan karena tidak dapat digunakan dengan baik oleh pengguna terutama dalam melakukan transaksi, dan *responsiveness* dari *customer service* dalam menanggapi pengguna dan masih perlu ditingkatkan. Oleh sebab itu OVO dan Dana perlu melakukan inovasi terhadap aplikasi juga dalam memberikan layanan pada pelanggan. Berdasarkan analisa data teks OVO dapat melakukan inovasi lebih utamanya dalam hal transaksi dan *upgrade* akun premier pengguna. Sedangkan Dana perlu meningkatkan kemampuan aplikasi dalam melakukan transfer karena banyak dialami gangguan oleh pelanggan, selain itu admin Dana perlu lebih tanggap dalam menangani dan melakukan tindak lanjut pada keluhan pelanggan

V. SARAN

Saran yang diberikan terhadap penelitian selanjutnya dengan berdasar pada penelitian yang telah dilakukan adalah dalam proses penggalian data bisa dilakukan lebih lama dan mengalami proses beberapa kali penggalian data sedikitnya 4 kali supaya mendapatkan data yang ingin diproses dan dianalisis semakin banyak dan diharapkan mendapat hasil klasifikasi lebih baik.

REFERENSI

- [1] Alamsyah, A., & Rachmadiansyah, I. (2018, Maret). Mapping Online Transportation Service Quality and Multiclass Classification Problem Solving Priorities. *Journal of Physics: Conference Series*, 971. doi:10.1088/1742-6596/971/1/012021
- [1] We Are Social & Hootsuite. (2020, Januari 30). Digital 2020: Global Internet User Accelerates. Dipetik November 20, 2020, dari wearesocial: <https://wearesocial.com/blog/2020/01/digital-2020-global-internet-use-accelerates>
- [2] Lin, T.C. 2016. Infinite Financial Intermediation. *Wake Forest Law Review*, 27.
- [3] Jayani, D.H, (9 September 2019). Metode Pembayaran E-Commerce Melalui E-Wallet Semakin Digemari. Dipetik 23 November 2020, dari Databoks: <https://databoks.katadata.co.id/datapublish/2019/09/05/metode-pembayaran-e-commerce-melalui-e-wallet-semakin-digemari>
- [4] Gil, P. (2019, November 9). What is Twitter & How Does It Work?. Diambil kembali dari Lifewire: <https://www.lifewire.com/what-exactly-is-twitter-2483331>
- [5] Barnhart, B. (2019, Agustus 22). Twitter Customer Service. Diakses dari sproutsocial: <https://sproutsocial.com/insights/twitter-customer-service>

- [6] Kusumasondjaja, S. 2018. The Roles of Message Appeals and orientation on social media brand communication effectiveness: An evidence from Indonesia. *Asia Pacific Journal of Marketing and Logistics*, 1135- 1158.
- [7] Hearst, M. (2003, Oktober 17). What Is Text Mining? Dipetik 25 November 2019, dari [Ischool Berkeley: http://people.ischool.berkeley.edu/~hearst/text-mining.html](http://people.ischool.berkeley.edu/~hearst/text-mining.html)
- [8] Liu, B (2012). *Sentimen Analysis and Opinion Mining*. Chicago: Morgan & Claypool Publishers.
- [9] Snijders, C., Matzat, U., & Reips, U.-D. 2012. 'Big Data': Big gaps of knowledge in the field of Internet. *International Journal of Internet Science*, 7,1-5.
- [10] Suyanto. 2017. *Data Mining Untuk Klasifikasi dan Klasterifikasi Data*. Bandung: Informatika.
- [11] Alamsyah, A., & Rachmadiansyah, I. (2018, Maret). Mapping Online Transportation Service Quality and Multiclass Classification Problem Solving Priorities. *Journal of Physics: Conference Series*, 971. doi:10.1088/1742-6596/971/1/012021
- [12] Liu, B. (2012, Mei). *Sentimen Analysis and Opinion Mining*. Synthesis Lectures on Human Language Technologies, 167. doi: 10.2200/S00416ED1V01Y201204HLT016
- [13] Thakkar, H., & Patel, D. 2015. Approaches for Sentiment Analysis on Twitter: A State-of-Art study. CoRr. Diakses dari <http://arxiv.org/abs/1512.01043>
- [14] Ramos, J.E. 2003. Using TF-IDF to Determine Word Relevance in Document Queries.
- [15] Shrimali, K. R. (2018, Juli 27). SVM using Scikit-learn in Phyton. Dipetik 16 November 2020, dari LearnOpenCV: <https://learnopencv.com/svm-using-scikit-learn-in-python/>
- [16] Medhat, W., Hassan, A., & Korashy, H. 2014. Sentimen analysis algorithms and applications: A Survey. *Ain Shams Engineering Journal*, 5(4), 1093-1113. doi:10.1016/j.asej.2014.04.001
- [17] Rafaeilzadeh, P., Tang, L., & Liu, H. (2009). Cross Validation. In LIU L., OZSU, M.T/ eds *Encyclopedoa of Database Sytem*: Boston : Springer

