

ANALISIS KEPUASAN PENGGUNA APPLE MUSIC PADA GOOGLE PLAY STORE MENGGUNAKAN METODE NLP

Aghies Nokiandi, Irwan Agus Sobari, Sri Rahayu

Informatika, Universitas Bina Sarana Informatika
Jl. Kramat Raya No. 98, Senen, Jakarta Pusat, Indonesia
aghiesn330@gmail.com

ABSTRAK

Apple Music adalah layanan pemutar video dan musik secara online yang dikembangkan oleh *Apple Inc.* dari tahun 2015. *Apple Music* merupakan versi berlangganan setiap bulan dari *iTunes*. Di *Apple Music*, pengguna dapat memutar musik atau video terbaru tanpa harus membayar terus-menerus disetiap konten baru. Dalam mendistribusikan konten yang ada dalam *Apple Music*, *Apple Inc.* merilis aplikasi tersebut diberbagai sistem operasi yaitu *MacOS*, *iOS*, *Android*, *Microsoft Windows*, dan pemutaran melalui *website*. Penelitian ini berfokus pada sistem operasi *Android*. Metode pengumpulan data dari penelitian ini adalah studi pustaka, *web scraping*, dan *aspect-based text classification*. Untuk metode analisa menggunakan *natural processing language* dengan model *IndoBERT*. Data pengumpulan yang didapatkan, bahwa pengguna merasa puas pada bagian “UI/UX” sebesar 4.2%, “Fitur” sebesar 45.4%, “UI/UX,Fitur” sebesar 7.2%, dan “Lainnya” sebesar 43.3%. Dengan *dataset* tersebut, dalam mengklasifikasikan kata, model mengalami kesulitan pada bagian “UI/UX” dan “UI/UX,Fitur”. Hasil *training* yang didapatkan dengan kedua *hyperparameter* berbeda, didapatkan hasil akurasi sebesar 91%. Namun dalam hasil pengujian manual, *hyperparameter* 1 mendapatkan nilai akurasi sebesar 85% dan *hyperparameter* 2 mendapatkan nilai akurasi sebesar 80%.

Kata kunci : Analisis Kepuasan, Natural Language Processing, IndoBERT, Kinerja, Apple Music, Google Play Store

1. PENDAHULUAN

Perkembangan teknologi dan internet sudah semakin pesat. Hanya dengan adanya internet, perangkat yang dipakai semakin fleksibel penggunaannya. Seperti, berkomunikasi, membeli makanan, dan mendengarkan musik. Saat ini, banyak yang sudah beralih ke platform pemutaran musik digital di kota-kota besar. Platform tersebut menawarkan dari gratis dan berbayar dengan fitur-fitur berbeda. Terkhusus *Apple Music* yang hanya menawarkan layanan berbayar.

Apple Music adalah layanan pemutar video dan musik secara online yang dikembangkan oleh *Apple Inc.*. *Apple Music* dibentuk pada tahun 2015. *Apple Music* ini merupakan versi berlangganan dari *iTunes*, dimana *iTunes* menawarkan pembelian musik secara retail. *Apple Music* sendiri menawarkan fitur unggul seperti kualitas suara yang lebih bagus dan karaoke [1].

Apple Music sendiri dapat dipasang pada perangkat *non-Apple*, seperti android yang dipasang melalui *Google Play Store*. Menurut data unduhan dari jurnal [2], telah didapatkan banyaknya pengguna yang mengunduh aplikasi *streaming* musik online, seperti *Spotify* dan *Apple Music*. Dengan jumlah unduhan *Spotify* mencapai 144 Juta dan *Apple Music* 68 Juta.

Sebagai aplikasi yang banyak digunakan setelah *Spotify*, ulasan pengguna dapat memberikan penilaian berdasarkan persepsi dari konsumen. Ulasan-ulasan ini dapat membantu *developer* untuk mengembangkan terus aplikasi tersebut. Pada tanggal 15 Maret 2025, *rating* aplikasi *Apple Music* mencapai 4.3 dengan 514.000 total ulasan. Dengan data ulasan sebesar itu,

tentu akan sulit mengolahnya dengan cara manual. Dengan metode *Natural Language Processing* dengan *pre-trained* model *IndoBERT*, dapat mempermudah pengolahan data. Namun, metode *Natural Language Processing* harus melewati aturan gramatika dan semantiknya, serta mengubah bahasa tersebut menjadi representasi formal yang dapat diproses oleh komputer [3]. Oleh karena itu, analisis kepuasan ini diperlukan. Analisis kepuasan ini akan diolah berdasarkan aspek, yaitu penilaian berdasarkan “UI/UX” dan “Fitur” sebagai label utamanya.

Model *IndoBERT* ini dipilih karena, kelebihan dalam sistematika proses dari *Bidirectional Encoder Representations from Transformers* (BERT), yaitu BERT memahami sebuah kalimat menggunakan sistem *self-attention*. Dengan sistem tersebut, memungkinkan model untuk melihat seluruh data secara bersamaan, kemudian menimbang nilai pada setiap kata terhadap hubungan antar kata. Dari penjelasan latar belakang diatas, maka penulis melakukan penelitian, Analisis Kepuasan Pengguna *Apple Music* pada *Google Play Store* Menggunakan NLP. Dengan hasil penelitian ini, diharapkan memberikan literasi dan panduan dalam penggunaan *Natural Language Processing*. Penelitian ini dilaksanakan untuk mengetahui akurasi dalam menggunakan *Natural Language Processing* dan untuk mengetahui faktor diberikannya ulasan negatif oleh konsumen.

Pada penelitian ini, berasal dari tiga penelitian terdahulu, yaitu. Pertama penelitian yang berjudul, “Analisis Sentimen Pada Ulasan Aplikasi *Home Kredit* dengan metode SVM dan KNN”. Masalah yang

dihadapi dari penelitian ini, yaitu inkonsistensi antara rating dan komentar, kebisingan pada data teks, dan karena kebisingan pada data menjadikan algoritma KNN kurang optimal. Pada penelitian ini bertujuan untuk membantu meninjau dan menguji performa klasifikasi semua ulasan dari aplikasi *Home Credit* pada *Google Play Store*. Kedua penelitian yang berjudul, “Analisis Sentimen Pengguna X Indonesia Terkait Kendaraan Listrik Menggunakan IndoBERT”. Masalah yang dihadapi dari penelitian ini, yaitu kualitas dan konsistensi data pada *Tweet* dan pemahaman konteks yang kompleks pada data *Tweet* yang diambil. Pada penelitian ini bertujuan untuk menguji performa klasifikasi semua pendapat dari sosial media X. Dan ketiga penelitian yang berjudul, “Analisis Sentimen Opini Masyarakat Terhadap Tech Winter Pada Twitter Menggunakan Natural Language Processing”. Masalah yang dihadapi pada penelitian ini, yaitu data teks dengan kata kunci ‘Tech Winter’ dan ‘layoff’ hanya berjumlah sedikit, kebisingan pada data teks, dan terlalu banyaknya metode yang dipakai. Pada penelitian ini bertujuan untuk membandingkan performa pada metode algoritma yang berbeda.

2. TINJAUAN PUSTAKA

2.1. Data Mining

Data mining merupakan kumpulan-kumpulan data yang sangat besar, selain itu dapat menjadi peran pendukung dalam suatu pengambilan keputusan. *Data mining* dapat diproses secara otomatis dan secara manual. Untuk mencari nilai tambah dari suatu basis data, data-data dan pola-pola harus saling dihubungkan. Keduanya dihubungkan untuk melakukan pengelompokan ke dalam satu atau lebih *cluster*, sehingga data-data yang berada dalam satu *cluster* akan mempunyai kesamaan yang tinggi[4].

2.2. Pengelompokan Data Mining

Pengelompokan *data mining* dapat dibagi menjadi beberapa kelompok berdasarkan tugasnya, yaitu Deskripsi, Estimasi, Prediksi, Klasifikasi, *Clustering*, Asosiasi.[4]

2.3. Web Scrapping

Web scraping adalah proses pengguna dalam mengambil informasi yang berasal dari internet secara otomatis. Hasil akhir dari Web Scrapping akan berupa file dengan format *comma-seperated values* (CSV).[5]

2.4. Data Pre-Processing

Data Pre-Processing merupakan tahap sebelum data itu digunakan untuk dianalisis. Data mentah akan disiapkan dan diubah menjadi format yang lebih terstruktur. Selain itu, tahap ini juga memastikan bahwa data siap digunakan secara efektif dalam model pembelajaran mesin. *Data Pre-Processing* yang dilakukan yaitu pembersihan data, *text normalization*, *text standardization*, dan *tokenizing*.[5]

2.5. Analisis Kepuasan

Analisis kepuasan merupakan analisa yang mengukur kualitas barang atau jasa, dimana karakteristik umumnya harus memenuhi syarat dan harapan konsumen. Dalam kepuasan sendiri terdapat faktor bukti langsung, perhatian individu dari karyawan kepada konsumen, kehandalan, dan jaminan.[6]

2.6. UI/UX

User interface (UI) atau antarmuka adalah semua elemen visual yang dipelajari dari bagian perencanaan dan desain. User interface dapat dianalisis dengan memperhatikan bagaimana orang-orang dan komputer itu bekerja sama secara efisien. User experience adalah pengalaman orang-orang dalam berinteraksi dengan komputer sesuai dengan antarmuka yang tersedia. Apakah antarmuka tersebut nyaman dan sesuai dengan kebutuhan, hal itulah yang dapat menjadi umpan balik dalam meningkatkan user interface.[7]

2.7. Fitur

Fitur merupakan semua elemen fungsi yang tersedia atau ditawarkan oleh suatu organisasi atau individu. Contoh fitur yang tersedia pada suatu aplikasi: melacak jam tangan secara real-time; kompatibilitas tinggi dengan perangkat lain; dan dapat mengubah dokumen menjadi gambar.[8]

2.8. Wordcloud

Wordcloud merupakan salah satu cara visualisasi data teks berdasarkan data yang paling sering disebut. Jika semakin sering disebut, semakin besar juga ukuran teksnya. Word cloud biasanya dipakai pada text mining, natural language processing, analisis sosial media dan marketing, dan lainnya.

2.9. IndoBERT

Bidirectional Encoder Representations from Transformers (BERT) adalah sebuah metode untuk melatih mesin dalam memahami bahasa manusia secara paralel, melalui dua sisi data, dan terdapat teknik tertentu untuk menjadi lebih efisien [9]. IndoBERT merupakan data-data BERT namun berfokus pada bahasa Indonesia. IndoBERT berisikan lebih dari 220 juta kata dengan sumber yang menggunakan bahasa Indonesia baku, seperti dari website resmi dan koran

2.10. Natural Language Processing (NLP)

Natural language processing merupakan pembelajaran mesin yang berfokus pada bahasa manusia dalam memahami, menafsirkan, dan menghasilkan bahasa alami manusia. *Natural language processing* juga merupakan salah satu bagian dari *artificial intelligence*. [10].

2.11. Analisis Sentimen Pada Ulasan Aplikasi *Home Credit* dengan Metode SVM dan KNN

Penelitian ini dilakukan oleh Adiansyah dan Wahyudin, dengan berfokus pada analisis sentimen dari aplikasi *Home Credit*. Data yang diambil untuk penelitian ini, yaitu sebanyak 2845, dan menggunakan metode SVM dan KNN. Hasil akurasi yang didapatkan, yaitu SVM sebesar 88% dan KNN sebesar 79%. Berdasarkan hasil tersebut, SVM dapat lebih baik dalam melakukan analisis sentimen.

2.12. Analisis Sentimen Pengguna X Indonesia Terkait Kendaraan Listrik Menggunakan IndoBERT

Penelitian ini dilakukan oleh Merdiansah, Siska, dan Ridha, dengan berfokus pada analisis sentimen kendaraan listrik pada platform sosial media X. Penelitian ini menggunakan IndoBERT sebagai alat untuk mengevaluasi performa model. Selain itu, digunakan juga dataset bawaan oleh IndoNLU. Dengan hasil jika menggunakan dataset bawaan, didapatkan akurasi sebesar 100% pada *training* dan 98% pada validasi. Sedangkan tanpa dataset bawaan, didapatkan akurasi sebesar 90% pada *training* dan 82% pada validasi. Hal ini menunjukkan hasil lebih baik pada menggunakan dataset bawaan IndoNLU.

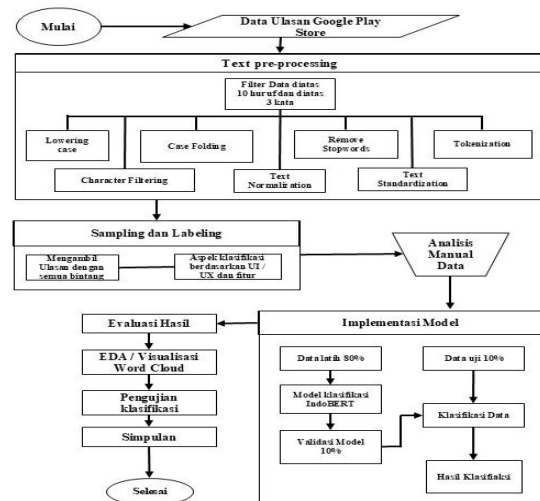
2.13. Analisis Sentimen Opini Masyarakat Terhadap Tech Winter Pada Twitter Menggunakan Natural Language Processing

Penelitian ini dilakukan oleh Fauzianto dan Supatman, dengan berfokus pada analisis sentimen *Tech Winter* dari platform sosial media X. Penelitian ini menggunakan beragam model, yaitu Regresi Logistik, SVM, *Random Forest*, *Neural Network*, dan

Naive Bayes. Hasil akurasi evaluasi yang didapatkan, yaitu regresi logistik mendapatkan nilai sebesar 83%, SVM mendapatkan nilai sebesar 94%, *Random Forest* mendapatkan nilai sebesar 78%, *Neural Network* mendapatkan nilai sebesar 89%, dan *Naive Bayes* mendapatkan nilai sebesar 94%.

3. METODE PENELITIAN

3.1. Kerangka Alur Penelitian



Gambar 1. Kerangka alur penelitian

- Pengambilan data, karena data mining selalu belajar dari data yang sudah ada [11]. Dataset yang diambil merupakan data sekunder, dengan data sebanyak 3051. dan dalam jangka waktu dari 11 September 2016 hingga 14 Juni 2025.

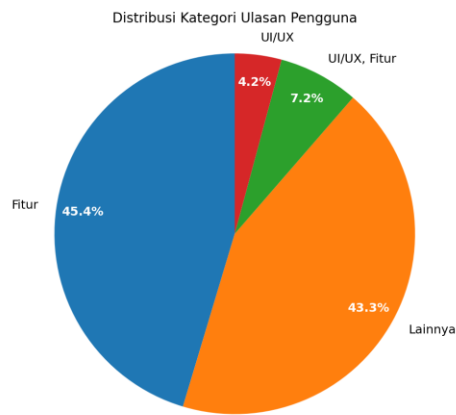
Tabel 1. Contoh dataset mentah

userName	score	at	content
id_00009	3	2025-06-09 06:49:30	BAGUSSSS BANGETTTT, aplikasi musik yang bagus selain spo**** karna katanya boikot?? jadi better pake ini aja duluu, tapi sayangnyaaa agak error gituu, musiknya ga bisa diputarr, udah di pencet pencet, pencet yang putar acaknya gabisa jugaa, ditunggu sekitar seharian masih gitu aja, jadii yaudah diapus teruss download lagi😁😁
id_00010	5	2025-06-07 19:40:05	waw
id_00011	3	2025-06-06 06:30:46	force close app when entering homepage
id_00012	3	2025-06-05 05:18:12	buruk
...			

- Text pre-processing*, ditahap ini telah dilakukan *filter data*, *lowering case*, *character filtering*, *case folding*, *text normalization*, *removing stopwords*, *text standardization*, and *tokenization*. [11].
- Sampling dan labelling*, sampling data dipilih dengan semua rating atau score. Untuk labeling menggunakan teknik *Aspect Based Sentiment Analysis (ABSA)*, dengan label 'UI/UX' dan 'Fitur'. [12]

Tabel 2. Contoh hasil *text pre-processing*

content	content_normalized_v3	tokens
force close app when entering homepage	paksa tutup aplikasi ketika memasuki layar utama	['paksa', 'tutup', 'aplikasi', 'ketika', 'memasuki', 'layar', 'utama']



Gambar 3. Grafik diagram label dataset

- d. Analisis manual data, dilakukan untuk memantau kembali data yang sudah diproses pada *text pre-processing*, *sampling*, dan *labeling*. Hal ini bertujuan untuk memeriksa kembali apakah ada kesalahan data atau duplikasi data. Terutama pada proses *labeling* dan standarisasi teks.
- e. Implementasi model, model yang dipilih yaitu model IndoBERT, dengan pembagian dataset menjadi 80% untuk data latih, 10% untuk data validasi, dan 10% untuk data uji. Dan berjalan selama 5 *epoch*.
- f. Evaluasi, disini akan mengeluarkan *classification report*, *confusion matrix*, grafik garis, dan *word cloud*. Kemudian memperhatikan hasil akhir dari training, yaitu *precision*, *recall*, *f1-score*, dan terutama besarnya akurasi yang didapatkan. Setelah dirasa sudah optimal, akan dilanjutkan dengan prediksi dengan data dari peneliti. Jika nilai akurasi dari prediksi tersebut memenuhi syarat, penelitian ini selesai.

4. HASIL DAN PEMBAHASAN

Memuat hasil, Pengujian dan pembahasan tentang skripsi yang telah dilakukan

4.1. Split Data

Split data pada dataset dilakukan dengan ketentuan 80% data akan digunakan untuk dataset latih, 10% untuk dataset validasi, dan 10% untuk dataset uji. Memberikan dataset latih lebih banyak, dapat membuat model yang digunakan semakin akurat. Akurat yang ditujukan yaitu belajar dalam mengenali pola, membuat prediksi, dan mengklasifikasikan data.

Tabel 3. Jumlah Split Data

Jumlah data <i>training</i>	1490
Jumlah data validasi	186
Jumlah data <i>testing</i>	187

4.2. Implementasi Model

Implementasi ini menggunakan IndoBERT sebagai landasan awal untuk pembentukan sebuah model. Dengan menggunakan "indobert-large-p2" dan

library yang digunakan yaitu *AutoTokenizer*, *AutoModelForSequenceClassification*, *TrainingArguments*, and *Trainer*. Penelitian ini menggunakan dua *hyperparameter*, untuk pengaturan *hyperparameter* pertama sesuai dengan tabel 4 dan hasil yang didapatkan sesuai dengan gambar 4. Untuk *hyperparameter* kedua sesuai dengan tabel 5 dan gambar 5

Tabel 4. Konfigurasi Hyperparameter 1

<i>per_device_eval_batch_size</i>	8
<i>per_device_train_batch_size</i>	8
<i>gradient_accumulation_steps</i>	0
<i>fp16</i>	false

Tabel 5. Konfigurasi Hyperparameter 2

<i>per_device_eval_batch_size</i>	2
<i>per_device_train_batch_size</i>	2
<i>gradient_accumulation_steps</i>	4
<i>fp16</i>	true

```
--- Classification Report Epoch Ini ---
```

	precision	recall	f1-score	support
UI/UX	1.00	0.78	0.88	9
Fitur	0.94	0.99	0.97	85
UI/UX, Fitur	1.00	0.64	0.78	14
Lainnya	0.96	1.00	0.98	78
accuracy			0.96	186
macro avg	0.98	0.85	0.90	186
weighted avg	0.96	0.96	0.95	186

Gambar 4. Hasil epoch lima hyperparameter 1

```
--- Classification Report Epoch Ini ---
```

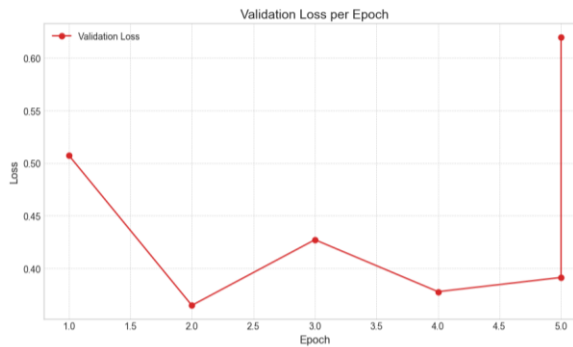
	precision	recall	f1-score	support
UI/UX	1.00	0.56	0.71	9
Fitur	0.91	0.99	0.95	85
UI/UX, Fitur	1.00	0.50	0.67	14
Lainnya	0.94	0.99	0.96	78
accuracy			0.93	186
macro avg	0.96	0.76	0.82	186
weighted avg	0.93	0.93	0.92	186

Gambar 5. Hasil epoch lima hyperparameter 2

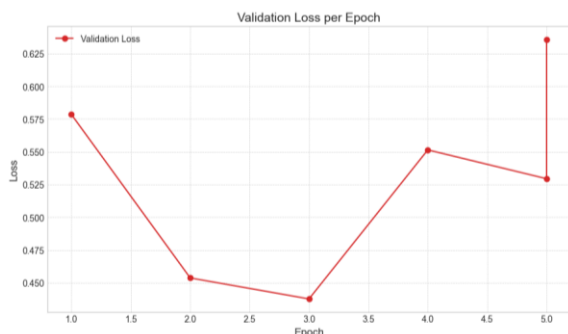
Hasil dari gambar 4 dan gambar 5 menunjukkan bahwa dengan menggunakan *hyperparameter* 1 akan membuat model lebih sensitif dalam mendeteksi pola kalimat, daripada menggunakan *hyperparameter* 2. Hal ini ditunjukkan dengan nilai akurasi yang mencapai 96%, serta nilai *precision*, *recall*, dan *f1-support* yang lebih tinggi.

4.3. Grafik garis

Gambar 6 dan gambar 7 menunjukkan bahwa *validation loss* di kedua *hyperparameter* hanya terdapat sedikit perbedaan. Namun, terlihat lebih stabil pada *hyperparameter* 1.

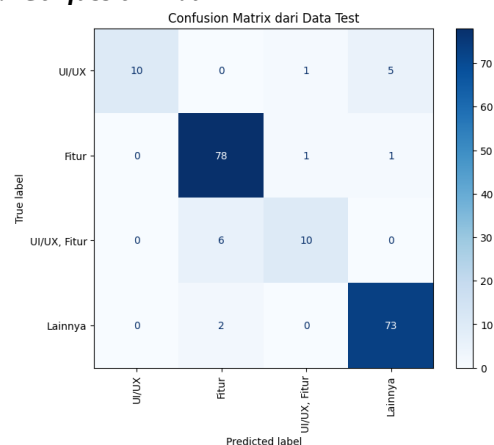


Gambar 6. Grafik garis hyperparameter 1

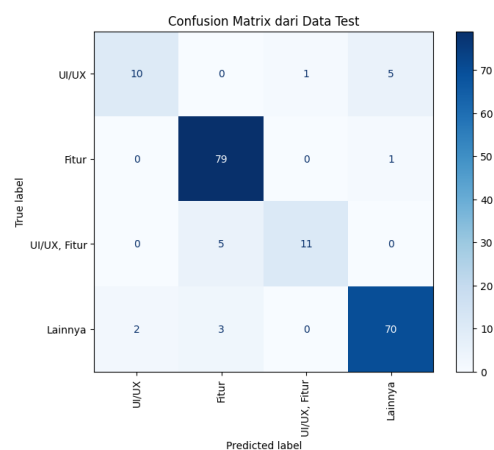


Gambar 7. Grafik garis hyperparameter 2

4.4. Confusion Matrix



Gambar 8. Confusion matrix hyperparameter 1

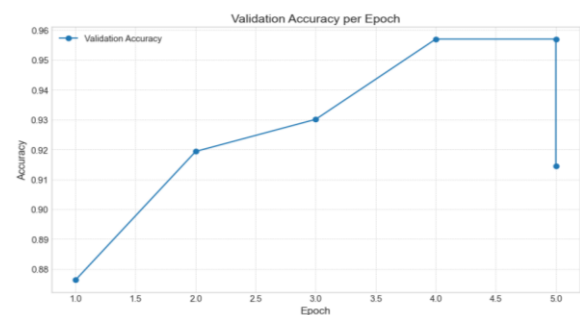


Gambar 9. Confusion matrix hyperparameter 2

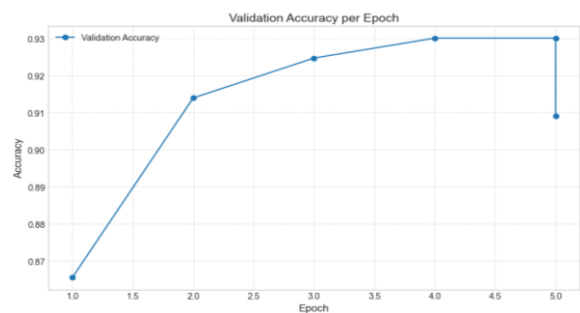
Pada gambar 8 dan gambar 9, menunjukkan bahwa performa dengan *hyperparameter* 1 lebih baik dalam mendeteksi label “Lainnya”. Sedangkan performa dengan *hyperparameter* 2, lebih baik dalam mendeteksi label “Fitur” dan “UI/UX, Fitur”.

4.5. Validation Accuracy

Pada gambar 10 dan gambar 11, menunjukkan bahwa hasil *Validation Accuracy* lebih tinggi pada *hyperparameter* 1. Namun, kedua *hyperparameter* menunjukkan peningkatan akurasi validasi di setiap *epoch* yang dilalui.

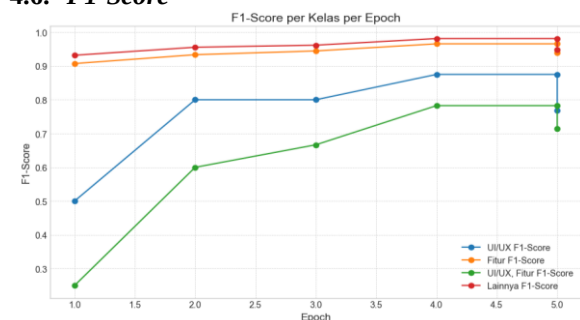


Gambar 10. Validation hyperparameter 1

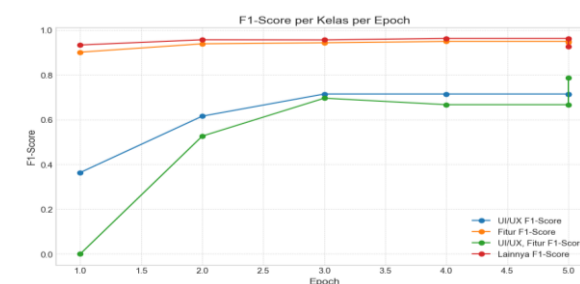


Gambar 11. Validation hyperparameter 2

4.6. F1-Score



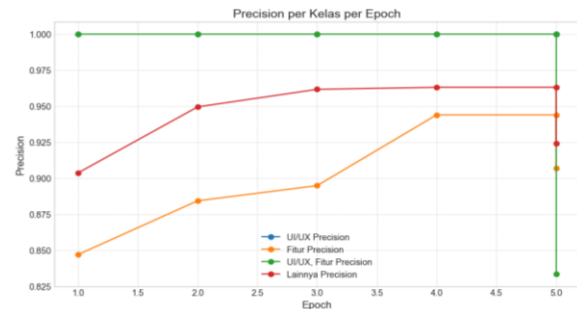
Gambar 12. F1-score hyperparameter 1



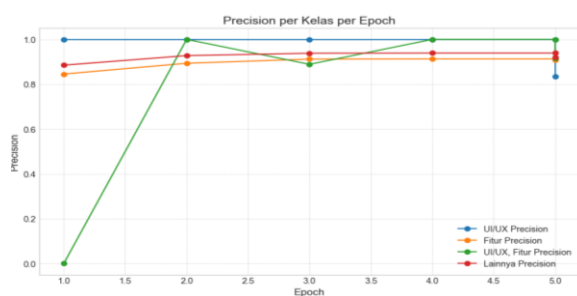
Gambar 13. F1-score hyperparameter 2

Pada gambar 12 dan 13 menunjukkan bahwa, saat mendeteksi pola di label “UI/UX” dan “UI/UX, Fitur” di *hyperparameter* 1 lebih sensitif daripada *hyperparameter* 2.

4.7. Precision



Gambar 14. Precision *hyperparameter* 1



Gambar 15. Precision *hyperparameter* 2

Pada gambar 14 dan gambar 15, Precision di *hyperparameter* 1 mendapat nilai sempurna pada label “UI/UX” dan “UI/UX, Fitur”, sedangkan *hyperparameter* 2 hanya pada label “UI/UX”.

4.8. Wordcloud



Gambar 16. Wordcloud *hyperparameter* 1



Gambar 17. Wordcloud *hyperparameter* 2

Pada hasil Wordcloud, *hyperparameter* 1 dapat mendeteksi perkataan lebih sensitif. Seperti pada kata “android” dan “mengunduh”.

4.9. Pengujian Manual

Pengujian manual ini dilakukan dengan menggunakan dataset langsung dari peneliti. Dataset yang diberikan berisikan 20 baris kalimat beserta label aslinya. Pengujian manual ini dilakukan pada tahap terakhir di setiap model yang telah melakukan *training*. Setelah pengujian ini akan diberikan *classification report* dan *confussion matrix*.

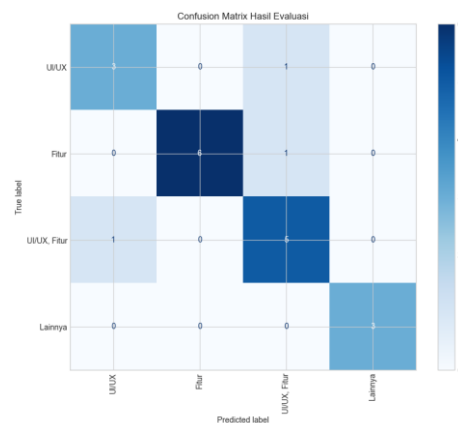
Tabel 1. Contoh Dataset dari Peneliti

komentar	label_asli
Tampilan aplikasinya bagus dan bersih gampang dipake.	UI/UX
Untuk download offlinenya sering error dan lambat.	Fitur
Secara desain sih oke tapi beberapa fungsi penting malah hilang di update terbaru.,	UI/UX, Fitur
Aplikasinya lumayan lah buat sehari-hari nggak ada yang spesial.	Lainnya
...	

4.10. Pengujian manual *hyperparameter* 1

Laporan Klasifikasi Detail:				
	precision	recall	f1-score	support
UI/UX	0.75	0.75	0.75	4
Fitur	1.00	0.86	0.92	7
UI/UX, Fitur	0.71	0.83	0.77	6
Lainnya	1.00	1.00	1.00	3
accuracy			0.85	20
macro avg	0.87	0.86	0.86	20
weighted avg	0.86	0.85	0.85	20

Gambar 18. Report manual *hyperparameter* 1

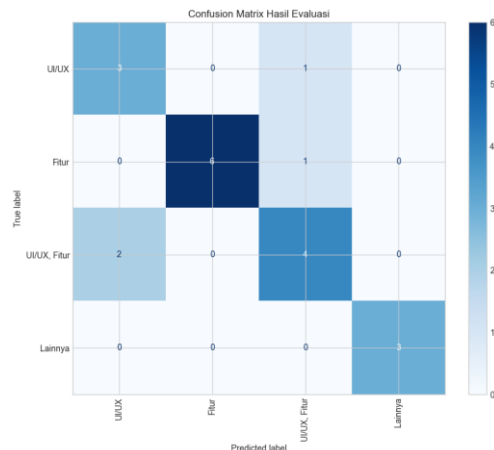


Gambar 19. Confussion matrix manual *hyperparameter* 1

4.11. Pengujian manual *hyperparameter* 2

Laporan Klasifikasi Detail:				
	precision	recall	f1-score	support
UI/UX	0.60	0.75	0.67	4
Fitur	1.00	0.86	0.92	7
UI/UX, Fitur	0.67	0.67	0.67	6
Lainnya	1.00	1.00	1.00	3
accuracy			0.80	20
macro avg	0.82	0.82	0.81	20
weighted avg	0.82	0.80	0.81	20

Gambar 20. Report manual *hyperparameter* 2



Gambar 21. Confusion matrix manual hyperparameter 2

Pada pengujian manual di kedua *hyperparameter*, *hyperparameter 1* lebih baik dalam mengenali label “UI/UX, Fitur”. Selain itu, nilai *accuracy*, *precision*, *recall*, dan *f1-score* lebih tinggi.

5. KESIMPULAN DAN SARAN

Analisis kepuasan pengguna menggunakan data yang berasal dari *Google Play Store*. Dengan mempertimbangkan aspek yang dapat dinilai dalam suatu aplikasi, didapatkan ‘UI/UX’ dan ‘Fitur’. Kedua label tersebut digunakan untuk memberikan *labelling manual* saat awal memulai *text pre-processing*. Dengan persentase yang didapatkan berdasarkan *labeling manual*, yaitu “UI/UX” sebesar 4.2%, “Fitur” sebesar 45.4%, “UI/UX,Fitur” sebesar 7.2%, dan “Lainnya” sebesar 43.3%. Dengan diterapkan model dasar IndoBERT sebagai landasan, hasil akurasi yang didapatkan bisa berada diatas 75%. Saran yang bisa dipertimbangkan untuk penelitian berikutnya, yaitu penggabungan antara *web scraping* dengan survey menggunakan metode *decision tree*, menggabungkan model lain, dan penggunaan *multi-label classification*.

DAFTAR PUSTAKA

- [1] Apple Inc., “Apple Music,” Apple, Jul. 2025. [Online]. Available: <https://www.apple.com/apple-music/>
- [2] Di. Widiagusti and T. A. Napitupulu, “856-Article Text-2601-1-10-20231107,” *Economics and Digital Business Review*, no. Volume 5 Issue 1 (2024) Pages 325-342, 2023, doi: doi.org/10.37531/ecotal.v5i1.856.
- [3] R. A. Fauzianto and Supatman, “Analisis Sentimen Opini Masyarakat Terhadap Tech Winter Pada Twitter Menggunakan Natural Language Processing,” *Jurnal Syntax Admiration*, vol. 3, no. 9, pp. 1577–1585, Sep. 2023, doi: [10.46799/jsa.v3i9.909](https://doi.org/10.46799/jsa.v3i9.909).
- [4] F. Zahra, M. A. Ridla, and N. Azise, “IMPLEMENTASI DATA MINING MENGGUNAKAN ALGORITMA APRIORI DALAM MENENTUKAN PERSEDIAAN BARANG (STUDI KASUS: TOKO SINAR HARAHAP),” *JUSTIFY: Jurnal Sistem Informasi Ibrahimy*, vol. 3, no. 1, pp. 55–65, Jul. 2024, doi: [10.35316/justify.v3i1.5335](https://doi.org/10.35316/justify.v3i1.5335).
- [5] A. Ardiansyah and Wahyudin, “Analisis Sentimen Pada Ulasan Aplikasi Home Credit dengan Metode SVM dan KNN,” *Jurnal Komputer Antartika*, no. Vol. 1 No. 4 (2023): Desember 2023, Oct. 2023, doi: <https://doi.org/10.70052/jka.v1i4.50>.
- [6] F. Fitrayana and K. Rizal, “PENERAPAN ALGORITMA C4.5 UNTUK MENENTUKAN KEPUASAN PELANGGAN DALAM KUALITAS PELAYANAN PENGIRIMAN BARANG (Studi Kasus : PT Nusantara Card Semesta Cabang Kemanggisian),” *Bianglala Informatika*, vol. 12, no. 1, pp. 1–8, Aug. 2024, doi: [10.31294/bi.v12i1.17263](https://doi.org/10.31294/bi.v12i1.17263).
- [7] M. Iqbal Sain, M. Asep Rizkiawan, M. A. Maulana Rahmat, and M. Sidik, “OPTIMALISASI PENGALAMAN PENGGUNA: REDESIGN UI/UX WEBSITE SIMAKIP UHAMKA DENGAN METODE DESIGN THINKING,” *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 1, Jan. 2025, doi: [10.23960/jitet.v13i1.5479](https://doi.org/10.23960/jitet.v13i1.5479).
- [8] A. D. Rawat, “Enhancing University-based Organic Product Initiatives through Digital Innovation: a Case Study of Dev Bhoomi Uttarakhand University,” *International Journal For Multidisciplinary Research*, vol. 6, no. 5, Oct. 2024, doi: [10.36948/ijfmr.2024.v06i05.28866](https://doi.org/10.36948/ijfmr.2024.v06i05.28866).
- [9] A. Warstadt et al., “Findings of the BabyLM Challenge: Sample-Efficient Pretraining on Developmentally Plausible Corpora,” in *Proceedings of the BabyLM Challenge at the 27th Conference on Computational Natural Language Learning*, Stroudsburg, PA, USA: Association for Computational Linguistics, Apr. 2023, pp. 1–6. doi: [10.18653/v1/2023.conll-babylm.1](https://doi.org/10.18653/v1/2023.conll-babylm.1).
- [10] R. Gupta, D. Srivastava, M. Sahu, S. Tiwari, R. K. Ambasta, and P. Kumar, “Artificial intelligence to deep learning: machine intelligence approach for drug discovery,” *Mol Divers*, vol. 25, no. 3, pp. 1315–1360, Aug. 2021, doi: [10.1007/s11030-021-10217-3](https://doi.org/10.1007/s11030-021-10217-3).
- [11] O. Maimon and L. Rokach, *Data Mining and Knowledge Discovery Handbook*, 2nd Edition. Boston, MA: Springer US, 2022. doi: [10.1007/978-0-387-09823-4](https://doi.org/10.1007/978-0-387-09823-4).
- [12] W. Zhang, X. Li, Y. Deng, L. Bing, and W. Lam, “A Survey on Aspect-Based Sentiment Analysis: Tasks, Methods, and Challenges,” *IEEE Trans Knowl Data Eng*, vol. 35, no. 11, pp. 11019–11038, Nov. 2023, doi: [10.1109/TKDE.2022.3230975](https://doi.org/10.1109/TKDE.2022.3230975).