

**IMPLEMENTASI ANALISIS SENTIMEN PADA ULASAN
APLIKASI DUOLINGO DI GOOGLE PLAYSTORE
MENGUNAKAN METODE NAÏVE BAYES.**

SKRIPSI

Diajukan untuk Memenuhi Tugas Akhir dan
Melengkapi Syarat Guna Memperoleh Gelar Sarjana
Strata Satu (S-1) dalam Teknologi Informasi



Diajukan oleh :

MELI APRILIYANI

NIM : 2108096057

**PROGRAM STUDI TEKNOLOGI INFORMASI
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI WALISONGO
SEMARANG**

2024

PERNYATAAN KEASLIAN

Yang bertandatangan dibawah ini:

Nama : Meli Apriliyani
NIM : 2108096057
Jurusan : Teknologi Informasi Menyatakan bahwa skripsi

yang berjudul:

**Implementasi Analisis Pada Ulasan Aplikasi Duolingo di
Google Playstore Menggunakan Metode Naïve Bayes**

Secara keseluruhan adalah hasil penelitian/karya saya sendiri,
kecuali bagian tertentu yang dirujuk sumbernya.

Semarang, 31 Oktober 2024

Pembuat Pernyataan,



METERAI
TEMPEL
K40 1579609

Meli Apriliyani

NIM : 2108096057



KEMENTERIAN AGAMA REPUBLIK INDONESIA

UNIVERSITAS ISLAM NEGERI WALISONGO SEMARANG

FAKULTAS SAINS DAN TEKNOLOGI

Jalan Prof. Dr. Hamka Kampus III Ngaliyan Semarang

Telp.024-7601295 Fax.7615387 e-mail:fst@walisongo.ac.id

PENGESAHAN

Naskah skripsi berikut ini:

Judul : Implementasi Analisis Sentimen Pada Ulasan Aplikasi Duolingo Di
Google Playstore Menggunakan Metode Naive Bayes

Penulis : Meli Apriliyani

NIM : 2108096057

Jurusan : Teknologi Informasi

Telah diujikan dalam sidang tugas akhir oleh Dewan Penguji Fakultas Sains dan Teknologi UIN Walisongo dan dapat diterima sebagai salah satu syarat memperoleh gelar sarjana dalam program studi Teknologi Informasi.

Semarang, 4 November 2024

DEWAN PENGUJI

Penguji I,

Hery Mustofa, M.kom

NIP. 198703 172019031007

Penguji II,

Siti Nur'aini, S.Kom., M.Kom

NIP. 198401312018012001

Penguji III,

Dr. Masy Ari Ulinuha, M.T.

NIP. 198108122011011007

Penguji IV,

Khoirul Ikil Mustofa M.Kom.

NIP. 198808072019031010

Pembimbing I,

Siti Nur'aini, S.Kom., M.Kom

NIP. 198401312018012001

Pembimbing II,

Dr. Khoirul Umam S.T., M. Kom

NIP. 197908272011011007

NOTA PEMBIMBING

Semarang, 14 Oktober 2024

Yth. Ketua Program Studi Teknologi Informasi
Fakultas Sains dan Teknologi
UIN Walisongo Semarang

Assalamu'alaikum. Wr. Wb.

Dengan ini diberitahukan bahwa saya telah melakukan bimbingan, arahan dan koreksi naskah skripsi dengan:

Judul	: Implementasi Analisis Sentimen Pada Ulasan Aplikasi Duolingo di Google Playstore Menggunakan Metode Naive Bayes
Nama	: Meli Apriliyani
NIM	: 2108096057
Jurusan	: Teknologi Informasi

Saya memandang bahwa naskah skripsi tersebut sudah dapat diajukan kepada Fakultas Sains dan Teknologi UIN Walisongo untuk diujikan dalam Sidang Munaqosyah.

Wassalamu'alaikum. Wr. Wb.

Pembimbing I,



Siti Nur'aini M. Kom

NIP. 198401312018012001

NOTA PEMBIMBING

Semarang, 14 Oktober 2024

Yth. Ketua Program Studi Teknologi Informasi
Fakultas Sains dan Teknologi
UIN Walisongo Semarang

Assalamu'alaikum. Wr. Wb.

Dengan ini diberitahukan bahwa saya telah melakukan bimbingan, arahan dan koreksi naskah skripsi dengan:

Judul	: Implementasi Analisis Sentimen Pada Ulasan Aplikasi Duolingo di Google Playstore Menggunakan Metode Naive Bayes
Nama	: Meli Apriliyani
NIM	: 2108096057
Jurusan	: Teknologi Informasi

Saya memandang bahwa naskah skripsi tersebut sudah dapat diajukan kepada Fakultas Sains dan Teknologi UIN Walisongo untuk diujikan dalam Sidang Munaqosyah.

Wassalamu'alaikum. Wr. Wb

Pembimbing II,



Dr. Khotibul Umam, ST., M.Kom

NIP. 197908272011011007

LEMBAR PERSEMBAHAN

Dengan mengucapkan Syukur alhamdulillah, laporan tugas akhir ini dapat penulis selesaikan. Karya kecil ini penulis persembahkan untuk :

1. Alm. Bapak Saya Nur Kholiq dan ibu saya Mahfudhoh serta Bapak saya Jumiren selaku orang tua Penulis.
2. Nur Fais Muhamad Helmi, S.E selaku calon suami saya yang senantiasa menemani penulis dan memberi motivasi besar bagi penulis menyelesaikan karya ini.
3. Kakak serta adik-adikku yang senantiasa mendo'akan penulis.
4. Sahabat serta teman-teman seperjuangan khususnya jurusan Teknologi Informasi 2021.
5. Seluruh Dosen Progam Studi Teknologi Informasi.
6. Almameter Universitas Islam Negeri Walisongo Semarang.

ABSTRAK

Penelitian ini menyelidiki analisis sentimen evaluasi Aplikasi Duolingo menggunakan metode Naive Bayes. Program Duolingo mencontohkan penggunaan teknologi data besar untuk pemrosesan data yang luas dan rumit. Google Play Store menawarkan fungsi peninjauan dan pemeringkatan yang dapat membantu pengembangan program dan perbaikan aspek yang tidak diinginkan. Proyek ini menggunakan teknik analisis sentimen yang secara otomatis menganalisis ulasan produk internet Indonesia dan mendapatkan informasi mengenai perasaan yang diungkapkan dalam ulasan tersebut. Metode Naive Bayes digunakan untuk menentukan klasifikasi ulasan menjadi positif atau negatif. Temuan penelitian menunjukkan bahwa kumpulan data yang terdiri dari 1000 data yang berasal dari ulasan program Duolingo di Google Play Store diberi label secara manual sebelum ke langkah prapemrosesan. Dari jumlah tersebut, 500 data memiliki sentimen positif, sedangkan 500 data memiliki sikap negatif. Selain itu, analisis sentimen menunjukkan tingkat akurasi sebesar 86%. Skor f1 menunjukkan nilai presisi 89% dan recall 83%, dengan hasil f1 pada klasifikasi sebesar 86%.

Kata Kunci : Aplikasi Duolingo, *Naive Bayes*, Analisis Sentimen

KATA PENGANTAR

Alhamdulillah, Puji Syukur atas kehadiran Allah SWT atas segala nikmat dan karunia-Nya sehingga penulis dapat menyelesaikan Skripsi yang berjudul “Implementasi Analisis Sentimen Pada Ulasan Aplikasi Duolingo di Google Playstore Menggunakan Metode *Naive Bayes*” dengan baik. Dengan tujuan dibuat skripsi ini sebagai salah satu bentuk syarat kelulusan pada program sarjana (S1) prodi teknologi informasi di Universitas Islam Negeri Walisongo Kota Semarang. Di sisi lain, penulis juga bertujuan untuk memberikan pengetahuan kepada pembaca.

Pada kesempatan ini penulis ingin mengucapkan banyak-banyak terimakasih kepada seluruh pihak yang memberi dukungan dan bantuan dari pelaksanaan skripsi hingga penyelesaian skripsi ini. Penulis mengakui bahwa apabila tanpa dibimbing, diberi arahan, serta dibina dan tanpa diberikan motivasi dari seluruh pihak, maka penulisan skripsi ini tidak akan berjalan dengan baik. Maka dari itu penulis mengucapkan terimakasih yang tak terhingga kepada :

1. Bapak Prof. Dr. Nizar, M.Ag, selaku Rektor Universitas Islam Negeri Walisongo Semarang.
2. Bapak Prof. Dr. H. Musahadi, M.Ag, selaku Dekan

Fakultas Teknologi Informasi Universitas Islam Negeri Walisongo Semarang.

3. Bapak Dr. Khotibul Umam, ST., M.Kom, selaku ketua program studi Teknologi Informasi Universitas Islam Negeri Walisongo Semarang.
4. Ibu Siti Nur'aini M.kom dan Bapak Dr. Khotibul Umam, ST., M.Kom, selaku dosen pembimbing skripsi saya yang selalu memberikan dukungan, arahan, bimbingan serta motivasi dalam pelaksanaan skripsi hingga pembuatan skripsi ini.
5. Staff, karyawan dan dosen di lingkungan Universitas Islam Negeri Walisongo Semarang.
6. Calon Suami saya yang selalu menemani dan memberi motivasi dan semangat dalam penulisan tugas akhir ini.
7. Orang tua serta keluarga yang selalu mendo'akan serta memberikan dukungan baik kepada penulis.
8. Teman-teman Teknologi Informasi yang selalu memberidukungan.
9. Semua pihak yang tidak bisa penulis sebutkan satu persatu yang terlibat dalam pembuatan tugas akhir ini sehingga dapat terselesaikan dengan baik.

Dalam pelaksanaan dan penyusunan skripsi, penulis menyadari bahwa tentunya masih jauh dari kata sempurna dan masih banyak kekurangan. Untuk itu, penulis sangat

mengharapkan kritik serta saran yang membangun demi kesempurnaan penulisan skripsi ini, dan semoga skripsi ini dapat bermanfaat untuk semua pihak.

Semarang, 24 Oktober 2024

Penulis

DAFTAS ISI

PERNYATAAN KEASLIAN	iii
LEMBAR PENGESAHAN	v
NOTA PEMBIMBING.....	vii
NOTA PEMBIMBING.....	ix
LEMBAR PERSEMBAHAN	xi
ABSTRAK	xiii
KATA PENGANTAR.....	xv
DAFTAS ISI	xix
DAFTAR TABEL.....	xxi
DAFTAR GAMBAR	xxiii
DAFTAR LAMPIRAN.....	xxv
BAB I	1
PENDAHULUAN	1
A. Latar Belakang.....	1
B. Rumusan Masalah	6
C. Batasan Masalah.....	6
D. Tujuan Penelitian	6
E. Manfaat Penelitian.....	7
F. Sistematika Penulisan	8
BAB II.....	10
LANDASAN PUSTAKA.....	10
A. Kajian Pustaka	10
1. Analisis Sentimen	10
2. Google Playstore.....	13
3. Naïve Bayes	17
4. Aplikasi Duolingo.....	19
B. Kajian Penelitian yang Relevan	22
BAB III.....	26
METODOLOGI PENELITIAN.....	26
A. Metode Pengumpulan Data	26
B. Perangkat Penelitian.....	26
C. Alur Penelitian	27
1. Metode Pengumpulan Data.....	28
2. Pelabelan	29
3. Data <i>Preprocessing</i>	30

4. Pemisahan Data	33
5. Pembobotan <i>TF-IDF</i>	34
6. Klasifikasi <i>Naïve Bayes</i>	35
7. Evaluasi	36
BAB IV	38
HASIL DAN PEMBAHASAN	38
A. Pengambilan Data	38
B. Pelabelan	39
C. Data Preprocessing	41
D. Pemisahan Data	49
E. Pembobotan TF-IDF	49
F. Klasifikasi dan Evaluasi <i>Naïve Bayes</i>	56
BAB V	61
KESIMPULAN DAN SARAN	61
A. Kesimpulan	61
B. Saran	62
DAFTAR PUSTAKA	63
DAFTAR LAMPIRAN	66
DAFTAR RIWAYAT HIDUP	110

DAFTAR TABEL

Tabel	Judul	Halaman
Tabel 2.1	Kajian Pustaka	22
Tabel 4.1	Hasil <i>Case folding</i>	43
Tabel 4.2	Hasil <i>Stopword</i>	46
Tabel 4.3	Hasil <i>Tokenizing</i>	48
Tabel 4.4	Perhitungan TF	50
Tabel 4.5	Perhitungan DF	52
Tabel 4.6	Perhitungan IDF	53
Tabel 4.7	Perhitungan TF-IDF	54

DAFTAR GAMBAR

Gambar	Judul	Halaman
Gambar 1.1	Aplikasi Duolingo	2
Gambar 2.1	Aplikasi Google Playstore	14
Gambar 3.1	Diagram Alur Penelitian	28
Gambar 3.2	Rumus <i>Confusion Matrix</i>	37
Gambar 4.1	Data Mentah Ulasan Positif dan Negatif	38
Gambar 4.2	Hasil Output <i>Scraping</i>	39
Gambar 4.3	Hasil Data CSV <i>scraping</i>	39
Gambar 4.4	Pelabelan NLTK	40
Gambar 4.5	Hasil Pelabelan NLTK	40
Gambar 4.6	Proses Instal Library	42
Gambar 4.7	<i>Source code cleansing dan casefolding</i>	43
Gambar 4.8	<i>Source code code stopwords</i>	46
Gambar 4.9	<i>Source code tokenizing</i>	48
Gambar 4.10	Hasil Output Kalsifikasi dan evaluasi	57
Gambar 4.11	Hasil Formulasi <i>Confusion Matrix</i>	58

DAFTAR LAMPIRAN

Lampiran	Judul	Halaman
Lampiran 1	Persetujuan pembimbing Proposal	66
Lampiran 2	LOA jurnal	67
Lampiran 3	History Jurnal	68
Lampiran 4	Coding	91
Lampiran 5	Dataset	106

BAB I

PENDAHULUAN

A. Latar Belakang

Pesatnya kemajuan teknologi telah memberikan banyak manfaat bagi Masyarakat di berbagai bidang kehidupan, memudahkan tugas sehari-hari, meningkatkan produktivitas, dan memperluas akses terhadap informasi dan layanan. Kemajuan teknologi yang pesat menghasilkan data dalam jumlah besar, yang dapat menghasilkan wawasan berharga jika ditangani dan dimanfaatkan secara efektif. Seiring dengan meningkatnya jumlah pengguna internet aktif, volume data yang dihasilkan juga meningkat. Teknologi *big data* memfasilitasi analisis sejumlah besar data yang rumit dan beragam, sehingga menghasilkan pengetahuan yang berharga (Agustina et al.).

Aplikasi Duolingo mencotohkan kemampuan teknologi dalam menangani Kumpulan data yang luas dan rumit. Aplikasi Duolingo adalah program tanpa biaya yang dikembangkan oleh Luis Von Ahn dan Severin Hacker pada bulan November 2011. Tagline organisasi ini adalah “Menyediakan pendidikan bahasa gratis kepada orang-orang di seluruh dunia”. Situs web perusahaan mengklaim

bahwa mereka memiliki lebih dari 30 juta orang yang telah mendaftar (Chohan et al.).



Gambar 1. 1 Aplikasi Duolingo

Aplikasi Duolingo mendapat rating terpuji menurut statistik yang diperoleh dari Google Playstore. Mengingat banyaknya ulasan pengguna untuk aplikasi Duolingo di Google Playstore, Sulit untuk memastikan proporsi pasti dari tanggapan baik atau negative. Memanfaatkan teknologi big data, aplikasi Duolingo memiliki kemampuan untuk menganalisis Kumpulan data yang luas rumit, termasuk data penggunaan aplikasi, data penggunaan kata, dan data penggunaan fitur. Dengan menggunakan data ini, aplikasi Duolingo dapat menghasilkan wawasan berharga, seperti frekuensi penggunaan kata, pola penggunaan berbagai fitur, dan

pola penggunaan aplikasi secara keseluruhan. Dengan memanfaatkan informasi ini, pengguna dapat meningkatkan pemanfaatan program, meningkatkan pengalaman belajar Bahasa, dan mengoptimalkan pemanfaatan kemampuan aplikasi.

Playstore adalah layanan yang disediakan Google dan menyediakan berbagai macam materi digital termasuk permainan, aplikasi, film, musik, dan buku, yang disusun dalam beberapa kategori berbeda(Nurian). Play Store memiliki fungsi penilaian dan ulasan yang memungkinkan pelanggan memberikan komentar dan penilaian terhadap barang-barang yang telah mereka gunakan. Melalui penggunaan kemampuan ini pengguna dapat berkontribusi pada peningkatan aplikasi dan memperbaiki fungsionalitas yang tidak berfungsi. Hal ini memudahkan konsumen untuk memilih program yang sesuai dengan preferensi dan kebutuhan spesifik mereka. Ada lebih dari 500 juta unduhan aplikasi Duolingo yang tersedia di YouTube Play Store. Melalui pemanfaatan pendekatan analisis sentimen, tujuan penelitian ini adalah untuk melakukan investigasi mendalam terhadap program Duolingo.

Analisis sentimen adalah salah satu pendekatan yang dapat diambil untuk mengatasi masalah pengkategorian evaluasi secara otomatis menjadi lebih

positif atau lebih negative(Nurian). Analisis sentimen sering digunakan dalam pemantauan media sosial untuk mendeteksi sentimen yang ada, termasuk opini positif dan negative (Dewi).

Tujuan dari penelitian ini adalah menggunakan metode *Naïve Bayes* untuk melakukan analisis sentimen pada evaluasi evaluasi aplikasi Duolingo. Namun demikian, penting untuk digarisbawahi bahwa hanya sejumlah kecil penelitian yang melakukan analisis sentimen pada program Duolingo. Jadi, tujuan dari penelitian ini untuk memberikan penjelasan mengenai pendekatan analisis sentimen yang menggunakan algoritma *Naïve Bayes*, yaitu teknik text mining yang diimplementasikan dalam bahasa komputer Python. Pendekatan *Naïve Bayes* dipilih karena penerapannya yang luas dalam penambangan teks dan akurasinya yang luar biasa. Menerapkan pendekatan *Naïve Bayes* dapat mengurangi durasi yang diperlukan untuk kategorisasi analisis sentiment (Ardiani et al.).

Permasalahan analisis sentimen terhadap ulasan aplikasi Duolingo ini perlu diperhatikan oleh kaum muslim karena Allah SWT berfirman dalam Al-Qur'an surah Al-Hujurat ayat 6 :

يَا أَيُّهَا الَّذِينَ آمَنُوا إِن جَاءَكُمْ فَاسِقٌ بِنَبَأٍ فَتَبَيَّنُوا أَن تُصِيبُوا قَوْمًا بِجَهَالَةٍ
فَتُصْحَبُوا عَلَىٰ مَا فَعَلْتُمْ نَالِمِينَ

Artinya : “ Hai orang-orang yang beriman, jika datang kepadamu orang fasik membawa suatu berita, maka periksalah dengan teliti, agar kamu tidak menimpakan suatu musibah kepada suatu kaum tanpa mengetahui keadaannya yang menyebabkan kamu menyesal atas perbuatanmu itu” (QS Al-Hujarat 49 : Ayat 6).

Dari ayat tersebut menjelaskan bahwa jika kita memperoleh suatu opini dari seorang yang belum diketahui kejelasannya maka diharuskan melakukan *tabayyun*. Apabila seseorang Tidak melakukan *tabayyun* dapat menimpa diri sendiri dan orang lain. Menurut penjelasan ayat tersebut, penulis melakukan penelitian dengan tujuan *tabayyun*, yaitu mengumpulkan berbagai ulasan pengguna dan menganalisisnya untuk menghasilkan kesimpulan yang bermanfaat bagi masyarakat. Karena itu, dalam penelitian ini, penulis menganalisis sentimen, yaitu sentimen positif dan negatif, menggunakan algoritma naive bayes pada data yang diambil pada ulasan aplikasi Duolingo di Google Playstore.

B. Rumusan Masalah

1. Bagaimana mengimplementasikan algoritma *Naïve Bayes* dalam membantu analisis sentimen ulasan pada aplikasi Duolingo di Google Playstore?
2. Bagaimana performa yang diberikan oleh algoritma *naïve bayes* pada analisis sentimen ulasan pada aplikasi Duolingo di Google Playstore?

C. Batasan Masalah

1. Data yang digunakan adalah data dari ulasan pada aplikasi Duolingo di Google Playstore berbahasa Inggris
2. Sentimen analisis dilakukan dengan klasifikasi dari metode *naive bayes*.
3. Tanggapan pada ulasan aplikasi Duolingo di Google Playstore akan diklasifikasikan menjadi dua sentimen yaitu sentimen positif dan negatif.
4. Menggunakan Bahasa pemrograman *python* dan *software google colab (jupyter notebook versi google)*
5. Jumlah data yang digunakan adalah 1000 ulasan.

D. Tujuan Penelitian

1. Menerapkan algoritma *naive bayes* untuk menganalisis tanggapan positif dan negatif ulasan pada aplikasi Duolingo di Google Playstore.
2. Mengetahui performa dari hasil perhitungan algoritma *naive bayes* untuk klasifikasi tanggapan

positif dan negatif terhadap ulasan pada aplikasi Duolingo di Google Playstore.

E. Manfaat Penelitian

1. Manfaat teoritis :

- a. Membantu untuk mengklasifikasi *ulasan/* tanggapan positif dan negatif pada aplikasi Duolingo.
- b. Mengetahui tingkat performa pada algoritma *naive bayes* dalam melakukan klasifikasi pada ulasan aplikasi Duolingo.
- c. Sebagai pijakan, bahan referensi dan pengembangan pada penelitian-penelitian selanjutnya yang berhubungan dengan penelitian sentimen analisis.

2. Manfaat praktis :

a. Bagi Pengguna Aplikasi

Dapat mengetahui jumlah respon/tanggapan user aplikasi Duolingo yang lebih cenderung ke hal yang positif atau negative.

b. Bagi Pemilik Aplikasi

Dapat menjadi bahan evaluasi pemilik dengan mengetahui sentimen publik mengenai ulasan positif dan negative pada aplikasi Duolingo di Google Playstore.

F. Sistematika Penulisan

Laporan penelitian ini secara keseluruhan terdiri dari beberapa bab, agar memudahkan para pembaca untuk memahami isi laporan maka peneliti menunjukkan sistematika penulisan. Berikut sistematika penulisan pada penelitian ini yaitu :

BAB I PENDAHULUAN

Dalam bab pendahuluan membahas latar belakang permasalahan untuk mengetahui alasan dilakukannya penelitian ini kemudian membahas rumusan masalah penelitian, tujuan penelitian, batasan masalah, manfaat penelitian dan sistematika penyusunan laporan.

BAB II LANDASAN PUSTAKA

Dalam bab landasan pustaka membahas kajian pustaka untuk mendukung dilakukan penelitian ini serta membahas teori-teori dasar yang berkaitan dengan penelitian ini yaitu analisis sentimen dengan menggunakan metode *Naive bayes*.

BAB III METODOLOGI PENELITIAN

Dalam bab metodologi penelitian membahas tentang cara peneliti memperoleh data dan kemudian membahas perangkat yang digunakan pada penelitian, kemudian membahas alur

pengerjaan penelitian serta gambaran umum yang terdapat pada uraian metodologi.

BAB IV HASIL DAN PEMBAHASAN

Dalam bab hasil dan pembahasan ini, membahas hasil penelitian yang dapat menjawab dari analisis permasalahan pada bab III metodologi penelitian.

BAB V KESIMPULAN DAN SARAN

Dalam bab kesimpulan dan saran, membahas dua point penting yaitu yang pertama, point kesimpulan yang didapatkan dari hasil bab IV yang diuraikan secara singkat dan jelas. Kemudian, yang kedua yaitu point saran yang berisikan saran-saran dari penulis untuk penelitian selanjutnya.

BAB II

LANDASAN PUSTAKA

A. Kajian Pustaka

1. Analisis Sentimen

Analisis sentimen adalah salah satu pendekatan yang dapat diambil untuk mengatasi masalah pengkategorian evaluasi secara otomatis menjadi lebih positif atau lebih negative. Analisis sentimen sering digunakan dalam pemantauan media sosial untuk mendeteksi sentimen yang ada, termasuk opini positif dan negative (Mursyid and Indriyanti). Metode ini menggunakan pemrosesan bahasa alami (NLP) dan pembelajaran mesin untuk menganalisis dan mengekstrak informasi tentang emosi atau pendapat yang ada dalam teks. Analisis sentimen tidak hanya dapat membedakan pendapat menjadi positif atau negatif, tetapi juga dapat menemukan aspek yang lebih kompleks dari perasaan, seperti netral atau ambivalen. Analisis sentimen sering digunakan dalam berbagai bidang, seperti pemantauan media sosial untuk mengidentifikasi sentimen yang ada, baik positif, negatif, atau netral, yang disampaikan oleh pengguna terhadap topik tertentu, seperti produk,

layanan, atau masalah Sosial. Dalam dunia bisnis, Perusahaan dapat menggunakan analisis sentiment untuk memahami persepsi public terhadap merek mereka, menemukan masalah yang mungkin muncul, dan membuat strategi pemasaran yang lebih baik.

Sebaliknya, analisis ini sering digunakan dalam dunia politik untuk mengamati pendapat publik tentang kebijakan atau kandidat tertentu, yang dapat berdampak pada komunikasi dan kampanye politik.

Analisis sentimen juga digunakan oleh sektor keuangan untuk mengamati pendapat pasar tentang saham, komoditas, dan aset lainnya(Nugroho). Data sentimen ini sering digunakan oleh investor dan analis untuk memperkirakan pergerakan pasar, menemukan tren, dan membuat keputusan investasi yang lebih tepat. Sentimen publik terhadap berita ekonomi dan peristiwa internasional juga dapat menunjukkan arah pasar di masa depan. Alat analisis sentimen semakin canggih berkat kemajuan teknologi, yang memungkinkan pengolahan data dalam skala besar secara real-time. Akibatnya, pengguna menerima informasi yang lebih akurat dan cepat. Investor di bidang keuangan menggunakan analisis sentimen untuk memahami reaksi pasar atau publik terhadap berita, laporan, atau peristiwa tertentu. Misalnya,

perubahan kebijakan ekonomi global atau pengumuman laporan pendapatan perusahaan dapat memicu banyak reaksi di media sosial dan berita. Investor dapat memprediksi arah pasar atau bahkan merespons lebih cepat terhadap peluang atau risiko yang muncul dengan mengamati sentimen terhadap topik-topik ini. Mengolah dan menganalisis jumlah data yang besar, termasuk posting di media sosial, artikel berita, dan laporan ekonomi, telah mempercepat proses pengambilan keputusan investasi dengan informasi yang lebih tepat waktu dan relevan.

Selain itu, analisis sentimen dapat digunakan di bidang lain, seperti kesehatan, di mana pendapat pasien tentang layanan kesehatan atau tanggapan masyarakat terhadap kebijakan kesehatan publik dapat memberikan informasi penting. Analisis sentimen, misalnya, dapat digunakan dalam pendidikan untuk mengevaluasi feedback siswa terhadap kurikulum, guru, dan metode pengajaran. Ini dapat membantu sekolah mengubah program mereka untuk meningkatkan pengalaman belajar dan keterlibatan siswa.

Alat analisis sentimen terus berkembang seiring kemajuan teknologi, terutama dengan kecerdasan buatan dan pembelajaran mesin, yang memungkinkan pemrosesan data yang lebih cepat, akurat, dan efisien. Selain itu, analisis ini sekarang dapat menangani berbagai bahasa dan dialek, dan dapat mengidentifikasi nuansa emosional yang lebih halus dalam teks yang sebelumnya sulit ditemukan. Analisis sentimen kini digunakan pada skala yang lebih besar untuk berbagai jenis data lainnya, seperti gambar, video, dan audio. Ini memperluas cakupan dan manfaatnya. Pada akhirnya, analisis sentimen akan menjadi penting untuk pengambilan keputusan strategis di berbagai sektor industri karena penggunaan teknologi canggih seperti kecerdasan buatan dan analisis big data.

2. Google Playstore

Google play store adalah sebuah layanan konten digital yang dimiliki oleh Google yang berisi produk-produk digital seperti aplikasi, music atau lagu, buku, game maupun pemutar media berbasis cloud(Saputra). Untuk pengguna perangkat Android, layanan ini menjadi platform utama untuk mengakses berbagai produk digital, seperti aplikasi, musik, buku,

game, dan pemutar media berbasis cloud. Pengguna dapat dengan mudah mengunduh dan memperbarui aplikasi, baik gratis maupun membayar, serta mengakses berbagai konten multimedia, seperti film dan majalah, melalui Play Store. Selain itu, Google Play Store memiliki fitur ulasan dan penilaian pengguna, yang memudahkan pengunduh untuk memilih aplikasi atau konten lainnya. Layanan ini, ketika diintegrasikan dengan akun Google, menjadi lebih mudah dan aman bagi pengguna untuk menyimpan histori pembelian, aplikasi, dan data penting lainnya di cloud. Para pengembang aplikasi juga dapat memanfaatkan Google Play Store untuk mendistribusikan dan memasarkan aplikasi mereka di seluruh dunia.



Gambar 2. 2 Aplikasi Google Playstore

Selain berfungsi sebagai platform distribusi, Google Play Store juga berfungsi sebagai ekosistem yang memungkinkan pengembang dan pengguna berinteraksi satu sama lain melalui pembaruan otomatis dan fitur umpan balik. Pengembang aplikasi dapat memperbarui aplikasi mereka dengan cepat dan langsung ke pengguna di seluruh dunia, yang

memungkinkan perbaikan bug, peningkatan fitur, dan pengoptimalan kinerja secara konsisten. Play Store menawarkan berbagai cara untuk menghasilkan uang dari perspektif bisnis, termasuk langganan, pembelian aplikasi, dan iklan yang terintegrasi. Google Play Pass juga memungkinkan pengguna berlangganan untuk mendapatkan akses ke berbagai aplikasi dan game premium tanpa terpengaruh oleh iklan dan pembelian dalam aplikasi, yang meningkatkan pengalaman pengguna.

Google Play Store juga menawarkan berbagai fitur yang meningkatkan pengalaman pengguna, seperti algoritma rekomendasi aplikasi yang disesuaikan dengan preferensi pengguna berdasarkan histori pencarian dan penggunaan aplikasi. Algoritma rekomendasi ini didukung oleh teknologi pembelajaran mesin, yang meningkatkan relevansi konten yang dilihat pengguna. Selain itu, Play Store menawarkan kemudahan akses otomatis untuk memperbarui aplikasi, yang menjamin bahwa pengguna selalu mendapatkan versi terbaru yang lebih aman dan efektif dari aplikasi yang mereka gunakan. Dari segi keamanan, Google Play Store menawarkan mekanisme yang disebut Google Play Protect, yang memindai aplikasi secara teratur untuk menemukan

malware atau ancaman keamanan. Selain itu, Google Play Protect memberikan pengguna kontrol untuk mengelola izin aplikasi, yang memungkinkan mereka untuk memilih data atau fitur mana yang dapat diakses oleh aplikasi tertentu, yang memberikan tingkat privasi yang lebih tinggi.

Selain itu, Google Play Store berfungsi sebagai platform untuk penyedia konten pendidikan, menyediakan berbagai aplikasi, buku, dan kursus online untuk siswa dan karyawan. Sifat ini menjadi semakin penting seiring meningkatnya kebutuhan akan keterampilan digital dan pembelajaran jarak jauh. Play Store terus mengembangkan berbagai kategori aplikasi dan konten, memungkinkan pengguna menemukan aplikasi inovatif yang membantu manajemen finansial, kesehatan, produktivitas, dan hiburan. Google Play Console menawarkan pengembang aplikasi berbagai alat dan fitur analitik yang memungkinkan mereka memantau kinerja aplikasi, memahami perilaku pengguna, dan menemukan area yang perlu ditingkatkan. Pengembang dapat mengoptimalkan produk mereka dan meningkatkan kepuasan pengguna dengan fitur seperti uji A/B dan pelaporan crash.

3. Naïve Bayes

Naïve *Bayes* adalah salah satu algoritma pembelajaran induktif yang paling efektif dan efisien untuk *machine learning* dan data mining(Syarli and Muin). Salah satu keunggulan utama algoritma Naïve Bayes adalah kemampuan untuk menangani dataset besar dengan cepat dan efisien. Ini dapat mengatasi volume data yang besar tanpa memerlukan sumber daya komputasi yang tinggi karena algoritma ini sederhana dan memiliki kompleksitas komputasi yang rendah. Selain itu, Naïve Bayes terkenal untuk bekerja dengan data yang mengandung ketidakpastian dan kebisingan, yang membuatnya sangat cocok untuk berbagai tugas.

Aplikasi praktis lainnya dari Naïve Bayes adalah pengenalan wajah, diagnosis medis, sistem rekomendasi, dan prediksi perilaku konsumen. Namun, algoritma ini masih dapat memberikan hasil klasifikasi yang akurat, terutama ketika distribusi data mendekati distribusi normal. Algoritma ini juga bekerja dengan pembelajaran online, yang memungkinkan model untuk diperbarui.

Selain keunggulannya, Naive Bayes juga dikenal karena kemampuannya untuk menghindari masalah overfitting, yang sering menjadi masalah

dalam banyak algoritma pengajaran mesin lainnya. Naïve Bayes sangat cocok untuk situasi di mana data pelatihan sulit diperoleh atau tidak ada sama sekali. Dia mampu memberikan performa yang stabil meskipun memiliki jumlah data latih yang terbatas. Selain itu, algoritma ini dapat dimodifikasi untuk berbagai variasi. Variasi Multinomial Naïve Bayes dan Gaussian Naïve Bayes adalah dua contoh variasi yang masing-masing digunakan untuk menangani berbagai jenis data.

Selain itu, Naïve Bayes sangat populer dalam bidang pemrosesan bahasa alami (NLP) dan analisis teks. Aplikasi utamanya adalah dalam klasifikasi teks, seperti analisis sentimen, deteksi spam, dan pengkategorian dokumen. Dalam analisis sentimen, Naive Bayes menggunakan fitur tertentu, seperti frekuensi kata, untuk menentukan apakah sebuah teks mengandung sentimen positif, negatif, atau netral. Ada banyak proyek analisis data berbasis teks yang memilih algoritma ini karena keandalannya dalam menangani data teks; ini terutama berlaku untuk data yang sangat besar. Algoritma ini juga berfungsi dengan baik untuk sistem deteksi spam. Naive Bayes dapat menentukan apakah pesan email adalah spam atau tidak dengan menggunakan fitur seperti kata kunci

atau pola teks. Sangat cocok untuk aplikasi karena dapat memproses banyak pesan dengan cepat dan efisien.

Meskipun Naive Bayes memiliki banyak keunggulan, juga memiliki keterbatasan. Salah satu keterbatasan utama algoritma ini adalah asumsi independensi yang dibuat olehnya. Menurut Naïve Bayes, setiap fitur atau variabel dalam dataset tidak saling bergantung satu sama lain, tetapi banyak fitur sebenarnya saling berhubungan. Meskipun asumsi ini membuat perhitungan lebih mudah, mereka dapat mengurangi akurasi di situasi di mana fitur sebenarnya saling berhubungan secara signifikan. Meskipun demikian, asumsi independensi ini seringkali cukup dekat, terutama ketika berkaitan dengan data besar, sehingga hasil yang dihasilkan tetap akurat dan dapat diandalkan. Variasi seperti Multinomial Naïve Bayes, yang sering digunakan untuk mengklasifikasikan teks dengan fitur frekuensi kata, dan Gaussian Naïve Bayes, yang baik untuk data yang didistribusikan secara kontinu, telah dikembangkan.

4. Aplikasi Duolingo

Aplikasi Duolingo adalah salah satu aplikasi yang akan membantu kita lebih mahir dan dirasa efektif untuk mengembangkan kemampuan, berbahasa asing,

yang tersedia dalam berbagai Bahasa, termasuk Inggris. Duolingo.com/id adalah salah satu aplikasi yang dapat membantu kita lebih mahir berbahasa asing(Aisyah and Hidayatullah). Untuk membuat Duolingo lebih menyenangkan dan mudah digunakan oleh semua kalangan umur, mereka sengaja menggunakan konsep "bermain sambil belajar". Duolingo memungkinkan Anda mempelajari banyak bahasa selain bahasa Inggris, seperti Indonesia, Spanyol, Perancis, Italia, Jerman, Portugis, dan Belanda. Ini memungkinkan Anda mempelajari semua bahasa yang ditawarkannya.

Aplikasi Duolingo tidak hanya menggunakan pendekatan pembelajaran konvensional, tetapi juga menggunakan fitur gamifikasi untuk meningkatkan motivasi pengguna untuk belajar. Sistem yang menghitung poin dan level mendorong pengguna untuk terus belajar. Selain itu, latihan interaktif dan beragam yang disediakan Duolingo, seperti mendengarkan audio, mengisi kalimat yang hilang, dan menerjemahkan kalimat dari dan ke berbagai bahasa, memungkinkan pembelajaran yang lebih mendalam. Aplikasi ini juga membantu pengguna tetap belajar bahasa melalui tes harian dan pengingat otomatis.

Selain fitur-fiturnya yang disebutkan di atas, Duolingo juga menawarkan berbagai mode latihan yang dirancang untuk membantu pengguna lebih memahami bahasa yang mereka pelajari. Salah satu mode yang paling terkenal adalah latihan berbicara, di mana pengguna diminta untuk melafalkan kalimat dalam bahasa yang mereka pelajari dan menerima umpan balik tentang cara mereka mengucapkannya. Ini membantu pengguna meningkatkan kemampuan mendengarkan dan berbicara serta memahami kosa kata dan tata bahasa. Selain itu, Duolingo memiliki fitur "Stories", atau cerita pendek, yang memungkinkan pengguna membaca dan mendengarkan cerita dalam bahasa yang mereka pelajari sambil menjawab pertanyaan yang didasarkan pada cerita tersebut. Fitur ini memberikan konteks nyata penggunaan bahasa, membantu pengguna memahami bagaimana bahasa digunakan dalam percakapan sehari-hari, dan meningkatkan pemahaman pengguna tentang membaca dan analisis.

Dengan adanya forum diskusi di mana pengguna dapat bertanya, berbagi pengalaman, atau berbicara tentang apa yang mereka pelajari, aplikasi ini membantu pembelajaran komunitas dan meningkatkan motivasi pengguna untuk belajar

bahasa. Dengan berbagai fiturnya, Duolingo telah berhasil membuat lingkungan pembelajaran yang interaktif, menyenangkan, dan universal.

B. Kajian Penelitian yang Relevan

Pada bagian kajian Pustaka ini, akan membahas penelitian-penelitian terkait sebelumnya dengan tujuan untuk memperkuat pelaksanaan penelitian ini.

Table 2.1 kajian pustaka

NO	Penulis	Judul	Hasil
1	(Birjali 2021)	<i>App Review in Google Play Store Using Naïve Bayes Algorithm</i>	Penelitian ini menemukan bahwa naïve bayes dapat membantu memecahkan masalah pengklasifikasian sentiment pada ulasan aplikasi Shopee dengan akurasi yang tinggi sebesar 96,667%
2	(Khoirul Insan 2020)	Analisis Sentimen Aplikasi	Penelitian ini tidak menjelaskan kecenderungan

		BRIMO pada Ulasan Pengguna di <i>Google Play</i> Menggunak an <i>Algoritma Naïve Bayes</i>	sentiment pada pengguna terhadap topik yang dibahas dengan menghasilkan akurasi yang cukup tinggi yaitu 84,25%
3	(Muham mad Zidan 2022)	Analisis sentiment kenaikan harga bahan bakar minyak (BBM) Berdasarka n respon pengguna media social twitter di	Pada penelitian ini menganalisis sentimen kenaikan harga BBM dengan dataset 1483 tweet yang diambil dari Twitter kemudian telah melalui tahap preprocessing dan mengklasifikasika nnya

		Indonesia menggunakan metode Naïve bayes.	menjadi 3 yaitu positif, netral dan negatif dengan perbandingan data latih dan data uji 80:20 yang diambil acak. Hasil analisis sentimen dengan metode Naï ve Bayes didapatkan nilai performa dengan tingkat akurasi sebesar 81%, presicion sebesar 83%, recall sebesar 81% dan f1 score sebesar 79%. Serta didapatkan sentimen negatif memiliki nilai
--	--	---	--

			persentase tertinggi.
--	--	--	-----------------------

Ketiga penelitian diatas menganalisis aplikasi perbankan atau e- commerce, sedangkan penelitian saya membahas aplikasi edukasi yaitu Duolingo yang konteks ulasannya berbeda, terutama dalam hal pengalaman pengguna dalam belajar Bahasa.

Penelitian terdahulu menggunakan metode sama yaitu Naïve Bayes, namun penelitian saya menambah dimensi baru dengan melakukan scraping data ulasan dari Google playstore dan labeling menggunakan Library NLTK.

BAB III

METODOLOGI PENELITIAN

Penelitian ini fokus pada pemeriksaan penilaian pengguna terhadap program Duolingo di Google Playstore. Dengan menggunakan *Algoritma Naïve Bayes* untuk memprediksi suatu ulasan apakah positif atau negatif, dengan kemampuan yang dapat mengatasi data dengan dimensi yang tinggi (Birjali et al.). *Algoritma Naïve Bayes* dipilih karena sangat akurat dan banyak digunakan dalam proses text mining (Wickramasinghe and Kalutarage). Selain itu, *Algoritma Naïve Bayes* adalah pengklasifikasi probabilistic yang paling sederhana dan mudah digunakan (Kaviani and Dhotre). Metode penelitian ini terdiri dari tujuh tahapan penting: pengumpulan data, anotasi manual, pemrosesan awal data, pembobotan TF-IDF, partisi data ke dalam set pelatihan dan pengujian, klasifikasi menggunakan *algoritma naïve bayes*. Dengan mengikuti proses ini secara sistematis, penelitian ini dapat menghasilkan hasil yang reliabel dan bermanfaat dalam memahami persepsi pengguna terhadap Aplikasi Duolingo.

A. Metode Pengumpulan Data

Adapun Teknik untuk mengumpulkan data adalah sebagai berikut :

a. Pengamatan (Observasi)

Observasi yaitu metode pengumpulan data dengan cara mengadakan tinjauan secara langsung ke objek yang diteliti. Data primer yang diperoleh penulis pada metode ini didapatkan dengan cara scraping pada ulasan aplikasi Duolingo yang terdapat pada Google Playstore.

b. Studi Pustaka

Untuk mendapatkan data sekunder yang bersifat teoritis, penulis mengumpulkan materi dengan menggunakan buku, jurnal, skripsi, website, dan sejenisnya yang berkaitan dengan masalah yang diteliti oleh penulis, baik teori dasar untuk analisis maupun dokumen laporan yang relevan.

B. Perangkat Penelitian

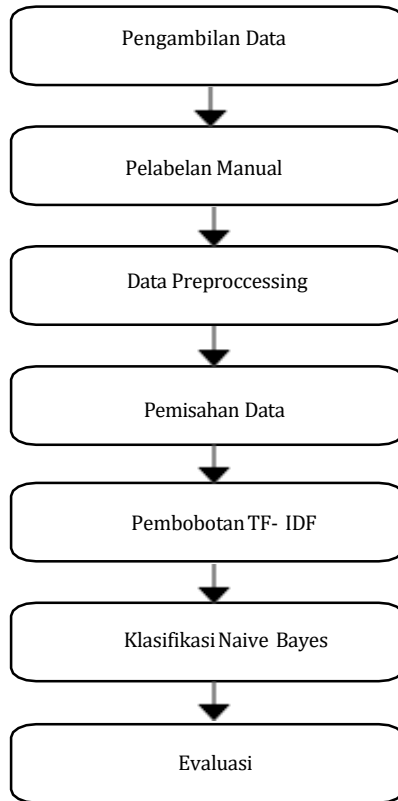
Spesifikasi perangkat yang digunakan pada penelitian ini membutuhkan beberapa perangkat keras dan perangkat lunak yang harus dipenuhi sebagai alat pendukung dalam berjalannya proses system. Beberapa spesifikasi kebutuhan yang diuraikan dalam proses system ini sebagai berikut :

1. Spesifikasi kebutuhan perangkat keras (*Hardware*)

- a. Device : Laptop Hp
 - b. Processor : Intel® Core™ i7-8550U
 - c. Memori (RAM) : 8192MB
 - d. Monitor : 14 Inch
2. Spesifikasi kebutuhan perangkat lunak (*software*)
- a. Sistem operasi : Windows 10
 - b. Bahasa pemrograman : python
 - c. Ms. Office : Ms. Word, Ms.Exel 2019
 - d. Google drive : Google Colab

C. Alur Penelitian

Adapun tahapan pengerjaan penelitian yang membahas mengenai gambaran umum alur penelitian menggunakan metode Naïve Bayes yang dilakukan penulis, bisa di lihat pada gambar 3.1 berikut :



Gambar 3. 1 Diagram Alur Penelitian

1. Metode Pengumpulan Data

Pengumpulan data dilakukan dengan menggunakan bahasa pemrograman Python dan alat scraping Google Play Store. Data yang diambil terdiri dari 1000 ulasan, yang mencakup ulasan dengan bintang dari 1 hingga 5, dan diambil secara otomatis melalui teknik scraping. Teknik ini

digunakan untuk mengambil data dari aplikasi Duolingo di Google Play Store. Data yang dikumpulkan kemudian disimpan dalam database yang terstruktur untuk memudahkan pengolahan lebih lanjut untuk analisis sentimen, evaluasi performa aplikasi, dan kebutuhan analitik lainnya. Data yang dikumpulkan termasuk berbagai informasi penting, seperti nilai bintang, teks ulasan, dan tanggal ulasan. Scraping ini dapat membantu menemukan pola dalam ulasan pengguna, meningkatkan kualitas aplikasi, dan memahami tren pengguna.

2. Pelabelan

Setelah data tinjauan diperoleh, library Natural Language Toolkit (NLTK) digunakan untuk memberikan label positif atau negatif. Pola kata, pemrosesan bahasa alami, dan polaritas ulasan dibantu oleh perpustakaan ini. Setelah pelabelan selesai, data ulasan dikelompokkan menjadi dua kategori berbeda: kategori positif terdiri dari ulasan dengan nilai Bintang 4 atau 5, dan kategori negatif terdiri dari ulasan dengan nilai Bintang 1, 2, atau 3. Analisis sentimen dapat dilakukan dengan lebih efisien dan akurat karena pembagian ini memudahkan proses analisis sentimen berikutnya,

memudahkan klasifikasi sentimen, dan memungkinkan fokus pada ulasan yang paling relevan dalam masing-masing kategori sentimen. Klasifikasi yang lebih jelas ini memudahkan proses analisis dan memungkinkan untuk menemukan masalah dan pendapat pengguna dengan lebih tepat. Akibatnya, analisis sentimen dapat dilakukan dengan lebih efisien dan akurat.

3. Data *Preprocessing*

Untuk memulai analisis sentimen, data harus diproses terlebih dahulu. Tujuan dari langkah ini adalah untuk mempersiapkan data mentah sebelum dianalisis untuk memastikan bahwa mereka dalam kondisi yang tepat dan relevan untuk analisis. Pra-pemrosesan mencakup sejumlah proses penting, seperti pembersihan teks, penghapusan karakter atau simbol tertentu, penghapusan stop words (kata-kata umum yang tidak mengandung sentimen), dan stemming atau lemmatization, yang mengubah kata-kata ke bentuk aslinya. Selain itu, data yang lebih relevan dipilih untuk meningkatkan struktur data yang akan digunakan dalam analisis. (Rakhmawati et al.). Selain itu, langkah ini melibatkan normalisasi data, yang mencakup mengubah teks menjadi huruf yang lebih kecil dan

menghilangkan spasi berlebih, sehingga hasil analisis lebih akurat. Dengan pra-pemrosesan yang baik, data menjadi lebih bersih, terstruktur, dan siap untuk dianalisis, yang memungkinkan hasil analisis sentimen yang lebih efektif dan terpercaya.

Prosedur penyiapan data terdiri dari empat tahapan yaitu *casefolding*, *tokenizing*, *stopwords*, dan *stemming*.

a. *Case folding*

Merupakan tahap pengolahan data atau mengkonversi semua huruf menjadi huruf kecil(Lestari and Saepudin). Proses ini mengurangi kompleksitas data dan mencegah pengulangan kata yang sama namun ditulis dalam format yang berbeda. Ini juga membantu menjaga teks konsisten sehingga tidak ada perbedaan yang tidak perlu antara kata yang ditulis dalam huruf besar atau kecil. Karena model analisis teks seringkali tidak membedakan huruf kapital, normalisasi huruf ini juga merupakan langkah penting dalam pra-pemrosesan data teks. Dengan menjaga data konsisten, hasil analisis yang dihasilkan pada tahap selanjutnya lebih akurat.

b. Stopwards

Adalah proses menghilangkan kata-kata yang tidak bermakna dari teks yang sedang diproses (Imron). Kata-kata seperti "saya", "dan", "atau", "di", dan sebagainya adalah kata-kata umum yang tidak memberikan nilai tambahan untuk analisis sentimen. Stopwords ini biasanya tidak memberikan makna khusus dalam teks dan, karena frekuensinya yang tinggi, dapat mengganggu proses analisis. Namun, mereka tidak relevan dengan sentimen yang ingin diukur. Untuk meningkatkan kualitas dan akurasi hasil analisis, penghapusan kata stopwords membantu fokus pada kata kunci yang lebih relevan untuk menentukan sentimen. Proses ini sangat penting karena mengurangi kebisingan dalam data. Ini juga memungkinkan model pembelajaran mesin atau metode analitik lainnya untuk lebih mudah menangkap pola-pola penting dalam teks.

c. Tokenizing

Adalah di mana kalimat dipecah menjadi beberapa kata yang disebut sebagai token yang memisahkannya (Sanjaya and Lhaksmana). Proses ini dikenal sebagai tokenizing, yaitu

Memecah teks menjadi bagian yang lebih kecil, seperti kata, frasa, atau bahkan kalimat, adalah proses yang disebut tokenization. Melakukannya membuatnya lebih mudah untuk menganalisis kemudian karena setiap bagian kecil, atau token, dapat dianalisis secara terpisah. Dengan melakukan ini, data teks dapat diubah menjadi format yang lebih terstruktur, yang memungkinkan analisis lebih mendalam, seperti menemukan pola, menghitung frekuensi kata, atau mengidentifikasi penen.

4. Pemisahan Data

Untuk tujuan perbandingan, data pra-pemrosesan dibagi menjadi kelompok tertentu, yang dikenal sebagai pemisahan data. Meskipun tujuan utamanya adalah mendapatkan kumpulan data yang cukup representatif untuk pelatihan model, kumpulan data pengujian digunakan untuk mengevaluasi kinerja model yang telah dilatih. Dengan melakukan pembagian ini, kinerja model dapat diuji pada data baru yang sebelumnya tidak dilihat oleh model selama pelatihan, yang membantu menilai seberapa baik model dapat menggeneralisasi atau bekerja dengan baik pada data baru. Untuk menghindari overfitting atau

underfitting dan untuk memastikan bahwa model dapat menghasilkan hasil yang memuaskan dalam situasi dunia nyata, pemisahan data yang tepat antara data pelatihan dan pengujian sangat penting. Selain itu, pemisahan data yang proporsional dan representatif membantu mengukur secara objektif akurasi, presisi, dan kinerja keseluruhan model sebelum diterapkan pada data yang sebenarnya.

5. Pembobotan *TF-IDF*

Prosedur pembobotan kata digunakan untuk mengubah data yang telah diproses sebelumnya, yang mungkin berupa kata-kata, menjadi representasi numerik. Metode pembobotan kata digunakan untuk menentukan seberapa besar kepentingan kata-kata yang digunakan dalam penelitian. Bobot suatu istilah dalam teks olahan berkorelasi langsung dengan frekuensi istilah tersebut. Tingkat frekuensi kata-balik dokumen (TF-IDF), yang memperhitungkan frekuensi kata yang muncul di seluruh korpus dan dalam dokumen tertentu, adalah salah satu teknik pembobotan yang paling umum digunakan. Jumlah kali kata tertentu muncul dalam dokumen tertentu tetapi jarang muncul di dokumen lain menunjukkan betapa pentingnya kata tersebut. Metode ini membantu

menemukan kata-kata yang lebih penting dalam konteks tertentu, sehingga analisis menjadi lebih efektif dan berkonsentrasi pada aspek yang lebih penting.

Akibatnya, Ketika jumlah dokumen yang diproses bertambah, semakin banyak pula karakteristik yang dihasilkan. Hal ini memungkinkan pemeriksaan yang lebih komprehensif dan berpotensi meningkatkan kaliber dan presisi temuan analitis yang dilakukan. Perhitungan bobot kata pada pendekatan “TF-IDF merupakan hasil penggabungan dua gagasan yaitu Term Frekuensi (TF) dan Inverse Document Frekuensi (IDF)”(Septian et al.). Salah satu definisi TF adalah berapa kali setiap kata muncul dalam dokumen tertentu. Nilai TF meningkat sebanding dengan jumlah kata yang ada disetiap dokumen dan sering muncul.

6. Klasifikasi *Naïve Bayes*

Proses klasifikasi merupakan suatu metode penambahan data yang dapat digunakan untuk membangun suatu model hingga semua data termasuk dalam kategori yang sama(Tangkelayuk). Teorema Bayes berfungsi sebagai premis fundamental yang digunakan dalam pendekatan

klasifikasi yang digunakan dalam algoritma Naïve Bayes. Teorema ini adalah yang pertama. Seorang ilmuwan Inggris. Akibatnya, ketika jumlah dokumen yang diproses bertambah, semakin banyak pula karakteristik yang dihasilkan. Hal ini memungkinkan. Pemeriksaan yang lebih komprehensif dan berpotensi meningkatkan kaliber dan presisi temuan analitis yang dilakukan. Perhitungan bobot kata pada pendekatan “TF-IDF merupakan hasil penggabungan dua gagasan yaitu Term Frekuensi (TF) dan Inverse Document Frekuensi (IDF)”(Syawal et al.). Salah satu definisi TF adalah berapa kali setiap kata muncul dalam dokumen tertentu. Nilai TF meningkat sebanding dengan jumlah kata yang ada disetiap dokumen dan sering muncul.

7. Evaluasi

Evaluasi merupakan fase yang mengukur efektivitas algoritma kategorisasi yang digunakan dalam penelitian(Prabowo and Kurniadi). Parameter seperti akurasi, presisi, recall, dan perolehan digunakan untuk mengevaluasi kinerja sistem. Langkah ini sangat penting untuk menentukan kemampuan model yang dilatih untuk mengklasifikasikan data dengan benar. Untuk

melakukan komputasi menggunakan benchmark ini, penting untuk menggunakan metodologi Confusion Matrix, yang terdiri dari komponen True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Komponen-komponen ini memberikan gambaran detail tentang prediksi model, termasuk di mana model melakukan kesalahan, seperti memprediksi positif saat sebenarnya negatif (FP) atau sebaliknya (FN).

Adapun Rumus Confusion Matrix untuk menghitung akurasi, presisi, dan recall adalah sebagai berikut:

$$\text{Accuracy} = \frac{TP + TN}{TP + FP + TN + FN}$$

$$\text{Precision} = \frac{TP}{TP + FP}$$

$$\text{Recall} = \frac{TP}{TP + FN}$$

Gambar 3. 2 Rumus confusion matrix

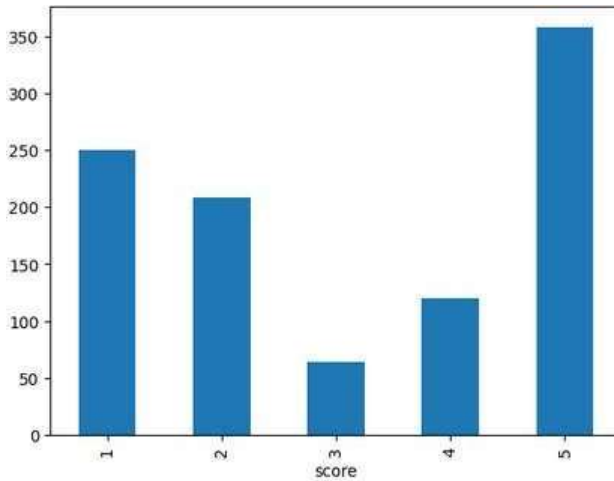
BAB IV

HASIL DAN PEMBAHASAN

Pada tahap ini akan menjelaskan tahapan yang berkaitan dengan prosedur yang dilakukan dalam penelitian, sebagaimana dituangkan dalam metodologi penelitian sebelumnya.

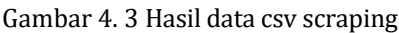
A. Pengambilan Data

Saat ini pengambilan data dengan memanfaatkan URL program Duolingo dari Playstore dimulai bulan Maret sampai bulan Mei 2024. Berikut ini disajikan hasil ekstraksi data yang dilakukan dengan menggunakan kata kunci “Duolingo” dan besarnya pencarian data dari 1000 kemunculan aplikasi Duolingo:



Gambar 4. 1 Data Mentah Ulasan Positif dan Negatif

Gambar 4. 2 Hasil output scraping



Setelah melakukan ekstraksi data, hasilnya menunjukkan bahwa aplikasi Duolingo memiliki beragam ulasan dari pengguna yang berbeda. Dari ulasan yang yang diberi Bintang 1, 2, 3, 4 dan 5.

B. Pelabelan

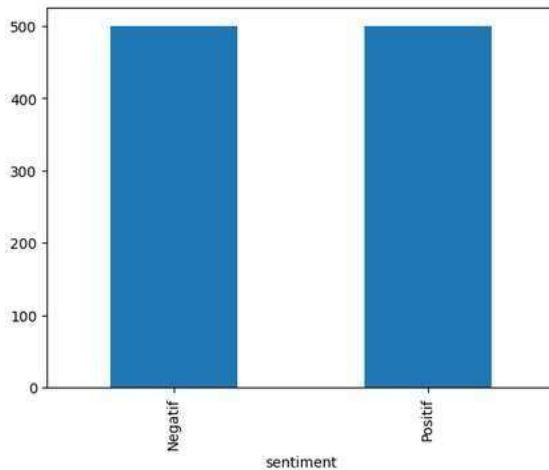
Setelah mengumpulkan data penelitian, dilanjutkan dengan melakukan pelabelan yang akan menggunakan library NLTK, NLTK (*Natural Language Toolkit*) adalah sebuah perpustakaan untuk pemrosesan bahasa alami

python yang dikembangkan oleh Stevan Bird dan Edward Loper. Berikut adalah hasil dari prosedur Pelabelan yang telah dilakukan :

```
import nltk
from textblob import TextBlob
all_polarity = 0;
status=[]
total_positif = total_negatif = total = 0

for ulasan in data['content']:
    analysis = TextBlob(ulasan)
    all_polarity += analysis.polarity
    if (analysis.sentiment.polarity > 0.5):
        total_positif += 1
        status.append("Positif")
    elif (analysis.sentiment.polarity <= 0.5):
        total_negatif += 1
        status.append("Negatif")
    total += 1
print(f"Hasil Analisis Data: \nPositif = {total_positif} \nNegatif = {total_negatif}")
print(f"Total Data = {total}")
```

Gambar 4. 4 Pelabelan NLTK



Gambar 4. 5 Hasil pelabelan NLTK

Hasil dari prosedur pelabelan NLTK ini

memberikan wawasan penting tentang cara pengguna memandang platform pembelajaran bahasa online Duolingo. Dengan membedakan ulasan menjadi kategori positif dan negatif. Pada analisis pada data ini menghasilkan analisis data positif sebanyak 500 dan data negative sebanyak 500.

C. *Data Preprocessing*

Sebelum menggunakan algoritma klasifikasi pada dokumen, tahap preprocessing harus diselesaikan. Data teks dibersihkan dan disiapkan untuk digunakan dalam analisis atau pemodelan lebih lanjut dalam tahapan pra-pemrosesan. Pada titik ini, berbagai proses dilakukan, seperti tokenisasi, penghilangan stopwords, stemming, dan lemmatisasi. Semua proses ini sangat penting untuk memastikan bahwa informasi penting tetap ada sementara data yang tidak diperlukan atau dapat mengganggu hasil analisis dibuang. Sementara penghilangan stopwords bertujuan untuk menghilangkan kata-kata umum seperti "dan", "di", atau "yang", proses tokenisasi membagi teks menjadi bagian yang lebih kecil. Tahapan pra-pemrosesan adalah sebagai berikut:

```
[ ] # import library
import pandas as pd
import string
import nltk
```

Gambar 4. 6 proses install library

a) *Case Folding*

Lebih lanjut, setiap huruf diubah menjadi huruf kecil untuk memastikan bahwa teks yang dihasilkan secara keseluruhan konsisten. Karena normalisasi ini menghilangkan perbedaan huruf besar dan kecil, kata-kata yang seharusnya sama tetap dianggap sebagai satu entitas. Hal ini sangat penting untuk pemrosesan teks, terutama ketika sistem perlu mengenali kata-kata yang memiliki arti yang sama dalam berbagai bentuk penulisan. Selain mengubah huruf menjadi lebih kecil, proses ini juga menghilangkan kata-kata yang tidak penting dari tanda baca atau karakter tertentu, seperti angka, simbol khusus, atau karakter yang tidak relevan secara semantik untuk analisis. Proses yang disebut case folding membuat teks lebih bersih dan seragam, yang mempermudah analisis dan meningkatkan kinerja algoritma yang digunakan. Berikut ini adalah proses *casefolding*

terdapat pada gambar 4.5 dan Tabel hasil saat menggunakan *casefolding* :

```
[ ] # Buat fungsi case folding
def clean_content(content):
    import string, re

    content = content.lower() # menjadikan lowercase
    content = re.sub("[^a-z]", "", content) # hapus semua karakter kecuali a-z
    content = re.sub("\t", " ", content) # mengganti tab dengan spasi
    content = re.sub("\n", " ", content) # mengganti new line dengan spasi
    content = re.sub("\s+", " ", content) # mengganti spasi > 1 dengan 1 spasi
    content = content.strip() # menghapus spasi di awal dan akhir

    return content

[ ] # case folding
data["content_clean"] = data["content"].apply(clean_content)
data.head()
```

Gambar 4.7 source code cleansing dan casefolding

Tabel 4. 1 Hasil Case Folding

Score	Content	At	Content Clean
3	I’his is the one. Ive tried several apps and pr...	2024-05-13 03:46:51	this is the one ive tried several apps and pro...
4	It has a decent variety of exercises. It does...	2024-05-13 10:02:02	it has a decent variety of exercises it does n...
3	Most impoitant it’s free, you can upgiade but...	2024-05-09 19:21:19	most impoitant it’s free, you can upgiade but i...

2	I really like the setup and the levels. I felt...	2024-05-04 11:01:20	i really like the setup and the levels. I felt...
	Fantastic app designed to keep you engaged wit...	2024-05-10 03:38:26	fantastic app designed to keep you engaged wit...

Hasil dari *case folding*, Semua huruf menjadi huruf kecil karena case folding. Ini menciptakan keseragaman dalam teks, yang membuat analisis lebih mudah. Dengan adanya keseragaman ini, algoritma yang menganalisis teks tidak perlu membedakan huruf besar dan kecil—yang pada dasarnya merujuk pada kata yang sama—sehingga mereka dapat melakukan analisis lebih lanjut. Selain itu, ini meningkatkan kemudahan pemrosesan teks dan mencegah kesalahan interpretasi yang dapat terjadi karena perbedaan penulisan. Misalnya, setelah case folding, kata "data" dan "data" akan dianggap sebagai satu entitas yang sama, yang meningkatkan akurasi hasil analisis. Keanekaragaman ini sangat penting untuk proses klasifikasi dokumen dan analisis teks lainnya, karena data yang lebih bersih dan

konsisten akan menghasilkan pemodelan yang lebih efisien dan akurat.

b) Stopwords

Stopwords adalah tahap berikutnya, yang merupakan bagian penting dari pra-pemrosesan teks. Dalam analisis, stopwords adalah kata-kata umum yang sering muncul dalam teks tetapi tidak signifikan. Tahap ini bertujuan untuk menghilangkan kata-kata yang tidak penting dan tidak bermakna dari daftar kata yang ditemukan dalam database stopwords, seperti "dan", "aku", "dari", "itu", "adalah", dan sebagainya. Menghapus stopwords sangat penting karena mereka tidak membantu memahami konteks atau tujuan utama analisis. Dengan menghilangkan kata-kata tersebut, teks menjadi lebih tajam dan informasi yang relevan dapat diekstraksi dengan lebih mudah.

Tahap berikutnya, stopwords, merupakan bagian penting dari pra-pemrosesan teks. Tujuan dari tahap ini adalah untuk menghilangkan kata-kata tidak penting dan tidak penting dari daftar kata yang ditemukan dalam database stopwords, seperti "dan", "aku", "dari", "itu", "adalah", dan sebagainya. Sangat penting untuk menghilangkan

stopwords karena mereka mempersulit pemahaman konteks atau tujuan utama analisis. Teks menjadi lebih jelas dengan menghilangkan kata-kata tersebut, dan informasi yang relevan dapat diekstraksi dengan lebih mudah. Berikut proses *stopword* dan Tabel hasil dari *stopword* adalah sebagai berikut :

```
# stopwords
import nltk.corpus
nltk.download('stopwords')
from nltk.corpus import stopwords
stop = stopwords.words('indonesian')
data['text_stopword'] = data['content_clean'].apply
(lambda x: ' '.join([word for word in x.split() if word not in (stop)]))
data.head(50)
```

Gambar 4.8 source code *stopword*

Tabel 4. 2 Hasil Stopward

<i>Content Clean</i>	<i>Text Stopwords</i>
This is the one ive tried several apps and pro..	this is the one ive tíied sevefál apps and píó...
it has a decent vafiety of exeícisesit does n...	it has a decent vafiety of exeícises itdoes n...
most impoítant it s fee you canupgíade but i...	most impoítant it s fee you can upgíade but i...
i feally like the setup and the levels ifelt ...	i feally like the

	setup and the levels ifelt ...
fantastic app designed to keep youengaged wit...	fantastic app designed to keep youengaged wit...

Penekanan pada kata-kata kunci yang lebih relevan dalam menentukan sentimen difasilitasi oleh hasil-hasil ini. Dengan demikian, proses ini meningkatkan kualitas analisis sentimen dengan menghilangkan informasi yang tidak relevan, yang memastikan hasilnya lebih akurat. Algoritma dapat berkonsentrasi pada kata-kata yang memiliki nilai sentimen, seperti kata-kata yang menunjukkan pendapat, emosi, atau evaluasi, dengan menghilangkan kata-kata yang tidak bermakna, seperti stopwords. Sebagai contoh, kata-kata seperti "bagus", "buruk", "suka", atau "benci" seringkali berfungsi sebagai petunjuk penting untuk menentukan apakah sebuah teks mengandung sentuhan positif atau negatif. Algoritma dapat lebih mudah menemukan pola emosi dan sikap dalam teks dengan menghilangkan kata-kata yang tidak membantu penilaian sentimen.

c) *Tokenizing*

Fase selanjutnya adalah *tokenizing*, yang melibatkan pembagian teks menjadi kata-kata individual untuk memfasilitasi proses stemming selanjutnya. Padatahap ini teks ulasan telah dibagi menjadi kata-kata individual, yang memungkinkan untuk melakukan analisis lebih lanjut secara terpisah pada setiapkata. Berikut proses dan hasil *tekonizing* :

```
# tokenizing
import nltk
nltk.download('punkt')
from nltk.tokenize import sent_tokenize, word_tokenize
data['text_tokens'] = data['text_StopWord'].apply(lambda x: word_tokenize(x))
data.head()
```

Gambar 4.9 source code tokenizing

Tabel 4. 3 Hasil Tokenizing

Text Stopward	Text Tokens
this is the one ive tried several appsand pio...	[this, is, the, one, ive, tried, several, apps...
it has a decent variety of exercisesit does n...	[it, has, a, decent, variety, of,exercises, i...
most important it s free you canupgrade but i...	[most, important, it, s, free, you,can, upgra..
i really like the setup and the levels ifelt ...	[i, really, like, the, setup, and, the,

	levels...
fantastic app designed to keep youengaged wit...	[fantastic, app, designed, to, keep, you, enga...

Dengan melalui tahap *tokenizing*, data siap untuk diproses lebih lanjut dalam analisis sentimen menggunakan algoritma klasifikasi.

D. Pemisahan Data

Saat ini, data dibagi menjadi dua kategori berbeda: data pengujian dan data pelatihan. Data dipartisi menjadi 20% dataset pengujian dan 80% dataset pelatihan dalam penelitian ini. Dengan membagi dataset menjadi 20% untuk pengujian dan 80% untuk pelatihan, penelitian ini memastikan bahwa model yang dibangun dapat diuji dengan baik pada data yang tidak terlihat sebelumnya, sementara masih memiliki sejumlah besar data untuk melatih model dengan baik. Hal ini memungkinkan untuk mendapatkan evaluasi yang lebih obyektif terhadap kinerja model dalam memprediksi sentimen ulasan aplikasi Duolingo.

E. Pembobotan TF-IDF

Saat ini, pendekatan pembobotan TF-IDF (Term Frequency-Invers Document Frequency) digunakan untuk mengubah data tekstual ke dalam format numerik. Kumpulan data memberikan bobot pada setiap kata berdasarkan frekuensinya dalam teks (TF) dan

kelangkaannya di seluruh dokumen dalam kumpulan data (IDF). Dengan menggunakan pembobotan TF-IDF, diperoleh numerik yang lebih akurat dari teks ulasan, yang dapat digunakan oleh algoritma klasifikasi untuk analisis sentimen.

Sebagai contoh penelitian menggunakan tiga ulasan untuk perhitungan manualisasi untuk pembobotan TFIDF sebagai berikut :

(Doc1) = “ *it’s pretty good, but they keep incentivizing*”

(Doc2) = “*I just love it everthing is good in it*”

(Doc3) = “*I am actually learning and speaking Italian*”

Pada metode TFIDF terdapat dua kata yaitu TF (*Term Frequency*) dan IDF (*Invers Document Frequency*), Dimana TF sendiri merupakan jumlah sebuah kata dari tiap dokumen, sedangkan IDF berfungsi mengurangi bobot suatu kata apabila kemunculannya banyak yang tersebar diseluruh dokumen. Berikut contoh perhitungan TF secara manualisasi yang terdapat pada table 4.4

Tabel 4. 4 Perhitungan TF

Token	TF		
	D1	D2	D3
It’s	1	0	0

Pretty	1	0	0
Good	1	1	0
But	1	0	0
They	1	0	0
Keep	1	0	0
Incentivizing	1	0	0
I	0	1	1
Just	0	1	0
Love	0	1	0
It	0	2	0
Everything	0	1	0
Is	0	1	0
In	0	1	0
Am	0	0	1
Actually	0	0	1
Learning	0	0	1
And	0	0	1
Speaking	0	0	1
Italian	0	0	1

- Doc 1 memiliki 7 kata
- Doc 2 memiliki 9 kata
- Doc 3 memiliki 8 kata

Selanjutnya perhitungan DF (*Document Frequency*)
terlebih dahulu, DF sendiri merupakan banyaknya

jumlah dokumen yang mengandung kata tersebut. Berikut contoh perhitungan DF secara manualisasi yang terdapat pada table 4.5

Tabel 4. 5 Perhitungan DF

Token	TF			DF
	D1	D2	D3	
It's	1	0	0	1
Pretty	1	0	0	1
Good	1	1	0	2
But	1	0	0	1
They	1	0	0	1
Keep	1	0	0	1
Incentivizing	1	0		1
I	0	1	1	2
Just	0	1	0	1
Love	0	1	0	1
It	0	2	0	2
Everything	0	1	0	1
Is	0	1	0	1
In	0	1	0	1
Am	0	0	1	1
Actually	0	0	1	1
Learning	0	0	1	1
And	0	0	1	1

Speaking	0	0	1	1
Italian	0	0	1	1

Setelah di dapatkan nilai DF, selanjutnya masuk ke perhitungan IDF. Diketahui jumlah dokumen (D) yang digunakan sebagai contoh disini yaitu 3 ulasan, sehingga $D = 3$. Berikut contoh perhitungan IDF pada tabel 4.6

Tabel 4. 6 Perhitungan IDF

Token	DF	D/DF (D=3)	IDF (Log(D/DF))	IDF+1
It's	1	3	Log 3= 0,477	1,477
Pretty	1	3	Log 3= 0,477	1,477
Good	2	1,5	Log 1,5=0,176	1,176
But	1	3	Log 3= 0,477	1,477
They	1	3	Log 3= 0,477	1,477
Keep	1	3	Log 3= 0,477	1,477
Incentivizing	1	3	Log 3= 0,477	1,477
I	2	1,5	Log 1,5=0,176	1,176
Just	1	3	Log 3= 0,477	1,477
Love	1	3	Log 3= 0,477	1,477
It	2	1,5	Log 1,5=0,176	1,176
Everything	1	3	Log 3= 0,477	1,477
Is	1	3	Log 3= 0,477	1,477

In	1	3	$\text{Log } 3 = 0,477$	1,477
Am	1	3	$\text{Log } 3 = 0,477$	1,477
Actually	1	3	$\text{Log } 3 = 0,477$	1,477
Learning	1	3	$\text{Log } 3 = 0,477$	1,477
And	1	3	$\text{Log } 3 = 0,477$	1,477
Speaking	1	3	$\text{Log } 3 = 0,477$	1,477
Italian	1	3	$\text{Log } 3 = 0,477$	1,477

Setelah nilai TF dan IDF didapatkan, Langkah selanjutnya menghitung TFIDF dengan mengalikan hasil nilai TF dengan nilai IDF. Berikut pada table 4.7 diperlihatkan perhitungan TFIDF.

Tabel 4. 7 Perhitungan TF-IDF

TF*IDF		
D1	D2	D3
1,477	0	0
1,477	0	0
1,176	1,176	0
1,477	0	0
1,477	0	0
1,477	0	0
1,477	0	0

0	1,176	1,176
0	1,477	0
0	1,477	0
0	2,352	0
0	1,477	0
0	1,477	0
0	1,477	0
0	0	1,477
0	0	1,477
0	0	1,477
0	0	1,477
0	0	1,477
0	0	1,477

```

Array([
[1,477, 1,477, 1,176, 1,477, 1,477, 1,477, 1,477, 0, 0, 0, 0,
0, 0, 0, 0, 0, 0, 0, 0],
[0, 0, 1,176, 0, 0, 0, 0, 1,176, 1,477, 1,477, 2,352, 1,477,
1,477, 1,477, 0, 0, 0, 0, 0, 0],
[0, 0, 0, 0, 0, 0, 0, 1,176, 0, 0, 0, 0, 0, 0, 1,477, 1,477, 1,477,
1,477, 1,477, 1,477]
])

```

Di mana setiap baris dan kolom menginterpretasikan setiap kata dalam teks secara keseluruhan. Dalam situasi seperti ini, representasi data biasanya disusun dalam matriks yang terdiri dari baris dokumen dan kolom kata-kata. Matriks ini memudahkan pemrosesan teks menjadi data numerik yang terstruktur, sehingga model pembelajaran mesin atau algoritma analitik lainnya dapat menganalisis hubungan antar kata dan dokumen. Nilai dalam matriks ini menunjukkan frekuensi atau bobot kemunculan kata tertentu dalam dokumen tertentu, tergantung pada teknik yang digunakan. Oleh karena itu, setiap kata yang ada dalam teks dapat dievaluasi berdasarkan pengaruh mereka terhadap perasaan, klasifikasi, atau pola tertentu yang ditemukan dalam berbagai dokumen.

F. Klasifikasi dan Evaluasi Naïve Bayes

Setelah pembobotan TF-IDF diterapkan, dengan menggunakan algoritma klasifikasi Naïve Bayes, komputer akan menggunakan data tersebut untuk tujuan pelatihan. Komputer akan menganalisis pola pada data pelatihan untuk membuat prediksi pada data baru berdasarkan pola yang diidentifikasi.

Selanjutnya, tahap penilaian menjadi sangat penting dalam menilai kemanjuran algoritma klasifikasi yang digunakan dalam penelitian. Evaluasi kinerja suatu

algoritma dapat dilakukan dengan pemanfaatan metrik seperti akurasi, presisi, dan perolehan. Pendekatan yang dikenal sebagai matriks konfusi digunakan untuk menghitung ukuran ini. Pendekatan ini terdiri dari komponen-komponen berikut: Ada empat jenis hipotesis: True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Seperti yang terlihat pada temuan keluaran penilaian Gambar dibawah ini :

```
MultinomialNB Accuracy: 0.86
MultinomialNB Precision: 0.8854166666666666
MultinomialNB Recall: 0.8333333333333334
MultinomialNB f1_score: 0.8585858585858586
confusion_matrix:
[[85 17]
 [11 87]]
=====
```

	precision	recall	f1-score	support
Negatif	0.89	0.83	0.86	102
Positif	0.84	0.89	0.86	98
accuracy			0.86	200
macro avg	0.86	0.86	0.86	200
weighted avg	0.86	0.86	0.86	200

Gambar 4.10 Hasil output evaluasi

Formulasi Confusion Matrix untuk menghitung akurasi , presisi dan Recall adalah sebagai berikut:

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN} = \frac{87 + 85}{87 + 11 + 85 + 17} = \frac{172}{200} = 0,86$$

$$Precision = \frac{TP}{TP + FP} = \frac{87}{87 + 11} = \frac{87}{98} = 0,89$$

$$Recall = \frac{TP}{TP + FN} = \frac{87}{87 + 17} = \frac{87}{104} = 0,83$$

Gambar 4.11 Hasil Formulasi *Confusion Matrix*

Keterangan :

- TP = Memprediksi kata positif dan itu benar
- TN = Memprediksi kata negative dan itu benar
- FP = Memprediksi kata positif dan itu salah
- FN = Memprediksi kata itu negative dan itu salah

Analisis sentimen evaluasi Aplikasi Duolingo diklasifikasikan menjadi sentimen positif dan negatif dengan menggunakan metode akurasi. Temuannya menunjukkan akurasi sebesar 0,86 atau 86%. Apabila dinyatakan seperti pada Gambar 9, sikap negatif menghasilkan skor presisi sebesar 89% dan hasil recall sebesar 83%. Skor f1 digunakan untuk menilai kualitas akurasi dan nilai recall. Kategorisasi tersebut

menghasilkan skor f1 sebesar 86%. Mengenai hasil sentimen positif, tingkat akurasinya adalah 84% dan tingkat recallnya adalah 89%. Skor f1 digunakan untuk menilai kualitas akurasi dan nilai recall. Kategorisasi tersebut menghasilkan skor f1 sebesar 86%.

Hasil ini menunjukkan bahwa model analisis sentimen yang digunakan dapat mengkategorikan ulasan aplikasi Duolingo dengan baik. Dengan tingkat akurasi keseluruhan yang tinggi, metode ini dapat dengan andal membedakan ulasan positif dan negatif. Namun, hasil presisi dan recall untuk sentimen negatif menunjukkan kinerja yang sangat baik, dan hasil recall yang sedikit lebih tinggi untuk sentimen positif menunjukkan bahwa model ini lebih mampu mengenali ulasan positif daripada ulasan negatif. Keseimbangan antara presisi dan recall diukur dengan skor f1 y. Ini sangat penting untuk mengukur kepuasan pengguna dan membantu pengembang aplikasi seperti Duolingo memahami bagaimana pengguna melihat layanan mereka.

Meskipun analisis sentimen sangat akurat, ada beberapa hal yang bisa ditingkatkan untuk lebih memahami ulasan pengguna. Salah satunya adalah analisis sentimen berbasis aspek. Dengan metode ini, ulasan pengguna tidak hanya dikategorikan sebagai

positif atau negatif secara keseluruhan, tetapi juga dievaluasi berdasarkan komponen tertentu dari aplikasi, seperti kinerja teknis, antarmuka pengguna, dan fitur gamifikasi. Metode ini akan memungkinkan pengembang untuk menemukan area aplikasi yang mendapatkan respons positif atau negatif, sehingga mereka dapat lebih berkonsentrasi pada memperbaiki aspek tertentu.

Penggunaan metode analisis yang lebih kompleks seperti *topic modeling* juga dapat membantu mengidentifikasi topik-topik utama yang sering dibahas dalam ulasan pengguna. Dengan demikian, pengembang aplikasi dapat memahami isu atau fitur mana yang paling menonjol dalam pikiran pengguna dan meresponsnya dengan lebih cepat. Menggabungkan berbagai metode ini akan memberikan analisis yang lebih kaya dan mendalam, membantu Duolingo terus meningkatkan layanan mereka sesuai kebutuhan dan harapan pengguna.

BAB V

KESIMPULAN DAN SARAN

A. Kesimpulan

Berdasarkan penelitian yang dilaksanakan, dapat disimpulkan bahwa :

1. Algoritma Naïve Bayes digunakan dalam analisis sentimen ulasan pada aplikasi Duolingo di Google Play Store. Sebelum analisis, dilakukan pelabelan manual pada 1000 ulasan untuk mengkategorikan sentimen positif dan negatif, dengan pembagian yang seimbang antara keduanya, yaitu masing-masing 500 ulasan positif dan 500 ulasan negatif. Setelah pelabelan, data diproses melalui tahap pre-processing sebelum akhirnya dianalisis menggunakan Naïve Bayes.
2. Hasil analisis menggunakan algoritma Naïve Bayes menunjukkan tingkat akurasi sebesar 86%. Ini menandakan bahwa algoritma Naïve Bayes mampu memprediksi sentimen ulasan dengan tingkat ketepatan yang cukup tinggi. Meskipun hasil ini memadai, masih terdapat ruang untuk peningkatan, seperti melalui penggunaan algoritma pembelajaran mesin yang lebih kompleks atau dataset yang lebih besar dan lebih beragam.
3. Penelitian ini memberikan gambaran yang cukup

menyeluruh tentang persepsi pengguna terhadap aplikasi Duolingo di Google Play Store. Analisis sentimen menunjukkan bahwa pengguna memiliki pandangan yang beragam terhadap aplikasi, dengan adanya keseimbangan antara sentimen positif dan negatif. Penelitian ini juga menunjukkan bahwa meskipun Naïve Bayes memberikan hasil yang baik, eksplorasi lebih lanjut terhadap metode lain dapat meningkatkan performa analisis.

B. Saran

Berdasarkan penelitian yang dilaksanakan, penulis berharap kepada peneliti selanjutnya untuk dapat dikembangkan dan terdapat beberapa saran, yaitu:

1. Menggunakan metode klasifikasi yang berbeda seperti *Support Vector Machine*(SVM), K-NN, *Decision Tree* dengan begitu dapat membandingkan hasil performa yang dilakukan untuk mencari metode klasifikasi yang terbaik.
2. Pada penelitian ini, data diambil dari ulasan Aplikasi Duolingo yang berada di Google Playstore, pada penelitian selanjutnya diharapkan bisa mengambil ulasan aplikasi lainnya.

DAFTAR PUSTAKA

- Agustina, Nurhaliza, et al. "Implementasi Algoritma Naive Bayes Untuk Analisis Sentimen Ulasan Shopee Pada Google Play Store." *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 2, no. 1, 2022, pp. 47–54, <https://doi.org/10.57152/malcom.v2i1.195>.
- Aisyah, Nur, and M. Hasan Hidayatullah. "Implementasi Aplikasi Duolingo Dalam Meningkatkan Kosa Kata Bahasa Inggris." *Bidayatuna Jurnal Pendidikan Guru Mandrasah Ibtidaiyah*, vol. 6, no. 1, 2023, pp. 44–59, <https://doi.org/10.54471/bidayatuna.v6i1.2015>.
- Ardiani, Lian, et al. "Implementasi Sentiment Analysis Tanggapan Masyarakat Terhadap Pembangunan Di Kota Pontianak." *Jurnal Sistem Dan Teknologi Informasi (Justin)*, vol. 8, no. 2, 2020, p. 183, <https://doi.org/10.26418/justin.v8i2.36776>.
- Bei, Fathurahman, and Saepudin Sudin. "Analisis Sentimen Aplikasi Tiket Online Di Play Store Menggunakan Metode Support Vector Machine (Svm)." *Sismatik*, vol. 01, no. 01, 2021, pp. 91–97.
- Birjali, Marouane, et al. "A Comprehensive Survey on Sentiment Analysis: Approaches, Challenges and Trends." *Knowledge-Based Systems*, vol. 226, 2021, p. 107134, <https://doi.org/10.1016/j.knosys.2021.107134>.
- Chohan, Saifurrachman, et al. "Analisis Sentimen Pengguna Aplikasi Duolingo Menggunakan Metode Naïve Bayes Dan Synthetic Minority Over Sampling Technique." *Paradigma - Jurnal Komputer Dan Informatika*, vol. 22, no. 2, 2020, pp. 139–44, <https://doi.org/10.31294/p.v22i2.8251>.
- Dewi, Ayu Kusuma. "Analisis Sentimen Ekspedisi Sicepat Dari Ulasan Google Play Menggunakan Algoritma Naïve Bayes." *JATISI (Jurnal Teknik Informatika Dan Sistem Informasi)*, vol. 9, no. 2, 2022, pp. 796–805, <https://doi.org/10.35957/jatisi.v9i2.1802>.
- Imron, Ali. "Analisis Sentimen Terhadap Tempat Wisata Di Kabupaten Rembang Menggunakan Metode Naive Bayes Classifier." *Teknik Informatika*, 2019, pp. 10–13,

<https://dspace.uii.ac.id/handle/123456789/14268>.

- Kaviani, Pouria, and Sunita Dhotre. "International Journal of Advance Engineering and Research Short Survey on Naive Bayes Algorithm." *International Journal of Advance Engineering and Research Development*, vol. 4, no. 11, 2017, pp. 607–11.
- Khoirul Insan, Moh Khoirul, et al. "Analisis Sentimen Aplikasi Brimo Pada Ulasan Pengguna Di Google Play Menggunakan Algoritma Naive Bayes." *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 7, no. 1, 2023, pp. 478–83, <https://doi.org/10.36040/jati.v7i1.6373>.
- Lestari, Sri, and Sudin Saepudin. "Analisis Sentimen Vaksin Sinovac Pada Twitter Menggunakan Algoritma Naive Bayes." *SISMATIK (Seminar Nasional Sistem Informasi Dan Manajemen Informatika)*, 2021, pp. 163–70.
- Mursyid, Rabhita, and Aries Dwi Indriyanti. "Perbandingan Akurasi Metode Analisis Sentimen Untuk Evaluasi Opini Pengguna Pada Platform Media Sosial (Studi Kasus: Twitter)." *Journal of Informatics and Computer Science (JINACS)*, vol. 06, 2024, pp. 371–83, <https://ejournal.unesa.ac.id/index.php/jinacs/article/view/61322%0Ahttps://ejournal.unesa.ac.id>.
- Nugroho, Agung. *Analisis Sentimen Pada Data Transaksi Keuangan Menggunakan Teknik Data Mining*. 2024, pp. 1–11.
- Nurian, Andriani. "Analisis Sentimen Ulasan Pengguna Aplikasi Google Play Menggunakan Naive Bayes." *Jurnal Informatika Dan Teknik Elektro Terapan*, vol. 11, no. 3s1, 2023, pp. 829–35, <https://doi.org/10.23960/jitet.v11i3s1.3348>.
- Prabowo, Abram Setyo, and Felix Indra Kurniadi. "Analisis Perbandingan Kinerja Algoritma Klasifikasi Dalam Mendeteksi Penyakit Jantung." *Jurnal SISKOM-KB (Sistem Komputer Dan Kecerdasan Buatan)*, vol. 7, no. 1, 2023, pp. 56–61, <https://doi.org/10.47970/siskom-kb.v7i1.468>.
- Rakhmawati, Nur Aini, et al. "Analisis Klasifikasi Sentimen Pengguna Media Sosial Twitter Terhadap Pengadaan Vaksin COVID-19." *Journal of Information Engineering and Educational Technology*, vol. 4, no. 2, 2020, pp. 90–92,

- <https://doi.org/10.26740/jieet.v4n2.p90-92>.
- Sanjaya, Gupta, and Kemas Muslim Lhaksmana. *Lexicon Based*). no. 3, 2020, pp. 9698–710.
- Saputra. *Landasan Teori Google Play Store*. 2019, pp. 1–23.
- Septian, Jeremy Andre, et al. “Analisis Sentimen Pengguna Twitter Terhadap Polemik Persepakbolaan Indonesia Menggunakan Pembobotan TF-IDF Dan K-Nearest Neighbor.” *Journal of Intelligent System and Computation*, vol. 1, no. 1, 2019, pp. 43–49, <https://doi.org/10.52985/insyst.v1i1.36>.
- Syarli, and Asrul Ashari Muin. “Metode Naive Bayes Untuk Prediksi Kelulusan.” *Jurnal Ilmiah Ilmu Komputer*, vol. 2, no. 1, 2020, pp. 22–26, <https://media.neliti.com/media/publications/283828-metode-naive-bayes-untuk-prediksi-kelulu-139fcfea.pdf>.
- Syawal, Rima Rahmawati, et al. “Klasifikasi Kualitas Ikan Nilem Berdasarkan Ukuran Menggunakan Algoritma Naive Bayes.” *J I M P - Jurnal Informatika Merdeka Pasuruan*, vol. 7, no. 2, 2023, p. 72, <https://doi.org/10.51213/jimp.v7i2.514>.
- Tangkelayuk, Aldi. “The Klasifikasi Kualitas Air Menggunakan Metode KNN, Naïve Bayes, Dan Decision Tree.” *JATISI (Jurnal Teknik Informatika Dan Sistem Informasi)*, vol. 9, no. 2, 2022, pp. 1109–19, <https://doi.org/10.35957/jatisi.v9i2.2048>.
- Wardana, Hartanto Kusuma, et al. “Sistem Pemeriksa Pola Kalimat Bahasa Indonesia Berbasis Algoritme Left-Corner Parsing Dengan Stemming.” *Jurnal Nasional Teknik Elektro Dan Teknologi Informasi (JNTETI)*, vol. 8, no. 3, 2019, p. 211, <https://doi.org/10.22146/jnteti.v8i3.515>.
- Wickramasinghe, Indika, and Harsha Kalutarage. “Naive Bayes: Applications, Variations and Vulnerabilities: A Review of Literature with Code Snippets for Implementation.” *Soft Computing*, vol. 25, no. 3, 2021, pp. 2277–93, <https://doi.org/10.1007/s00500-020-05297-6>.

DAFTAR LAMPIRAN

LAMPIRAN 1 : PERSETUJUAN PEMBIMBING PROPOSAL

PERSETUJUAN PEMBIMBING

Proposal Skripsi ini telah disetujui oleh Pembimbing
untuk dilaksanakan.

Disetujui pada

Hari : Rabu

Tanggal : 18 September 2024

Pembimbing I,



Siti Nur Aini, S.Kom., M.Kom
NIP. 198401312018012001

Pembimbing II,



Dr. Khothibul Umam S.T., M. Kom
NIP. 197908272011011007

Mengetahui,

Ketua Jurusan Teknologi Informasi



Dr. Khothibul Umam S.T., M. Kom
NIP. 197908272011011007

LAMPIRAN 2 : LOA JURNAL



LETTER OF ACCEPTANCE

Dear Meli Apriyiani, Mirza Izzal Musyaffaq, Siti Nur'Aini, Maya Rini Handayani, Khotibul Umam,

I hope this letter finds you well. On behalf of the editorial board of the Journal of AITI, I am pleased to inform you that your manuscript entitled "**Implementasi Analisis Sentimen pada Ulasan Aplikasi Duolingo di Google Playstore menggunakan Algoritma Naïve Bayes**" has been **accepted** for publication in *Volume 22 No. 1* issue and will be available online at <https://ejournal.uksw.edu/aiti> on *March 30, 2025*.

Your research was thoroughly reviewed by our expert panel of reviewers. We believe that your contribution will greatly enhance the content and value of our journal. To proceed with the publication process, please provide the final version of your manuscript in accordance with our journal's formatting guidelines and make sure to address any comments or suggestions provided by the reviewers.

Once again, congratulations on the acceptance of your manuscript. We look forward to publishing your work in the Journal of AITI

Thank you for choosing our journal as the platform for sharing your valuable research.

Sincerely,

A blue ink signature of Dr. Indrastanti R. Widiastari, M.T. is written over the AITI logo.

Dr. Indrastanti R. Widiastari, M.T.
Editor in Chief

LAMPIRAN 3 : HISTORY JURNAL

Submit : 25 April 2024

Artikel

Pra- : 25 April 2024

Review

Uploud : 30 April 2024

Revisi

PraReview

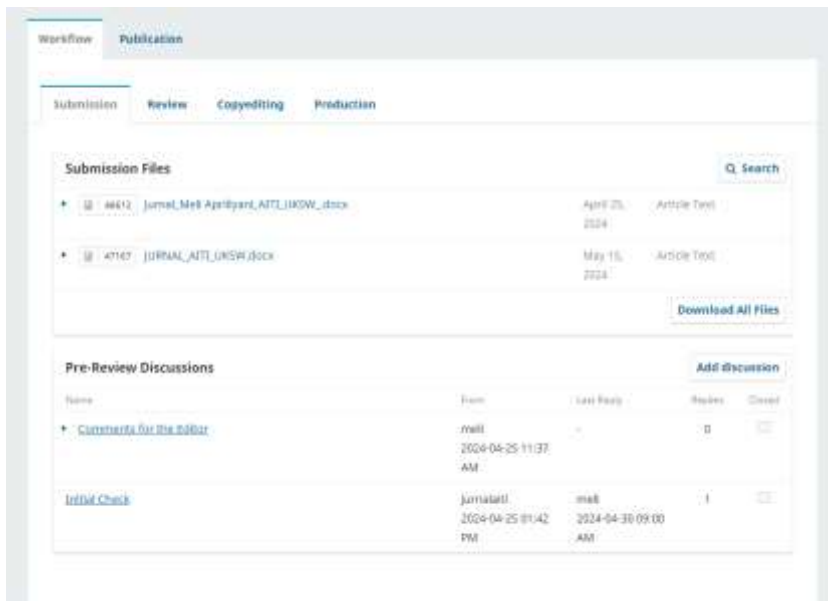
Review : 17 Mei 2024

Uploud : 27 Mei 2024

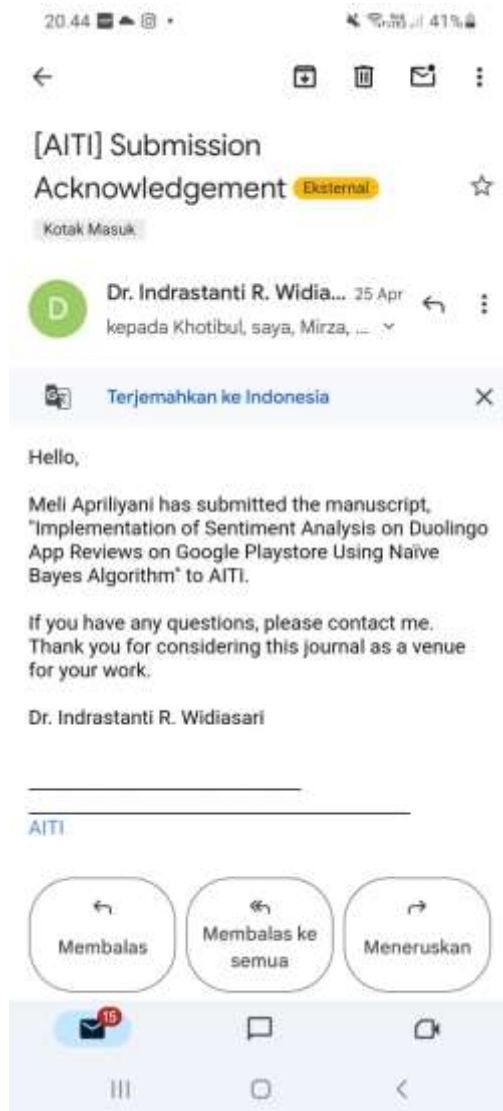
Revisi

Publikasi : 30 September 2024

Link : <https://ejournal.uksw.edu/aiti/article/view/12100>



Proses Submit Artikel



Gmail Naskah Diterima

Initial Check

Participants

JURNAL AITI (jurnalaiti)
 Mel Apriyanti (net)

Messages

From	To
Dear Author,	jurnalaiti 2024-04-25 05:42 PM
Terima kasih telah submit naskah pada jurnal AITI. Setelah dilakukan pengecekan plagiasi hasil turnitinnya menunjukkan sebesar 7%, maksimal 15%. Mohon dapat dihilangkan kata "kita", "kami", "peneliti", "penulis" atau kata-kata lain yang mengujuk pada identitas penulis, tidak diperkenankan dipergunakan di dalam naskah. Serta kata di bawah atau di atas bisa digantikan dengan langsung menyebut Gambar/Tabel yang dimaksud, misalnya Gambar 1, Tabel 1, dst. Silahkan mengirimkan perbaikan melalui discussion ini dan kami tunggu sampai dengan tanggal 9 Mei 2024. Terima kasih  12100_AITI_naskah.pdf Terimakasih kak atas koreksinya, tjd mengirimkan naskah saya yang sudah saya perbaiki  Mel Apriyanti_JURNAL_AITLUGSW.docx Delete	mes 2024-04-30 08:00 AM

Add Message

Proses Pra Revisi

Workflow
Publications

Submission
Review
Copyrighting
Production

Round 1
Round 2

Round 2 Status

Submission accepted.

Reviewer's Attachments

Q Search

No Files

Revisions

Q Search
Upload File

No Files

Review Discussions

Add discussion

Name	Date	Last Reply	Replies	Owner
Proses Review	evangs 2024-05-15 09:39 PM	-	0	
Proses Revisi	evangs 2024-05-17 12:29	mail 2024-05-27 10:38	1	

Proses Pra Revisi



Pemberitahuan untuk Revisi

Participants

Evangs Mailoa (evangs)

Meli Apriliyani (meli)

Messages

Note	From
Dear Author, Naskah Anda sudah diberi masukan oleh para reviewer. Mohon perbaiki dan submit kembali untuk diperiksa oleh reviewer. Mohon unggah file ke sini ya... Beri tanda highlight kuning untuk bagian yang sudah Anda revisi sesuai masukan dari masing-masing reviewer. Salam,	evangs 2024-05-17 12:29 PM
• Selamat Pagi bapak/ibu, saya sudah memperbaiki naskah saya sesuai dengan revisi bapak/ibu: Untuk Reviewer A : 1.saya telah mencantumkan bulan dan tahun dalam pengambilan data saya. 2.data sudah saya seimbangkan yaitu 500 positif dan 500 negatif. 3. pelabelan saya menggunakan library NLTK dari phyton. 4 akurasi sudah menjadi 86% dimana sudah seharusnya memenuhi syarat untuk Reviewer B : 1. Sudah saya tambahkan state of the art pada tambahan referensi saya. 2.kesimpulan sudah saya buat menjadi 1 paragraf dan saya tambahi saran 3. untuk reviewer C : 1. sudah saya tambahi alasan mengapa saya memilih algoritma naive bayes pada metodologi penelitian. 2. sudah saya tambahi riset terdahulu. dan saya mengganti gambar hasil case folding, tokenizing dan stopword menjadi tabel agar lebih jelas .	meli 2024-05-27 10:28 AM

Proses Revisi

[AITI] Editor Decision

Eksternal



Kotak Masuk



JURNAL AITI Kemarin

kepada saya, Mirza, Siti, Maya... ▾



Terjemahkan ke Indonesia

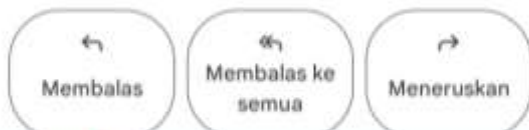


Meli Apriliyani, Mirza Izzal Musyaffaq, Siti Nur'Aini, Maya Rini Handayani, Khotibul Umam:

The editing of your submission, "Implementasi analisis sentimen pada ulasan aplikasi Duolingo di Google Playstore menggunakan algoritma Naïve Bayes," is complete. We are now sending it to production.

Submission URL: <https://ejournal.uksw.edu/aiti/authorDashboard/submission/12100>

AITI



Pemberitahuan Naskah Terbit

Workflow Publication

Submission Review Copyediting Production

Copyediting Discussions [Add discussion](#)

Name	Room	Last Reply	Replies	Closed
Pembayaran Article Processing Charge	awangs	Jurnalati 2024-05-30 07:48 AM	2	

Copyedited [Q Search](#)

5137 10_12100_copyedit.docx	September 30, 2024	Article Text
-----------------------------	--------------------	--------------

Naskah telah terbit

atj JURNAL TEKNOLOGI INFORMASI

Volume 10 Nomor 1 2024

Implementasi analisis sentimen pada ulasan aplikasi Desain di Google Playstore menggunakan algoritma Naive Bayes

Full Article
[Download Full Article](#)
[Download Full Article](#)
[Download Full Article](#)

Abstract
 Penelitian ini bertujuan untuk menganalisis sentimen ulasan pengguna terhadap aplikasi Desain di Google Playstore menggunakan algoritma Naive Bayes. Hasil penelitian menunjukkan bahwa mayoritas ulasan pengguna bersifat positif, dengan persentase sebesar 75%. Penelitian ini diharapkan dapat memberikan informasi yang berguna bagi pengembang aplikasi untuk meningkatkan kualitas produk mereka.

Keywords
 Sentiment Analysis, Naive Bayes, Google Playstore, Aplikasi Desain

DOI [https://doi.org/10.24127/atj.v10i1.12100](#)

Keywords Sentiment Analysis, Naive Bayes, Google Playstore, Aplikasi Desain

Full Article
[Download Full Article](#)
[Download Full Article](#)
[Download Full Article](#)

Abstract
 Penelitian ini bertujuan untuk menganalisis sentimen ulasan pengguna terhadap aplikasi Desain di Google Playstore menggunakan algoritma Naive Bayes. Hasil penelitian menunjukkan bahwa mayoritas ulasan pengguna bersifat positif, dengan persentase sebesar 75%. Penelitian ini diharapkan dapat memberikan informasi yang berguna bagi pengembang aplikasi untuk meningkatkan kualitas produk mereka.

Keywords
 Sentiment Analysis, Naive Bayes, Google Playstore, Aplikasi Desain

DOI [https://doi.org/10.24127/atj.v10i1.12100](#)

Keywords Sentiment Analysis, Naive Bayes, Google Playstore, Aplikasi Desain



Sertifikat jurnal



Implementasi analisis sentimen pada ulasan aplikasi Duolingo di Google Playstore menggunakan algoritma Naïve Bayes

Meli Apriliyani ¹⁾, Mirza Izzal Musyaffaq ¹⁾, Siti Nur'Aini ¹⁾,
Maya Rini Handayani ¹⁾, Khotibul Umam ²⁾

^{1,2)}Program Studi Teknologi Informasi, Fakultas Sains dan Teknologi,
UIN Walisongo Semarang
Jl. Prof. Dr. Hamka, Semarang 50185, Jawa Tengah, Indonesia.
Email: ²⁾khotibul_umam@walisongo.ac.id

	Riwayat Artikel	
Diterima:	Direvisi:	Disetujui:
25-04-2024	27-05-2024	28-05-2024

Abstract

This study investigates sentiment analysis to evaluate the Duolingo Application using the Naive Bayes method. The Duolingo program exemplifies the use of big data technology for processing vast and complex data. Google Play Store offers review and rating functions that can help with program development and fixing undesirable aspects. This project uses sentiment analysis techniques to automatically analyze Indonesian internet product reviews and obtain information about the feelings expressed in these reviews. The Naïve Bayes method is used to classify reviews as positive or negative. The research findings show that a dataset consisting of 1000 pieces of data originating from reviews of the Duolingo program on the Google Play Store was manually labeled before the preprocessing step. Of this number, 500 data have positive sentiments, while 500 have negative attitudes. Additionally, sentiment analysis shows an accuracy rate of 86 persen. The f1 score precision is 89 persen and recall 83 persen values, with f1 results for classification of 86 persen.

Keywords: Duolingo App, Naïve Bayes, sentiment analysis

Abstrak

Penelitian ini menyelidiki analisis sentimen evaluasi Aplikasi Duolingo menggunakan metode Naive Bayes. Program Duolingo mencontohkan penggunaan teknologi data besar untuk pemrosesan data yang luas dan rumit. Google Play Store menawarkan fungsi peninjauan dan pemberitahuan yang dapat membantu pengembangan program dan perbaikan aspek yang tidak diinginkan. Proyek ini menggunakan teknik analisis sentimen yang secara otomatis menganalisis ulasan produk internet Indonesia dan mendapatkan informasi mengenai perasaan yang diungkapkan dalam ulasan tersebut. Metode Naive Bayes digunakan untuk menentukan klasifikasi ulasan menjadi positif atau negatif. Temuan penelitian menunjukkan bahwa kumpulan data yang terdiri dari 1000 data yang berasal dari ulasan program Duolingo di Google Play Store diberi label secara manual sebelum ke langkah prapemrosesan. Dari jumlah tersebut, 500 data memiliki sentimen positif, sedangkan 500 data memiliki sikap negatif. Selain

itu, analisis sentimen menunjukkan tingkat akurasi sebesar 86 persen. Skor *f1* menunjukkan nilai presisi 89 persen dan *recall* 83 persen, dengan hasil *f1* pada klasifikasi sebesar 86 persen.

Kata kunci: Aplikasi Duolingo, *Naïve Bayes*, Analisis Sentimen

Pendahuluan

Pesatnya kemajuan teknologi telah memberikan banyak manfaat bagi masyarakat di berbagai bidang kehidupan, memudahkan tugas sehari-hari, meningkatkan produktivitas, dan memperluas akses terhadap informasi dan layanan. Kemajuan teknologi yang pesat menghasilkan data dalam jumlah besar, yang dapat menghasilkan wawasan berharga jika ditangani dan dimanfaatkan secara efektif. Seiring dengan meningkatnya jumlah pengguna internet aktif, volume data yang dihasilkan juga meningkat. Teknologi *big data* memfasilitasi analisis sejumlah besar data yang rumit dan beragam, sehingga menghasilkan pengetahuan yang berharga [1].

Aplikasi Duolingo mencontohkan kemampuan teknologi dalam menangani kumpulan data yang luas dan rumit. Aplikasi Duolingo adalah program tanpa biaya yang dikembangkan oleh Luis Von Ahn dan Severin Hacker pada bulan November 2011. Tagline organisasi ini adalah "Menyediakan pendidikan bahasa gratis kepada orang-orang di seluruh dunia." Situs web perusahaan mengklaim bahwa mereka memiliki lebih dari 30 juta orang yang telah mendaftar [2]. Aplikasi Duolingo mendapat rating terpuji menurut statistik yang diperoleh dari Google Playstore. Mengingat banyaknya ulasan pengguna untuk aplikasi Duolingo di Google Playstore, sulit untuk memastikan proporsi pasti dari tanggapan baik atau negatif. Memanfaatkan teknologi *big data*, aplikasi Duolingo memiliki kemampuan untuk menganalisis kumpulan data yang luas dan rumit, termasuk data penggunaan aplikasi, data penggunaan kata, dan data penggunaan fitur. Dengan menggunakan data ini, aplikasi Duolingo dapat menghasilkan wawasan berharga, seperti frekuensi penggunaan kata, pola penggunaan berbagai fitur, dan pola penggunaan aplikasi secara keseluruhan. Dengan memanfaatkan informasi ini, pengguna dapat meningkatkan pemanfaatan program, meningkatkan pengalaman belajar bahasa, dan mengoptimalkan pemanfaatan kemampuan aplikasi.

Play Store adalah layanan yang disediakan Google dan menyediakan berbagai macam materi digital termasuk permainan, aplikasi, film, musik, dan buku, yang disusun dalam beberapa kategori berbeda [3]. Play Store memiliki fungsi penilaian dan ulasan yang memungkinkan pelanggan memberikan komentar dan penilaian terhadap barang-barang yang telah mereka gunakan. Melalui penggunaan kemampuan ini pengguna dapat berkontribusi pada peningkatan aplikasi dan memperbaiki fungsionalitas yang tidak berfungsi. Hal ini memudahkan konsumen untuk memilih program yang sesuai dengan preferensi dan kebutuhan spesifik mereka. Ada lebih dari 500 juta unduhan aplikasi Duolingo yang tersedia di YouTube Play Store. Melalui

pemanfaatan pendekatan analisis *sentiment*, tujuan penelitian ini adalah untuk melakukan investigasi mendalam terhadap program Duolingo.

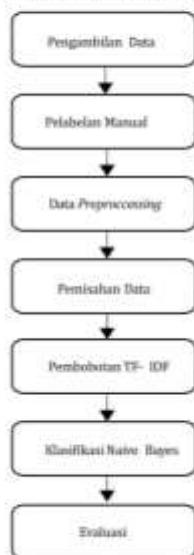
Analisis sentimen adalah salah satu pendekatan yang dapat diambil untuk mengatasi masalah pengkategorian evaluasi secara otomatis menjadi lebih positif atau lebih negatif [4]. Analisis sentimen sering digunakan dalam pemantauan media sosial untuk mendeteksi sentimen yang ada, termasuk opini positif dan negatif [5]. Tujuan dari penelitian ini adalah menggunakan metode Naïve Bayes untuk melakukan analisis sentimen pada evaluasi aplikasi Duolingo. Namun demikian, penting untuk digarisbawahi bahwa hanya sejumlah kecil penelitian yang melakukan analisis sentimen pada program Duolingo. Tujuan dari penelitian ini adalah untuk memberikan penjelasan mengenai pendekatan analisis sentimen yang menggunakan algoritma Naïve Bayes, yaitu teknik *text mining* yang diimplementasikan dalam bahasa komputer Python. Untuk mengklasifikasikan ulasan sebagai baik atau negatif, metode Naïve Bayes akan digunakan. Pendekatan Naïve Bayes dipilih karena penerapannya yang luas dalam penambahan teks dan akurasi yang luar biasa. Menerapkan pendekatan Naïve Bayes dapat mengurangi durasi yang diperlukan untuk kategorisasi analisis sentimen [6].

Penelitian sebelumnya yang membahas tentang analisis sentimen dan menggunakan data pada Google Play Store yaitu penelitian tentang *App Review in Google Play Store*. Penelitian ini menemukan bahwa Naïve Bayes dapat membantu memecahkan masalah pengklasifikasian sentimen pada ulasan aplikasi Shopee dengan akurasi yang tinggi sebesar 96,667 persen [7]. Penelitian lain yang membahas analisis sentimen menggunakan data Google Play adalah untuk menganalisis sentimen aplikasi BRIMO. Penelitian ini tidak menjelaskan kecenderungan sentimen pada pengguna terhadap topik yang dibahas dengan menghasilkan akurasi yang cukup tinggi yaitu 84,25 persen [8]. Kedua penelitian ini memiliki kesamaan dalam penggunaan algoritma dan hanya fokus pada pengklasifikasian ulasan dan akurasi metodenya tidak menjelaskan kecenderungan sentimen pengguna terhadap topik yang dibahas.

Metode Penelitian

Penelitian ini fokus pada pemeriksaan penilaian pengguna terhadap program Duolingo di Google Playstore. Dengan menggunakan Algoritma Naïve Bayes untuk memprediksi suatu ulasan apakah positif atau negatif, dengan kemampuan yang dapat mengatasi data dengan dimensi yang tinggi [9]. Algoritma Naïve Bayes dipilih karena sangat akurat dan banyak digunakan dalam proses *text mining* [10]. Selain itu, algoritma Naïve Bayes adalah pengklasifikasi probabilistik yang paling sederhana dan mudah digunakan [11]. Metode penelitian ini terdiri dari tujuh tahapan penting: pengumpulan data, anotasi manual, pemrosesan awal data, pembobotan TF-IDF, partisi data ke dalam set pelatihan dan pengujian, klasifikasi menggunakan algoritma *Naïve Bayes*, dan

evaluasi. Dengan mengikuti proses ini secara sistematis, penelitian ini dapat menghasilkan hasil yang reliabel dan bermanfaat dalam memahami persepsi pengguna terhadap Aplikasi Duolingo. Gambar 1 menunjukkan alur penelitian.



Gambar 1 Alur penelitian

Pengumpulan Data. Proses pengumpulan data dilakukan melalui bahasa pemrograman Python dan *scraping tool* Google Play Store. Data yang diambil adalah 1000 data mulai *rating* bintang 1 sampai 5. Dengan menggunakan teknik *scraping*, data ulasan aplikasi Duolingo di Google Playstore diambil secara Otomatis. data yang dikumpulkan termasuk evaluasi pengguna, berupa nilai Bintang dan ulasan. Data yang telah dikumpulkan akan disimpan dalam *database* untuk penggunaan kasus yang akan datang.

Pelabelan. Setelah data tinjauan diperoleh, tugas pemberian label positif atau negatif menggunakan *library* NLTK. Setelah proses pelabelan selesai, data ulasan dibagi menjadi dua kategori berbeda: data ulasan positif yang mencakup ulasan dengan *rating* Bintang 4 atau 5 dan data ulasan negatif yang terdiri dari ulasan dengan *rating* bintang 1, 2, atau 3. Pembagian ini dilakukan untuk menyederhanakan prosedur analisis sentimen berikutnya dan memungkinkan fokus pada ulasan yang paling relevan dalam masing-masing kategori sentimen. Dengan demikian, analisis sentimen

dapat dilakukan dengan lebih efisien dan benar untuk lebih memahami opini pengguna Duolingo.

Data Preprocessing. Untuk melakukan analisis sentimen penting untuk melalui langkah *pretreatment* data terlebih dahulu. Pada langkah pra-pemrosesan dilakukan pemilihan data untuk menyempurnakan struktur data yang akan digunakan [12]. Prosedur penyiapan data terdiri dari empat tahapan yaitu *casefolding*, *tokenizing*, *stopwords*, dan *stemming*.

Case folding merupakan tahap pengolahan data atau mengkonversi semua huruf menjadi huruf kecil [13]. Proses ini membantu dalam menjaga konsistensi dalam teks, sehingga tidak ada perbedaan yang tidak perlu antara kata yang ditulis dalam huruf besar atau kecil. *Stopwords* adalah proses menghilangkan kata-kata yang tidak bermakna dari teks yang sedang diproses [14]. Kata *stopwords* adalah kata umum yang tidak memberikan nilai tambah pada analisis sentimen, contohnya "saya", "dan", "atau", "di", dan lainnya. Penghapusan kata *stopwords* membantu dalam fokus pada kata-kata kunci yang lebih relevan dalam menentukan sentiment. *Tokenizing* adalah langkah di mana kalimat dipecah menjadi beberapa kata yang disebut sebagai token yang memisahkannya [15]. Teks dibagi menjadi unit-unit yang lebih kecil dengan prosedur ini, yang membuat analisis di masa depan lebih mudah dilakukan. *Stemming* terdiri dari proses menghilangkan kata-kata yang mempunyai imbuhan untuk diubah menjadi kata-kata sederhana [16]. Proses ini membantu dalam mengurangi variasi kata yang sebenarnya memiliki arti sama, sehingga memperkuat analisis sentimen dengan memperlakukan kata-kata tersebut secara konsisten. Dengan melalui tahapan-tahapan tersebut proses *preprocessing* dapat dipastikan data siap diolah lebih lanjut dalam analisis sentimen, sehingga dapat menghasilkan hasil yang lebih akurat.

Pemisahan data. Pemisahan data yang juga dikenal sebagai produksi data pelatihan dan data pengujian dilakukan untuk membagi data pra-pemrosesan menjadi subset tertentu untuk tujuan perbandingan. Meskipun tujuan utamanya adalah memperoleh kumpulan data yang cukup representatif untuk tujuan pelatihan model, kumpulan data pengujian digunakan untuk tujuan mengevaluasi efektivitas model yang dilatih. Dengan melakukan hal ini, dapat dinilai sejauh mana model menunjukkan kemampuan generalisasi terhadap data baru yang sebelumnya tidak teramati. Memastikan perbandingan yang tepat antara data pelatihan dan data pengujian sangat penting untuk menjamin bahwa model dapat menghasilkan hasil yang memuaskan dalam skenario dunia nyata.

Pembobotan TF-IDF. Informasi yang telah diproses sebelumnya, yang mungkin berupa kata-kata, akan diubah menjadi representasi numerik dengan menggunakan prosedur pembobotan kata. Teknik pembobotan kata berupaya untuk menentukan proporsi kepentingan setiap kata yang dijadikan ciri dalam penelitian. Frekuensi suatu

istilah dalam teks olahan berkorelasi langsung dengan bobotnya. Akibatnya, ketika jumlah dokumen yang diproses bertambah, semakin banyak pula karakteristik yang dihasilkan. Hal ini memungkinkan pemeriksaan yang lebih komprehensif dan berpotensi meningkatkan kaliber dan presisi temuan analitis yang dilakukan. Perhitungan bobot kata pada pendekatan TF-IDF merupakan hasil penggabungan dua gagasan yaitu *Term Frequency* (TF) dan *Inverse Document Frequency* (IDF) [17]. Salah satu definisi TF adalah berapa kali setiap kata muncul dalam dokumen tertentu. Nilai TF meningkat sebanding dengan jumlah kata yang ada di setiap dokumen dan sering muncul.

Klasifikasi Naïve Bayes. Proses klasifikasi merupakan suatu metode penambahan data yang dapat digunakan untuk membangun suatu model hingga semua data termasuk dalam kategori yang sama [18]. Teorema Bayes berfungsi sebagai premis fundamental yang digunakan dalam pendekatan klasifikasi yang digunakan dalam algoritma Naïve Bayes. Teorema ini adalah yang pertama. Seorang ilmuwan Inggris bernama Thomas Bayes lah yang menemukan Naïve Bayes [19]. Teknik Naïve Bayes sering digunakan oleh beberapa akademisi dalam analisis sentimen. Klasifikasi Naïve Bayes memprediksi kemungkinan kepemilikan suatu kelas dengan mengasumsikan bahwa karakteristik yang digunakan tidak tergantung satu sama lain. Metode Naïve Bayes memiliki keunggulan efisiensi, memungkinkan prosedur analisis sentimen lebih cepat dan efisien. Hal lebih lanjut yang perlu dipertimbangkan adalah bahwa algoritma Naïve Bayes memiliki kecenderungan untuk memperoleh akurasi tinggi dengan jumlah data pelatihan yang terbatas.

Evaluasi. Evaluasi merupakan fase yang mengukur efektivitas algoritma kategorisasi yang digunakan dalam penelitian. Kinerja sistem dinilai menggunakan ukuran seperti akurasi, presisi, dan perolehan. Untuk melakukan komputasi menggunakan benchmark ini, penting untuk menggunakan metodologi *Confusion Matrix* yang terdiri dari komponen *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN). Adapun rumus *Confusion Matrix* untuk menghitung akurasi, presisi, dan *recall* adalah seperti yang terlihat pada Persamaan (1) sampai (3).

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN} \quad (1)$$

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

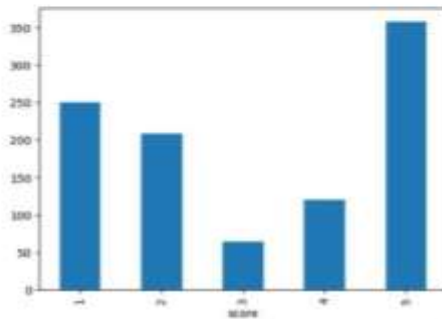
$$Recall = \frac{TP}{TP + FN} \quad (3)$$

Hasil dan Pembahasan

Pada tahap ini akan dijelaskan tahapan yang berkaitan dengan prosedur yang dilakukan dalam penelitian, sebagaimana dituangkan dalam metodologi penelitian

sebelumnya.

Tahap pertama adalah tahap pengambilan data. Pengambilan data dilakukan dengan memanfaatkan URL program Duolingo dari Playstore dimulai bulan Maret 2024. Gambar 2 dan Gambar 3 menyajikan hasil ekstraksi data yang dilakukan dengan menggunakan kata kunci "Duolingo" dan besarnya pencarian data dari 1000 kemunculan aplikasi Duolingo.



Gambar 2 Data mentah ulasan positif dan negatif

```
score
5    358
1    250
2    200
4    120
3     64
Name: count, dtype: int64
```

Gambar 3 Hasil output *scraping*

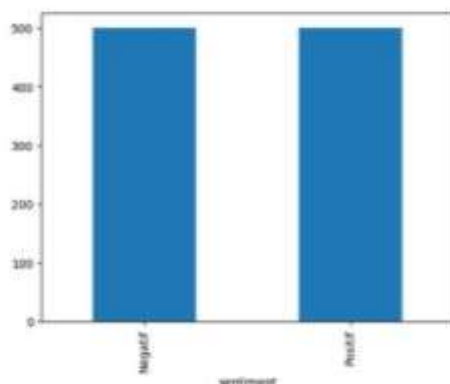
Setelah melakukan ekstraksi data, hasilnya menunjukkan bahwa aplikasi Duolingo memiliki beragam ulasan dari pengguna yang berbeda. Dari ulasan yang yang diberi bintang 1, 2, 3, 4 dan 5.

Setelah mengumpulkan data penelitian, dilanjutkan dengan melakukan pelabelan yang akan menggunakan *library* NLTK. NLTK (*Natural Language Toolkit*) adalah sebuah perpustakaan untuk pemrosesan bahasa alami Python yang dikembangkan oleh Stevan Bird dan Edward Loper [20]. Gambar 4 adalah hasil dari prosedur pelabelan yang telah dilakukan.

```
import nltk
from textblob import TextBlob
all_polarity = []
status=[]
total_positif = total_negatif = total = 0

for ulasan in data['content']:
    analysis = TextBlob(ulasan)
    all_polarity += analysis.polarity
    if (analysis.sentiment.polarity > 0.0):
        total_positif += 1
        status.append("Positif")
    elif (analysis.sentiment.polarity <= 0.0):
        total_negatif += 1
        status.append("Negatif")
    total += 1
print(f"Hasil Analisis Data: \nPositif = {total_positif} \nNegatif = {total_negatif}")
print(f"Total Data = {total}")
```

Gambar 4 Pelabelan NLTK



Gambar 5 Hasil Pelabelan NLTK

Hasil dari prosedur pelabelan NLTK ini memberikan wawasan penting tentang cara pengguna memandang platform pembelajaran bahasa online Duolingo. Dengan membedakan ulasan menjadi kategori positif dan negatif. Pada analisis pada data ini menghasilkan analisis data positif sebanyak 500 dan data negatif sebanyak 500.

Tahap selanjutnya adalah tahap *Preprocessing* yang merupakan langkah krusial yang harus diselesaikan sebelum menggunakan algoritma klasifikasi pada dokumen. Tahapan pra-pemrosesan adalah sebagai berikut:

Case Folding. Untuk menjamin keseluruhan teks yang dihasilkan konsisten, langkah ini dilakukan untuk mengubah semua huruf menjadi huruf kecil. Selama

prosedur ini, kata-kata yang tidak diperlukan dibersihkan dari tanda baca atau karakter tertentu. Tabel 1 merupakan hasil saat menggunakan *case folding*. Hasil dari *case folding*, semua huruf menjadi huruf kecil. Hal ini memberikan keseragaman dalam teks dan mempermudah analisis selanjutnya.

Tabel 1 Hasil *case folding*

Score	Content	At	Content clean
3	This is the one. I've tried several apps and pr...	2024-05-13 03:46:51	this is the one ive tried several apps and pro...
4	It has a decent variety of exercises. it does ...	2024-05-13 10:02:02	it has a decent variety of exercises it does n...
3	Most important it's free, you can upgrade but ...	2024-05-09 19:21:19	most important it s free you can upgrade but i...
2	I really like the setup and the levels. I felt...	2024-05-04 11:01:20	i really like the setup and the levels i felt ...
4	Fantastic app designed to keep you engaged wit...	2024-05-10 03:38:26	fantastic app designed to keep you engaged wit...

Stopwords. Selanjutnya adalah tahapan *stopwords*, *Stopwords* merupakan tahap yang penting juga yaitu penghapusan kata-kata yang tidak penting dan tidak bermakna berdasarkan kata yang terdapat pada database *stopword*, seperti "dan", "aku", "dari", dan lain-lain. Tabel 2 merupakan hasil dari *stopwords*. Penekanan pada kata-kata kunci yang lebih relevan dalam menentukan sentimen difasilitasi oleh hasil-hasil ini. Dengan demikian, proses ini membantu dalam meningkatkan kualitas analisis sentimen dengan membuang informasi yang tidak relevan, sehingga memastikan bahwa hasil analisis lebih akurat.

Tabel 2 Hasil *stopwords*

Content clean	Text stopwords
this is the one ive tried several apps and pro...	this is the one ive tried several apps and pro...
it has a decent variety of exercises it does n...	it has a decent variety of exercises it does n...
most important it s free you can upgrade but i...	most important it s free you can upgrade but i...
i really like the setup and the levels i felt ...	i really like the setup and the levels i felt ...
fantastic app designed to keep you engaged wit...	fantastic app designed to keep you engaged wit...

Tokenizing. Fase selanjutnya adalah *tokenizing*, yang melibatkan pembagian teks menjadi kata-kata individual untuk memfasilitasi proses *stemming* selanjutnya. Pada tahap ini teks ulasan telah dibagi menjadi kata-kata individual, yang memungkinkan untuk melakukan analisis lebih lanjut secara terpisah pada setiap kata. Tabel 3 merupakan hasil *tokenizing*. Dengan melalui tahap *tokenizing*, data siap untuk diproses lebih lanjut dalam analisis sentimen menggunakan algoritma klasifikasi.

Tabel 3 Hasil *tokenizing*

Text stopword	Text tokens
this is the one ive tried several apps and pro...	[this, is, the, one, ive, tried, several, apps...
it has a decent variety of exercises it does n...	[it, has, a, decent, variety, of, exercises, i...
most important it s free you can upgrade but i...	[most, important, it, s, free, you, can, upgra...
i really like the setup and the levels i felt ...	[i, really, like, the, setup, and, the, levels...
fantastic app designed to keep you engaged wit...	[fantastic, app, designed, to, keep, you, engu...

Stemming. *Stemming* merupakan suatu metode yang digunakan untuk menghilangkan imbuhan, preposisi, kata ganti, dan kata keterangan sesuai dengan standar yang ditetapkan oleh KBBI [21]. Prosedur *stemming* sastra Indonesia ini dilakukan dengan bantuan Perpustakaan Sastrawi untuk kepentingan penelitian ini. Ilustrasi hasil dari metode *stemming* dapat dilihat pada Gambar 6. Proses ini membantu dalam mengurangi variasi kata yang sebenarnya memiliki arti yang sama, sehingga memperkuat analisis sentimen dengan memperlakukan kata-kata tersebut secara konsisten.

```

1 : this : this
2 : is : is
3 : the : the
4 : one : one
5 : ive : ive
6 : tried : tried
7 : several : several
8 : apps : apps
9 : and : and
10 : programs : programs
11 : to : to
12 : help : help
13 : me : me
14 : learn : learn
15 : japanese : japanese
16 : but : but

```

Gambar 6 Hasil output *stemming*

Tahap kedua adalah tahap pemisahan data. Saat ini, data dibagi menjadi dua kategori berbeda: data pengujian dan data pelatihan. Data dipartisi menjadi 20 persen dataset pengujian dan 80 persen dataset pelatihan dalam penelitian ini. Dengan membagi dataset menjadi 20 persen untuk pengujian dan 80 persen untuk pelatihan, penelitian ini memastikan bahwa model yang dibangun dapat diuji dengan baik pada data yang tidak terlihat sebelumnya, sementara masih memiliki sejumlah besar data untuk melatih model dengan baik. Hal ini memungkinkan untuk mendapatkan evaluasi yang lebih obyektif terhadap kinerja model dalam memprediksi sentimen ulasan aplikasi Duolingo.

Tahap ketiga adalah tahap pembobotan TF-IDF. Saat ini, pendekatan pembobotan TF-IDF (*Term Frequency-Invers Document Frequency*) digunakan untuk mengubah data tekstual ke dalam format numerik. Kumpulan data memberikan bobot pada setiap kata berdasarkan frekuensinya dalam teks (TF) dan kelangkaannya di seluruh dokumen dalam kumpulan data (IDF). Dengan menggunakan pembobotan TF-IDF, diperoleh numerik yang lebih akurat dari teks ulasan, yang dapat digunakan oleh algoritma klasifikasi untuk analisis sentimen.

Tahap keempat adalah tahap klasifikasi dengan Naive Bayes. Setelah pembobotan TF-IDF diterapkan, dengan menggunakan algoritma klasifikasi Naive Bayes, komputer akan menggunakan data tersebut untuk tujuan pelatihan. Komputer akan menganalisis pola pada data pelatihan untuk membuat prediksi pada data baru berdasarkan pola yang diidentifikasi.

Tahap kelima adalah tahap evaluasi. Selanjutnya, tahap penilaian menjadi sangat penting dalam menilai kemandirian algoritma klasifikasi yang digunakan dalam penelitian. Evaluasi kinerja suatu algoritma dapat dilakukan dengan pemanfaatan metrik seperti akurasi, presisi, dan perolehan. Pendekatan yang dikenal sebagai matriks konfusi digunakan untuk menghitung ukuran ini. Pendekatan ini terdiri dari komponen-komponen berikut: *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN). Temuan keluaran penilaian terlihat pada Gambar 9.

```

MultinomialNB Accuracy: 0.86
MultinomialNB Precision: 0.8541666666666666
MultinomialNB Recall: 0.8733333333333334
MultinomialNB f1_score: 0.8638888888888889
confusion_matrix:
[[85 17]
 [11 87]]
=====
              precision    recall  f1-score   support

negatif         0.88         0.83         0.86         102
positif         0.84         0.89         0.86          98

accuracy         0.86         0.86         0.86         200
macro avg        0.86         0.86         0.86         200
weighted avg     0.86         0.86         0.86         200

```

Gambar 9 Hasil output evaluasi

Formulasi *Confusion Matrix* untuk menghitung akurasi, presisi, dan *recall* adalah sebagai berikut:

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN} = \frac{87 + 85}{87 + 11 + 85 + 17} = \frac{172}{200} = 0,86$$

$$Precision = \frac{TP}{TP + FP} = \frac{87}{87 + 11} = \frac{87}{98} = 0,89$$

$$Recall = \frac{TP}{TP + FN} = \frac{87}{87 + 17} = \frac{87}{104} = 0,83$$

dimana TP memprediksi kata positif dan itu benar, TN memprediksi kata negatif dan itu benar, FP memprediksi kata positif dan itu salah, dan FN memprediksi kata itu negatif dan itu salah.

Analisis sentimen evaluasi Aplikasi Duolingo diklasifikasikan menjadi sentimen positif dan negatif dengan menggunakan metode akurasi. Temuannya menunjukkan akurasi sebesar 0,86 atau 86 persen. Apabila dinyatakan seperti pada Gambar 9, sikap negatif menghasilkan skor presisi sebesar 89 persen dan hasil *recall* sebesar 83 persen. Skor *f1* digunakan untuk menilai kualitas akurasi dan nilai *recall*. Kategorisasi tersebut menghasilkan skor *f1* sebesar 86 persen. Mengenai hasil sentimen positif, tingkat akurasinya adalah 84 persen dan tingkat *recall*nya adalah 89 persen. Skor *f1* digunakan untuk menilai kualitas akurasi dan nilai *recall*. Kategorisasi tersebut menghasilkan skor *f1* sebesar 86 persen.

Simpulan

Berdasarkan penelitian data yang dikumpulkan terdiri dari 1000 data yang diambil dari evaluasi program Duolingo di Google Play Store, seperti yang ditunjukkan oleh temuan penelitian. Sebelum tahap *pre-processing*, pelabelan manual mengungkapkan bahwa ulasan positif lebih banyak daripada ulasan negatif, dengan 500 data menunjukkan sentimen positif dan 500 data menunjukkan sentimen negatif. Selain itu, analisis sentimen mencapai tingkat akurasi sebesar 86 persen. Adapun saran untuk penelitian berikutnya adalah penggunaan algoritma lain untuk mendapatkan hasil perbandingan yang lebih baik, serta menerapkan analisis sentimen dengan menggunakan *tools* seperti *machine learning* atau yang lainnya.

Daftar Pustaka

- [1] 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.
- [2] S. Chohan, A. Nugroho, A. M. B. Aji, and W. Gata, "Analisis Sentimen Pengguna Aplikasi Duolingo Menggunakan Metode Naive Bayes dan Synthetic Minority Over Sampling Technique," *Paradig. - J. Komput. dan Inform.*, vol. 22, no. 2, pp. 139–144, 2020, doi: 10.31294/p.v22i2.8251.
- [3] A. Nurian, "Analisis Sentimen Ulasan Pengguna Aplikasi Google Play Menggunakan Naive Bayes," *J. Inform. dan Tek. Elektro Terap.*, vol. 11, no. 3s1, pp. 829–835, 2023, doi: 10.23960/jitet.v11i3s1.3348.
- [4] E. Indrayuni, "Klasifikasi Text Mining Review Produk Kosmetik Untuk Teks

- Bahasa Indonesia Menggunakan Algoritma Naive Bayes," *J. Khatulistiwa Inform.*, vol. 7, no. 1, pp. 29–36, 2019, doi: 10.31294/jki.v7i1.1.
- [5] A. K. Dewi, "Analisis Sentimen Ekspedisi Sicepat Dari Ulasan Google Play Menggunakan Algoritma Naive Bayes," *JATISI (Jurnal Tek. Inform. dan Sist. Informasi)*, vol. 9, no. 2, pp. 796–805, 2022, doi: 10.35957/jatisi.v9i2.1802.
- [6] L. Ardiani, H. Sujaini, and T. Tursina, "Implementasi Sentiment Analysis Tanggapan Masyarakat Terhadap Pembangunan di Kota Pontianak," *J. Sist. dan Teknol. Inf.*, vol. 8, no. 2, p. 183, 2020, doi: 10.26418/justin.v8i2.36776.
- [7] D. Pratmanto, R. Rousyati, F. F. Wati, A. E. Widodo, S. Suleman, and R. Wijianto, "App Review Sentiment Analysis Shopee Application in Google Play Store Using Naive Bayes Algorithm," *J. Phys. Conf. Ser.*, vol. 1641, no. 1, 2020, doi: 10.1088/1742-6596/1641/1/012043.
- [8] M. K. Khoirul Insan, U. Hayati, and O. Nurdian, "Analisis Sentimen Aplikasi Brimo Pada Ulasan Pengguna Di Google Play Menggunakan Algoritma Naive Bayes," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 7, no. 1, pp. 478–483, 2023, doi: 10.36040/jati.v7i1.6373.
- [9] M. Birjali, M. Kasri, and A. Beni-Hssane, "A comprehensive survey on sentiment analysis: Approaches, challenges and trends," *Knowledge-Based Syst.*, vol. 226, p. 107134, 2021, doi: 10.1016/j.knosys.2021.107134.
- [10] I. Wickramasinghe and H. Kalutara, "Naive Bayes: applications, variations and vulnerabilities: a review of literature with code snippets for implementation," *Soft Comput.*, vol. 25, no. 3, pp. 2277–2293, 2021, doi: 10.1007/s00500-020-05297-6.
- [11] P. Kaviani and S. Dhotre, "International Journal of Advance Engineering and Research Short Survey on Naive Bayes Algorithm," *Int. J. Adv. Eng. Res. Dev.*, vol. 4, no. 11, pp. 607–611, 2017.
- [12] N. A. Rakhmawati, M. I. Aditama, R. I. Pratama, and K. H. U. Wiwaha, "Analisis Klasifikasi Sentimen Pengguna Media Sosial Twitter Terhadap Pengadaan Vaksin COVID-19," *J. Inf. Eng. Educ. Technol.*, vol. 4, no. 2, pp. 90–92, 2020, doi: 10.26740/jieet.v4n2.p90-92.
- [13] S. Lestari and S. Saepudin, "Analisis Sentimen Vaksin Sinovac Pada Twitter Menggunakan Algoritma Naive Bayes," *SISMAITIK (Seminar Nas. Sist. Inf. dan Manaj. Inform.)*, pp. 163–170, 2021.
- [14] A. Imron, "Analisis Sentimen Terhadap Tempat Wisata di Kabupaten Rembang Menggunakan Metode Naive Bayes Classifier," *Tek. Inform.*, pp. 10–13, 2019, [Online]. Available: <https://dspace.uii.ac.id/handle/123456789/14268>

- [15] G. Sanjaya and K. M. Lhaksana, "Lexicon Based)," vol. 7, no. 3, pp. 9698–9710, 2020.
- [16] F. Bei and S. Sudin, "Analisis Sentimen Aplikasi Tiket Online Di Play Store Menggunakan Metode Support Vector Machine (Svm)," *Sismatik*, vol. 01, no. 01, pp. 91–97, 2021.
- [17] J. A. Septian, T. M. Fachrudin, and A. Nugroho, "Analisis Sentimen Pengguna Twitter Terhadap Polemik Persepakbolaan Indonesia Menggunakan Pembobotan TF-IDF dan K-Nearest Neighbor," *J. Intell. Syst. Comput.*, vol. 1, no. 1, pp. 43–49, 2019, doi: 10.52985/insyst.v1i1.36.
- [18] A. Tangkelayuk, "The Klasifikasi Kualitas Air Menggunakan Metode KNN, Naïve Bayes, dan Decision Tree," *JATISI (Jurnal Tek. Inform. dan Sist. Informasi)*, vol. 9, no. 2, pp. 1109–1119, 2022, doi: 10.35957/jatisi.v9i2.2048.
- [19] R. R. Syawal, L. Pratama, T. Wahyuni, and D. G. Passa, "Klasifikasi kualitas Ikan Nilem Berdasarkan Ukuran Menggunakan Algoritma Naive Bayes," *J I M P - J. Inform. Merdeka Pasuruan*, vol. 7, no. 2, p. 72, 2023, doi: 10.51213/jimp.v7i2.514.
- [20] A. Nuzulia, "漢語No Title No Title No Title," *Angew. Chemie Int. Ed.* 6(11), 951–952, no. 2015, pp. 5–24, 1967.
- [21] H. K. Wardana, I. Swanita, and B. W. Yohanes, "Sistem Pemeriksa Pola Kalimat Bahasa Indonesia berbasis Algoritme Left-Corner Parsing dengan Stemming," *J. Nus. Tek. Elektro dan Teknol. Inf.*, vol. 8, no. 3, p. 211, 2019, doi: 10.22146/jnteti.v8i3.515.

LAMPIRAN 4 : CODING

Pengambilan data

```
#install library google play scraper
!pip install google-play-scraper

#import library
from google_play_scraper import app
import pandas as pd
import numpy as np

#scraping data
from google_play_scraper import Sort, reviews

result, continuation_token = reviews(
    'com.duolingo',
    lang='us',
    country='us',
    sort=Sort.MOST_RELEVANT,
    count=250,
    filter_score_with=1
)

#membuat dataframe dari hasil scraping
data = pd.DataFrame(np.array(result),
    columns=['review'])
data = data.join(pd.DataFrame(data.pop('review').tolist()
    ))
data
```

```

#filter hanya ambil score dan content
data = data[['score', 'content', 'at']]
data.head()

data.to_csv("data1.csv", index = False)

#scraping data
from google_play_scraper import Sort, reviews

result, continuation_token = reviews(
    'com.duolingo',
    lang='us',
    country='us',
    sort=Sort.MOST_RELEVANT,
    count=208,
    filter_score_with=2
)

#membuat dataframe dari hasil scraping
data = pd.DataFrame(np.array(result),
columns=['review'])
data = data.join(pd.DataFrame(data.pop('review').tolist()
))
data

#filter hanya ambil score dan content
data = data[['score', 'content', 'at']]
data.head()

```

```

data.to_csv("data2.csv", index = False)

#scraping data
from google_play_scraper import Sort, reviews

result, continuation_token = reviews(
    'com.duolingo',
    lang='us',
    country='us',
    sort=Sort.MOST_RELEVANT,
    count=64,
    filter_score_with=3
)

#membuat dataframe dari hasil scraping
data = pd.DataFrame(np.array(result),
columns=['review'])
data = data.join(pd.DataFrame(data.pop('review').tolist
()))
data

#filter hanya ambil score dan content
data = data[['score', 'content', 'at']]
data.head()

data.to_csv("data3.csv", index = False)

#scraping data
from google_play_scraper import Sort, reviews

result, continuation_token = reviews(

```



```

        'com.duolingo',
        lang='us',
        country='us',
        sort=Sort.MOST_RELEVANT,
        count=120,
        filter_score_with=4
    )

#membuat dataframe dari hasil scraping
data = pd.DataFrame(np.array(result),
columns=['review'])
data =
data.join(pd.DataFrame(data.pop('review').tolist
()))
data

#filter hanya ambil score dan content
data = data[['score', 'content', 'at']]
data.head()

data.to_csv("data4.csv", index = False)

#scraping data
from google_play_scraper import Sort, reviews

result, continuation_token = reviews(
    'com.duolingo',
    lang='us',
    country='us',
    sort=Sort.MOST_RELEVANT,
    count=358,
    filter_score_with=5

```

```
)
```

```
#membuat dataframe dari hasil scraping
data = pd.DataFrame(np.array(result),
columns=['review'])
data =
data.join(pd.DataFrame(data.pop('review').tolist
()))
data
```

```
#filter hanya ambil score dan content
data = data[['score', 'content', 'at']]
data.head()
```

```
data.to_csv("data5.csv", index = False)
```

```
import pandas as pd
data = pd.read_csv('/content/data(2).csv')
data
```

```
#melihat jumlah data tiap score(1-5)
data['score'].value_counts()
```

```
from google.colab import drive
drive.mount('/content/drive')
```

```
#menampilkan jumlah dalam grafik
data.groupby(['score']).size().plot(kind = "bar")
```

Data Preprocessing

```
# import library
import pandas as pd
import string
import nltk

# buat fungsi case folding
def clean_content(content):
    import string, re

    content = content.lower() # menjadikan
lowercase
    content = re.sub("[^a-z]", " ", content) #
hapus semua karakter kecuali a-z
    content = re.sub("\t", " ", content) #
mengganti tab dengan spasi
    content = re.sub("\n", " ", content) #
mengganti new line dengan spasi
    content = re.sub("\s+", " ", content) #
mengganti spasi > 1 dengan 1 spasi
    content = content.strip() # menghapus spasi
di awal dan akhir

    return content

# case folding
data["content_clean"] =
data["content"].apply(clean_content)
data.head()

# stopwords
```

```

import nltk.corpus
nltk.download('stopwords')
from nltk.corpus import stopwords
stop = stopwords.words('indonesian')
data['text_StopWord'] =
data['content_clean'].apply(lambda x: '
'.join([word for word in x.split() if word not in
(stop)]))
data.head(50)

```

```

# tokenizing
import nltk
nltk.download('punkt')
from nltk.tokenize import sent_tokenize,
word_tokenize
data['text_tokens'] =
data['text_StopWord'].apply(lambda x:
word_tokenize(x))
data.head()

```

```

# install library
!pip install Sastrawi

```

```

# stemming
from Sastrawi.Stemmer.StemmerFactory import
StemmerFactory
factory = StemmerFactory()
stemmer = factory.create_stemmer()

```

```

from Sastrawi.Stemmer.StemmerFactory import
StemmerFactory
#import swifter

```

```

# buat stemmer

```

```

factory = StemmerFactory()
stemmer = factory.create_stemmer()

# stemmed
def stemmed_wrapper(term):
    return stemmer.stem(term)

term_dict = {}
hitung=0

for document in data['text_tokens']:
    for term in document:
        if term not in term_dict:
            term_dict[term] = ' '

print(len(term_dict))
print("-----")
for term in term_dict:
    term_dict[term] = stemmed_wrapper(term)
    hitung+=1
    print(hitung, ":", term, ":" , term_dict[term])

print(term_dict)
print("-----")

# penerapan stemmed term ke dataframe
def get_stemmed_term(document):
    return [term_dict[term] for term in document]

#script ini bisa dipisah dari eksekusinya setelah
#pembacaan term selesai
data['text_steamindo'] =
data['text_tokens'].apply(lambda x: '
'.join(get_stemmed_term(x)))
data.head(20)

```

```

# save hasil preprocessing ke csv
data.to_csv('hasil_PreProcessing.csv',      index=
False)

#disini kita importkan library re, kemudian kita
lakukan praproses
import re
def praproses(text):
    text = re.sub('<[^>]*>', '', text)
    emoticons = re.findall('(?:[:|;|=)() (?:-
)?(?:\)|\(|D|P)',
                                text)
    text = (re.sub('[\W]+', ' ', text.lower()) +
            ' '.join(emoticons).replace('-', ' '))
    return text

```

Labeling

```
import nltk
from textblob import TextBlob
all_polarity = 0;
status=[]
total_positif = total_negatif = total = 0

for ulasan in data['content']:
    analysis = TextBlob(ulasan)
    all_polarity += analysis.polarity
    if (analysis.sentiment.polarity > 0.0):
        total_positif += 1
        status.append("Positif")
    elif (analysis.sentiment.polarity <= 0.0):
        total_negatif += 1
        status.append("Negatif")
    total += 1
print(f"Hasil Analisis Data: \nPositif = {total_positif} \nNegatif = {total_negatif}")
print(f"Total Data = {total}")

status = pd.DataFrame({'sentiment' : status})
data['sentiment'] = status

#menghitung jumlah sentimen positif negatif
netral
data['sentiment'].value_counts()

data.groupby(['sentiment']).size().plot(kind =
"bar")
```

Pemisahan Data

```
# memecah data ttest 20% dari keseluruhan
from sklearn.model_selection import
train_test_split
X_train, X_test, y_train, y_test =
train_test_split(data['content'],
data['sentiment'],

                test_size = 0.20,

                random_state = 0)
```


Pembobotan TF-IDF

```
# pembobotan tf-idf
from sklearn.feature_extraction.text import
TfidfVectorizer

tfidf_vectorizer = TfidfVectorizer()
tfidf_train = tfidf_vectorizer.fit_transform(X_train)
tfidf_test = tfidf_vectorizer.transform(X_test)

print(X_train.shape)
print(y_train.shape)
print(X_test.shape)
print(y_test.shape)

from sklearn.feature_extraction.text import
CountVectorizer

vectorizer = CountVectorizer()
vectorizer.fit(X_train)

X_train = vectorizer.transform(X_train)
X_test = vectorizer.transform(X_test)
```

Klasifikasi Naive Bayes

```
from sklearn.naive_bayes import MultinomialNB

nb = MultinomialNB()
nb.fit(tfidf_train, y_train)

X_train.toarray()

y_pred = nb.predict(tfidf_test)
```

Evaluasi

```
from sklearn.metrics import accuracy_score

accuracy = accuracy_score(y_test, y_pred)

import pandas as pd
from sklearn.feature_extraction.text import TfidfVectorizer
from sklearn.naive_bayes import MultinomialNB
from sklearn.metrics import accuracy_score, precision_score, recall_score, f1_score
from sklearn.model_selection import train_test_split
from sklearn.metrics import classification_report
from sklearn.metrics import confusion_matrix

clf = MultinomialNB()
clf.fit(X_train, y_train)
predicted = clf.predict(X_test)

print("MultinomialNB Accuracy:",
accuracy_score(y_test, predicted))
print("MultinomialNB Precision:",
precision_score(y_test, predicted,
average="binary", pos_label="Negatif"))
print("MultinomialNB Recall:",
recall_score(y_test, predicted, average="binary",
pos_label="Negatif"))
print("MultinomialNB f1_score:",
f1_score(y_test, predicted, average="binary",
pos_label="Negatif"))

print(f'confusion_matrix:\n
{confusion_matrix(y_test, predicted)}')
print('=====')
```

```
=====\n')
print(classification_report(y_test,    predicted,
zero_division=0))

# Load dataset
data = pd.read_csv('hasil_PreProcessing.csv')
```

LAMPIRAN 5 : DATASET

Score	Content	AT
5	duolingo makes learning fun and easy with many options to help learn at you pace and level it's easy to use and understand, it's like a game of learning. I would recommend to anyone who wants to learn a language for free and without any delay. You will get bored BUT once you have a good streak and start forgetting to do your lessons and get the the message that you streak will end trust me you will practice to not lose that streak. Having the streak feels really good	5/21/2024 10:44
5	It's great for some languages but no so for others. Overall enjoying it though and my favourite part is the way it's set up (practices, challenges, streaks ect) as it's incredible for keeping consistency. One min a day for a year is much better than doing a few long study sessions.	5/20/2024 10:56
4	I've been using it for a litrelly for more than half of a year .its awesome but only thing that I am irritated is updates after every few days. I mean no problem with updates but they update it without any notification. we have to try a lot . After a long time, I've started understanding that if it doesn't open it needs update it's also not for begginers they don't even explain the answers IYKYK. that's all	4/22/2024 14:58
1	Frustrating since my last update. The app keeps telling me I'm offline when I'm not.	4/4/2024 12:33

	Somehow, in my offline state I can still redo former levels. It's frustrating. Then there's the shutting off when I'm in the middle of a class. Please fix this.	
1	AVOID!!! Not suitable for beginners, only for someone who has basic grammar skills. I really wanted to learn Spanish but after a 800 day streak I have given up as this method of no explanations, just repetitions just doesn't work for me as a beginner and frustrated me. Even the full version doesn't allow you to learn in your own time as you get punished if you don't do it daily and miss out on boosts if you can't learn during their dictated times. No help when stuck/no response on emails. Poor!	4/29/2024 7:34
5	A good way of learning a completely new language for this beginner. Lots of nudges to keep you engaged. Every now and then there are words that are not explained before they are used, but most of the time you can click on them to see what they mean or you soon learn them. The app functions well, I've had no crashes or bugs. Adverts are not too intrusive as they appear between lessons, and there is a way to pay for an advert free version with many other benefits. Yo soy muy feliz con Duolingo.	2/28/2024 7:33
2	The app itself is well laid out and easy to navigate. Not too sure about the teaching style. You learn grammar and sentences before you learn the basics. Seems strange. Wasted 3 hours 'learning' the	3/19/2024 20:05

	same phrases over and over. While repetition works for SOME people, this gets BORING very very fast.	
5	still the best app for language learning at your own pace. it has been years since I had used this app, but I find myself drawn back to learning, and this app holds true to satisfying results and learning fun new words is various ways. Just like with all lessons, redundancy happens, but I feel duolingo has found a way to keep it interesting enough that the feeling of "having too" do this and "I want too" do this is balanced out. Fun/challenging/relaxing âœœðŸŒ¼â~ðŸŒ¼ðŸŒ¼-ðŸŒ¼	5/19/2024 10:03
2	Those new profile drawn things replacing pictures is strange and shouldn't be forced. Especially when you're Paying for using the app. Also would be better if we can learn more useful everyday words and not "ants newspaper" and "foxes donut". A few voice recordings were impossible to understand and sounded very bad. Timed exercises do not extend time when you go from just letters to full on sentences and are impossible to complete on time allocated making it an anxiety induced lesson and not fun	4/10/2024 8:25
1	I used to love Duolingo. I have a 1,383 day streak. I even got Super for two years because I genuinely felt like it was worth it. The big update was annoying at first and there's still some changes I really	3/22/2024 15:59

	don't like, but some things felt more intuitive at the same time. I've stuck with it. However, upon learning that the company has laid off contract workers in favor of using AI to write lessons, I have canceled my Duolingo Super subscription and I will be looking for a new app to learn on.	
--	---	--

DAFTAR RIWAYAT HIDUP

A. Identitas Diri

1. Nama Lengkap : Meli Apriliyani
2. Tempat & tgl Lahir : Demak, 26 April 2003
3. Alamat Rumah : Jetak Rt.03 Rw. 01 Bonang Demak
4. HP : 088238655489
5. Email : melyapril26@gmail.com

B. Riwayat Pendidikan

1. Pendidikan Formal

- a. TK Kinasih Bonang Demak
- b. MI Tsamrotul 1 Bonang Demak
- c. MTs NU Serangan Bonang Demak
- d. SMK N 1 Tuntang Semarang
- e. Universitas Negeri Islam Walisongo

2. Pendidikan Non Formal

- a. Pondok Pesantren Al-Kautsar Bonang Demak

C. Prestasi Akademik

- a. -
- b. -

D. Karya Ilmiah

- a. Buku Mengudara Tak Sempat Bersayap : Alumni Kelas Menulis LPW PMII Koms. UIN Waslisongo Semarang 2022

Hormat Saya

Meli Apriliyani
NIM. 2108096057