

Sistem Rekomendasi Keaslian Barang Berdasarkan Klasifikasi Ulasan Tokopedia Menggunakan Metode *LSTM*

1st Gabriel Azarya Krishna Lokajaya
Program Studi S1 Informatika
Universitas Telkom Surabaya
Kota Surabaya, Indonesia
gabrielazarya@student.telkomuniversity.ac.id

2nd Alqis Rausanfitra, S.Kom., M.Kom.
Program Studi S1 Informatika
Universitas Telkom Surabaya
Kota Surabaya, Indonesia
alqisfita@telkomuniversity.ac.id

3rd Pima Hani Safitri, S.Kom., M.Kom.
Program Studi S1 Informatika
Universitas Telkom Surabaya
Kota Surabaya, Indonesia
phanisafitri@telkomuniversity.ac.id

Abstrak — Perkembangan *marketplace* membuat pengguna semakin bergantung pada ulasan untuk melihat keaslian suatu produk. Namun, ulasan yang panjang dan tidak terstruktur sering membuat proses penilaian menjadi sulit. Penelitian ini bertujuan membangun sistem rekomendasi keaslian barang berdasarkan analisis ulasan Tokopedia menggunakan metode *Long Short-Term Memory (LSTM)*. Proses penelitian dimulai dari pengambilan data ulasan, *preprocessing* tanpa *stopword removal*, serta penerapan teknik *oversampling* untuk menyeimbangkan kelas sentimen. Setelah itu dilakukan pelabelan manual dan pelatihan model dengan variasi *hyperparameter* melalui 10-fold *cross-validation*. Model terbaik diperoleh pada kombinasi $K = 10$, *fold* ke-10, *epoch* 20, *batch size* 16, dan *dropout* 0.3. Model ini menghasilkan akurasi 89.50%, presisi 87.15%, *recall* 86.34%, dan *F1-score macro* 84.41% pada data pengujian. Hasil tersebut menunjukkan bahwa model mampu memberikan prediksi yang cukup stabil dan dapat digunakan untuk melihat kecenderungan keaslian produk berdasarkan sentimen ulasan. Sistem rekomendasi yang dibangun diharapkan dapat membantu pengguna agar lebih mudah dalam mempertimbangkan kualitas dan keaslian produk sebelum membeli.

Kata kunci — Sistem Rekomendasi, Keaslian Produk, Tokopedia, *LSTM*, Analisis Ulasan

I. PENDAHULUAN

Perkembangan *e-commerce* di Indonesia mengalami peningkatan yang signifikan dan menjadi bagian penting dalam perekonomian digital. Pertumbuhan tersebut mendorong perubahan perilaku konsumen dalam melakukan transaksi jual beli secara daring melalui platform *marketplace*. Tokopedia sebagai salah satu *marketplace* terbesar di Indonesia menyediakan fitur rating dan ulasan konsumen yang berperan penting dalam membantu calon pembeli dalam mengambil keputusan pembelian [1], [2].

Rating dan ulasan konsumen memiliki pengaruh yang signifikan terhadap tingkat kepercayaan serta keputusan pembelian pengguna pada platform *e-commerce*. Ulasan positif dapat meningkatkan kepercayaan konsumen,

sedangkan ulasan negatif cenderung menurunkan minat beli terhadap suatu produk [1], [3]. Oleh karena itu, ulasan konsumen menjadi sumber data yang bernilai untuk dianalisis lebih lanjut, khususnya dalam mengidentifikasi kualitas dan keaslian barang yang dijual di *marketplace*.

Seiring dengan meningkatnya transaksi daring, peredaran barang palsu atau *counterfeit* juga menjadi permasalahan yang serius. Konsumen sering mengalami kesulitan dalam membedakan barang asli dan barang palsu karena keterbatasan informasi visual serta deskripsi produk yang tidak selalu mencerminkan kondisi sebenarnya [4]. Dalam kondisi tersebut, ulasan konsumen menjadi indikator penting yang mencerminkan pengalaman nyata pembeli terhadap produk yang diterima.

Ulasan konsumen pada *marketplace* umumnya berbentuk teks tidak terstruktur dan menggunakan bahasa yang beragam, sehingga memerlukan pendekatan *Natural Language Processing (NLP)* (*NLP*) agar dapat dianalisis secara otomatis [5]. Melalui penerapan *NLP*, informasi penting yang terkandung dalam ulasan konsumen dapat diekstraksi dan dimanfaatkan sebagai dasar pengambilan keputusan.

Dalam pengolahan teks berbasis urutan kata, metode *Recurrent Neural Network (RNN)* (*RNN*) dan *Long Short-Term Memory (LSTM)* (*LSTM*) banyak digunakan karena kemampuannya dalam menangkap konteks dan dependensi jangka panjang pada data sekuensial. Berbagai penelitian menunjukkan bahwa *LSTM* memiliki performa yang baik dalam klasifikasi teks, khususnya pada analisis sentimen dan ulasan pengguna [6], [7], [8]. Berdasarkan permasalahan tersebut, penelitian ini mengusulkan sistem rekomendasi keaslian barang berdasarkan klasifikasi ulasan Tokopedia menggunakan metode *LSTM* guna membantu konsumen dalam menentukan keaslian produk sebelum melakukan pembelian.

II. KAJIAN TEORI

Kajian teori disusun untuk memberikan landasan konseptual dan teoritis yang relevan dengan penelitian yang

dilakukan. Pembahasan difokuskan pada penelitian terdahulu yang berkaitan dengan analisis ulasan konsumen, konsep keaslian barang, serta penerapan *Natural Language Processing*, *word embedding*, dan metode *Long Short-Term Memory* dalam klasifikasi teks. Landasan teori ini digunakan sebagai acuan dalam perancangan metode penelitian dan analisis hasil yang diperoleh.

A. Penelitian Terdahulu

Penelitian oleh Tri Hermanto, Setyanto, dan Luthfi menunjukkan bahwa kombinasi *Long Short-Term Memory* dan *Convolutional Neural Network* dengan *Word2Vec* mampu meningkatkan performa klasifikasi sentimen teks berita berbahasa Indonesia dibandingkan model tunggal, meskipun masih dibatasi oleh ukuran dataset yang kecil [9]. Sementara itu, penelitian oleh Amelia dan Santoso membuktikan bahwa metode *LSTM* sangat efektif dalam klasifikasi teks panjang pada domain sinopsis film dengan tingkat akurasi yang sangat tinggi, sehingga menegaskan kemampuan *LSTM* dalam memahami konteks sekuensial teks secara mendalam [10].

Penelitian tinjauan sistematis yang dilakukan oleh Roy dan Dutta mengungkapkan bahwa sistem rekomendasi berbasis *deep learning* dan pendekatan hibrid semakin mendominasi berbagai domain aplikasi, meskipun masih menghadapi tantangan seperti *cold start* dan ketidakseimbangan data [11]. Di sisi lain, penelitian oleh Arwindarti, Setiawan, dan Imron menunjukkan bahwa penerapan *Word2Vec* dan *LSTM* pada data opini publik media sosial mampu menghasilkan performa klasifikasi sentimen yang stabil, menegaskan efektivitas *LSTM* dalam pengolahan teks tidak terstruktur [7].

Penelitian oleh Huang menekankan pentingnya representasi semantik dalam sistem rekomendasi berbasis konten melalui integrasi *word embedding* untuk meningkatkan akurasi rekomendasi [12]. Selanjutnya, penelitian oleh Ichwanto, Fredlina, dan Putra membuktikan bahwa kombinasi *RNN-LSTM* dan *Synthetic Minority Oversampling Technique (SMOTE)* efektif dalam klasifikasi teks opini dengan data tidak seimbang, meskipun performa dapat menurun saat diuji pada data baru [13]. Kedua penelitian ini memperkuat relevansi penggunaan *LSTM*, *word embedding*, dan teknik penyeimbangan data dalam penelitian sistem rekomendasi keaslian barang berbasis ulasan konsumen.

B. Barang Asli dan Barang Palsu

Barang palsu atau *counterfeit* merupakan produk tiruan yang menyerupai barang bermerek asli namun tidak diproduksi oleh pemegang merek resmi. Fenomena jual-beli barang palsu dipengaruhi oleh faktor harga yang lebih murah dan rendahnya kesadaran konsumen terhadap risiko pembelian barang tidak asli [3]. Dalam konteks *marketplace*, perbedaan barang asli dan palsu sulit dikenali secara visual sehingga diperlukan pendekatan berbasis data.

C. Natural Language Processing (NLP)

Natural Language Processing merupakan cabang kecerdasan buatan yang berfokus pada interaksi antara komputer dan bahasa manusia. Pendekatan *NLP* memungkinkan pengolahan teks tidak terstruktur menjadi data yang dapat dianalisis secara komputasional. Dalam

penelitian ini, *NLP* digunakan untuk memproses ulasan konsumen Tokopedia sebelum dilakukan klasifikasi keaslian barang [4].

D. Preprocessing Teks

Tahap *preprocessing* bertujuan untuk meningkatkan kualitas data teks dengan menghilangkan *noise* serta menyeragamkan format tulisan. Proses *preprocessing* meliputi *case folding*, *cleaning*, *tokenization*, *normalization*, dan *stemming*. Tahapan ini digunakan untuk meningkatkan performa model klasifikasi [4].

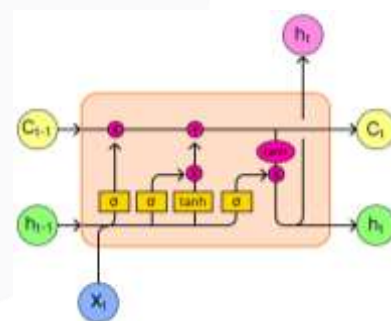
E. Word Embedding

Word embedding merupakan teknik representasi kata ke dalam bentuk vektor numerik yang merepresentasikan hubungan semantik antar kata. Setiap kata dipetakan ke dalam ruang vektor berdimensi tertentu sehingga kata-kata dengan makna serupa memiliki representasi yang berdekatan [8].

Penggunaan *word embedding* banyak diterapkan dalam penelitian klasifikasi teks dan analisis sentimen karena mampu mempertahankan konteks makna kata dibandingkan pendekatan berbasis frekuensi. Dalam penelitian ini, *word embedding* digunakan sebagai masukan utama bagi model *LSTM* [8].

F. Long Short-Term Memory (LSTM)

Long Short-Term Memory merupakan pengembangan dari *Recurrent Neural Network* yang dirancang untuk mengatasi permasalahan *vanishing gradient* pada pemrosesan data sekuensial. Arsitektur *LSTM* memiliki mekanisme memori jangka panjang yang dikendalikan oleh beberapa gerbang (*gate*), yaitu *forget gate*, *input gate*, *cell state*, dan *output gate*. Mekanisme ini memungkinkan model mempertahankan informasi penting dan mengabaikan informasi yang tidak relevan dalam urutan data teks [5], [6]. Arsitektur *LSTM* ditunjukkan pada GAMBAR 1.



GAMBAR 1

(ARSITEKTUR LONG SHORT TERM MEMORY)

Gerbang pertama adalah *forget gate* yang berfungsi menentukan informasi mana dari *cell state* sebelumnya yang perlu dipertahankan atau dibuang, sebagaimana dirumuskan pada Persamaan (1).

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f) \quad (1)$$

Nilai f_t berada pada rentang 0 hingga 1, di mana nilai mendekati 0 menunjukkan informasi yang dilupakan,

sedangkan nilai mendekati 1 menunjukkan informasi yang dipertahankan.

Selanjutnya, *input gate* mengatur informasi baru yang akan disimpan ke dalam *cell state*. Proses ini melibatkan dua tahap, yaitu perhitungan nilai gerbang input pada Persamaan (2) dan kandidat memori baru pada Persamaan (3).

$$i_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_i) \quad (2)$$

$$\tilde{c}_t = \tanh(W_c \cdot [h_{t-1}, x_t] + b_c) \quad (3)$$

Nilai i_t menentukan seberapa besar informasi baru $\tilde{c}_t$ akan ditambahkan ke memori.

Pembaruan *cell state* dilakukan dengan mengombinasikan informasi lama dan informasi baru, sebagaimana ditunjukkan pada Persamaan (4).

$$c_t = f_t \cdot c_{t-1} + i_t \cdot \tilde{c}_t \quad (4)$$

Persamaan ini memungkinkan LSTM mempertahankan konteks penting dalam jangka panjang.

Gerbang terakhir adalah *output gate* yang menentukan keluaran (*hidden state*) dari LSTM pada setiap waktu, sebagaimana dirumuskan pada Persamaan (5) dan Persamaan (6).

$$o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o) \quad (5)$$

$$h_t = o_t \cdot \tanh(c_t) \quad (6)$$

Nilai h_t digunakan sebagai representasi sekuensial yang menjadi masukan bagi lapisan selanjutnya atau proses klasifikasi.

III. METODE

Bab ini menjelaskan metode penelitian yang digunakan untuk membangun sistem rekomendasi keaslian barang berdasarkan analisis ulasan konsumen Tokopedia. Metode penelitian disusun secara sistematis agar setiap tahapan saling terintegrasi dan dapat dipertanggungjawabkan secara ilmiah. Penjelasan pada bab ini menjadi dasar utama dalam memahami proses eksperimen, analisis hasil, serta pembahasan yang disajikan pada bab selanjutnya.

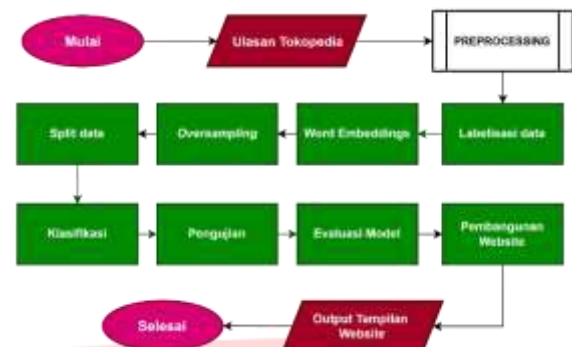
A. Pendekatan dan Jenis Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan metode eksperimen untuk menguji kemampuan model dalam mengklasifikasikan ulasan konsumen berdasarkan indikasi keaslian barang. Pendekatan kuantitatif dipilih karena penelitian berfokus pada pengukuran performa model menggunakan metrik evaluasi yang bersifat numerik. Metode eksperimen digunakan untuk mengamati pengaruh penerapan *Long Short-Term Memory* terhadap hasil klasifikasi ulasan konsumen secara terkontrol.

B. Perancangan Sistem

Perancangan sistem menggambarkan alur kerja sistem rekomendasi keaslian barang secara menyeluruh, mulai dari data masukan hingga keluaran akhir berupa hasil klasifikasi dan rekomendasi. Alur sistem ini dirancang untuk memastikan bahwa setiap tahapan pemrosesan data dilakukan secara berurutan, terintegrasi, dan sesuai dengan

tujuan penelitian. Gambaran umum perancangan sistem ditunjukkan pada GAMBAR 2 sebagai representasi visual dari keseluruhan proses yang dijalankan.



GAMBAR 2
(PERANCANGAN SISTEM)

Berdasarkan pada GAMBAR 2, sistem diawali dengan proses pengumpulan data berupa ulasan teks konsumen Tokopedia melalui teknik *web scraping*. Data yang diperoleh masih bersifat mentah dan mengandung berbagai variasi bahasa serta noise, sehingga perlu melalui tahap *preprocessing* teks. Tahap ini berfungsi untuk membersihkan dan menormalisasi teks agar siap diproses oleh model pembelajaran mesin.

Setelah tahap *preprocessing*, data ulasan dilabeli secara manual ke dalam kelas barang asli dan barang palsu. Data yang telah dilabeli kemudian melalui proses penyeimbangan kelas untuk mengatasi ketimpangan distribusi data. Selanjutnya, data teks direpresentasikan ke dalam bentuk vektor numerik menggunakan teknik *word embedding* agar dapat dipahami oleh model *Long Short-Term Memory*. Output dari model *LSTM* berupa hasil klasifikasi digunakan sebagai dasar dalam sistem rekomendasi keaslian barang, sehingga alur sistem dari input hingga output dapat berjalan secara logis dan konsisten.

C. Pengumpulan Data

Pengambilan data ulasan komentar dilakukan melalui proses *scraping* 50 produk pada bulan Juli - Agustus 2025 di Tokopedia yang akan dijadikan objek klasifikasi sistem rekomendasi dalam penelitian ini. Data yang diperoleh disimpan dalam bentuk dataset teks mentah *csv* untuk diproses pada tahap selanjutnya sesuai alur sistem yang telah dirancang. Hasil *Scraping* ditunjukkan pada TABEL 1 berikut.

TABEL 1
(DATA KOMENTAR HASIL *SCRAPPING*)

No.	Komentar
1.	kalo ingin pakek kaos kaki tebal saran up satu size, kalo kaos kaki tipis size normal aja
2.	pengiriman cepat, kualitas sepatu bagus, jahitannya rapih. ukuran sepatu akurat
3.	Terimakasih kak sepatunya bagus banget sukaaaa 🥰🥰🥰
4.	Barangnya bagus ga nyesel beli disiniiiii 🥰
5.	Maaf baru kasi ulasan hpnya top banget mantap. Mendarat dengan selamat dalam kondisi 🙌🙌🙌 InsyaAllah kalau ada rezeki order lagi, terimakasih

D. Preprocessing Teks

Tahap *preprocessing* teks bertujuan untuk meningkatkan kualitas data dengan menghilangkan noise dan menyeragamkan struktur teks ulasan. Proses yang dilakukan meliputi *case folding*, *cleaning*, *tokenization*, *normalization*, dan *stemming*. Pada penelitian ini, proses *stopword removal* tidak diterapkan agar konteks kalimat tetap terjaga, mengingat makna ulasan konsumen sering kali bergantung pada struktur kalimat secara utuh seperti pada GAMBAR 3.



GAMBAR 3
(ALUR PREPROCESSING)

Setelah melalui tahapan preprocessing, teks ulasan yang semula tidak terstruktur menjadi lebih bersih dan seragam sehingga siap untuk diproses pada tahap selanjutnya. Hasil perubahan teks sebelum dan sesudah preprocessing ditampilkan pada TABEL 2, yang menunjukkan bahwa kata-kata tidak relevan, simbol, dan variasi penulisan berhasil diminimalkan tanpa menghilangkan makna utama ulasan.

TABEL 2
(CONTOH HASIL PREPROCESSING)

Step	Sebelum di proses	Sesudah di proses
Data Cleaning	"Barang ini!! ASLI banget... #recommended"	"Barang ini ASLI banget recommended"
Case Folding	"Barang Ini ASLI Banget recommended"	"barang ini asli banget recommended"
Tokenization	"barang ini asli banget recommended"	[barang, ini, asli, banget, recommended]
Normalization	[barang, ini, asli, banget, recommended]	[barang, ini, asli, sekali, direkomendasikan]
Stemming	[barang, ini, asli, sekali, direkomendasikan]	[barang, ini, asli, sekali, rekomendasi]

E. Pelabelan dan Penyeimbangan Data

Pelabelan data dilakukan secara manual dengan membaca dan memahami konteks keseluruhan ulasan konsumen seperti pada TABEL 3 dengan menggunakan dua label, yaitu "asli" dan "palsu". Pendekatan ini dipilih karena indikasi barang asli atau palsu sering kali tersirat dan tidak selalu dapat diidentifikasi hanya dari kata kunci tertentu. Setelah proses pelabelan, dilakukan penyeimbangan data menggunakan teknik *SMOTE* untuk menghindari bias model terhadap kelas mayoritas.

TABEL 3
(LABELISASI DATA)

Komentar	Label
barang asli dan direkomendasikan untuk beli bagus	Asli
	Asli

harga mahal sesuai kualitas sangat bagus dan <i>original</i>	Asli
bau lem menyengat	Palsu
jelek	Palsu

F. Word embedding

Data ulasan yang telah diproses dan diseimbangkan kemudian direpresentasikan ke dalam bentuk vektor numerik menggunakan teknik *word embedding*. Representasi ini memungkinkan sistem untuk menangkap hubungan semantik antar kata dalam satu ulasan. Hasil *word embedding* menjadi masukan utama bagi model *Long Short-Term Memory* dalam proses pembelajaran pola teks ulasan konsumen.

G. Klasifikasi dan Pengujian Model LSTM

Model *Long Short-Term Memory* digunakan untuk mengklasifikasikan ulasan konsumen ke dalam kategori barang asli dan barang palsu. Model ini dipilih karena kemampuannya dalam menangkap dependensi jangka panjang pada data sekuensial teks. Proses pelatihan dilakukan dengan variasi *hyperparameter* yang diuji meliputi jumlah *epoch* (10, 20, 30), ukuran *batch* (16, 32, 64), dan nilai *dropout* (0.2, 0.3, 0.5). Penelitian ini juga menggunakan *k fold cross validation*, sehingga dalam pelatihan dan pengujian model memakai beberapa nilai *k* yaitu *k* (3, 5, 10). Kombinasi dari parameter-parameter tersebut dicoba secara sistematis untuk menemukan konfigurasi optimal yang memberikan hasil evaluasi terbaik berdasarkan metrik *accuracy*, *precision*, *recall*, dan *F1-score*.

H. Evaluasi dan Keterkaitan dengan Hasil Pembahasan

Evaluasi performa model dilakukan untuk mengetahui sejauh mana performa model telah mengenali pola dan belajar dari dataset. Evaluasi dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score*, serta dianalisis melalui *confusion matrix*. Hasil evaluasi ini secara langsung menjadi dasar pembahasan pada bab hasil dan pembahasan, di mana setiap metrik yang disajikan merefleksikan kinerja sistem sesuai alur yang telah dirancang.

IV. HASIL DAN PEMBAHASAN

Bab ini menyajikan hasil implementasi metode *Long Short-Term Memory (LSTM)* dalam mengklasifikasikan ulasan konsumen Tokopedia untuk mendukung sistem rekomendasi keaslian barang. Pembahasan difokuskan pada proses pengolahan dataset, evaluasi performa model, serta analisis hasil terbaik yang diperoleh dari pengujian yang dilakukan. Hasil eksperimen dianalisis secara kuantitatif menggunakan metrik evaluasi yang relevan dan disajikan dalam bentuk tabel dan visualisasi untuk memperjelas interpretasi hasil. Dengan demikian, bab ini menjadi dasar dalam menilai efektivitas metode yang diusulkan dalam menjawab tujuan penelitian.

A. Pengolahan Dataset dan Penerapan Model

Total data hasil *scrapping* ulasan yang di dapat berjumlah 12.252 data dari total 50 produk. Dataset diproses melalui tahapan *preprocessing* yang meliputi *case folding*, *cleaning*, *tokenization*, *normalization*, dan *stemming* untuk menghilangkan noise serta menyeragamkan struktur teks. Setelah melalui *preprocessing*, data ulasan diberi label secara manual oleh tiga *annotator* independen untuk dijadikan *ground truth* dalam pelatihan model. Proses ini

melibatkan pembacaan setiap ulasan yang telah dibersihkan dan penentuan label berdasarkan kandungan makna teks. Suatu ulasan diberi label "Asli" jika mengandung indikasi pujian terhadap kualitas atau keaslian, sedangkan label "Palsu" diberikan jika ulasan mengandung keluhan atau indikasi pemalsuan, dengan hasil akhir berdasarkan kesepakatan ketiga penilai. Hasil akhir label disepakati dari voting label yang paling dominan. Hasil kesepakatan pelabelan datanya seperti pada

TABEL 4 berikut.

TABEL 4
(HASIL LABELISASI DATA)

<i>cleaned</i>	<i>Annotator 1</i>	<i>Annotator 2</i>	<i>Annotator 3</i>	<i>label</i>
saya belum dapat paket itu mungkin salah kirim	asli	asli	palsu	asli
pesan tidak sesuai ijo stabilo ukur l datang warna item ukur xl	asli	palsu	asli	asli
real	asli	asli	asli	asli
hp nya dah nyampe lumayan bagus hanya beda warna psn yang hitam dtng yang merah	palsu	asli	asli	asli
nyesel beli hp ini wa nya g bisa d pake	palsu	palsu	palsu	palsu

Setelah *preprocessing* dan labelisasi, data teks direpresentasikan ke dalam bentuk numerik menggunakan teknik *word embedding*. Setiap kata dipetakan ke dalam vektor berdimensi tertentu sehingga hubungan semantik antar kata dapat dipertahankan. Representasi vektor ini kemudian menjadi masukan utama bagi model *Long Short-Term Memory (LSTM)*, yang dirancang untuk memproses data sekuensial dan menangkap dependensi jangka panjang dalam ulasan konsumen.

Untuk mengatasi permasalahan ketidakseimbangan jumlah data antar kelas, penelitian ini menerapkan metode *Synthetic Minority Oversampling Technique (SMOTE)*. Penerapan *SMOTE* dengan komposisi *balancing data* 1:1 antara palsu dan juga asli. Setelah seluruh tahapan pengolahan selesai, masuk ke dalam proses pelatihan dan pengujian model *LSTM*. Proses pelatihan dilakukan dengan variasi *hyperparameter* yang diuji meliputi jumlah *epoch*, *batch size*, dan nilai *dropout*.

B. Analisis Performa Model

Analisis performa model dilakukan untuk mengetahui pengaruh variasi parameter pelatihan terhadap hasil klasifikasi. Beberapa kombinasi parameter diuji, meliputi variasi jumlah *epoch* (10, 20, 30), ukuran *batch* (16, 32, 64), dan nilai *dropout* (0.2, 0.3, 0.5). Kemudian *k fold cross validation* yang memakai beberapa nilai *k* yaitu *k* (3, 5, 10). Setiap kombinasi parameter dievaluasi menggunakan metrik akurasi, presisi, *recall*, dan *F1-macro* (nilai rata-rata dari *F1-score* label "asli" dan "palsu"). Hasil pengujian performa model untuk setiap kombinasi parameter ditampilkan pada TABEL 5.

TABEL 5
(Tabel Klasifikasi Hasil Kombinasi)

Kombinasi	<i>K</i>	<i>Fold</i>	<i>Epoch</i>	<i>Batch</i>	<i>Dropout</i>	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Macro</i>
21	3	1	30	16	0.5	84.83%	88.35%	86.97%	85.38%
	3	2	30	16	0.5	88.98%	87.21%	86.27%	84.41%
	3	3	30	16	0.5	85.83%	87.64%	86.27%	84.59%

Kombinasi	<i>K</i>	<i>Fold</i>	<i>Epoch</i>	<i>Batch</i>	<i>Dropout</i>	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Macro</i>
47	5	1	30	16	0.3	87.58%	89.14%	88.25%	86.67%
	5	2	30	16	0.3	81.26%	86.43%	82.91%	81.50%
	5	3	30	16	0.3	86.27%	85.95%	85.10%	83.01%
	5	4	30	16	0.3	89.02%	88.10%	86.82%	85.18%
	5	5	30	16	0.3	86.87%	88.88%	87.77%	86.19%
65	10	1	20	16	0.3	84.62%	88.03%	86.25%	84.72%
	10	2	20	16	0.3	87.69%	87.64%	86.83%	85%
	10	3	20	16	0.3	84.09%	88.31%	86.06%	84.65%
	10	4	20	16	0.3	86.85%	87.84%	86.35%	84.73%
	10	5	20	16	0.3	86%	86.97%	84.62%	83.06%
	10	6	20	16	0.3	86.95%	87.90%	87.29%	85.41%
	10	7	20	16	0.3	84.94%	87.84%	85.95%	84.41%
	10	8	20	16	0.3	83.88%	89.02%	87.77%	86.24%
	10	9	20	16	0.3	86.42%	88.69%	87.20%	85.68%
	10	10	20	16	0.3	89.50%	87.15%	86.34%	84.41%

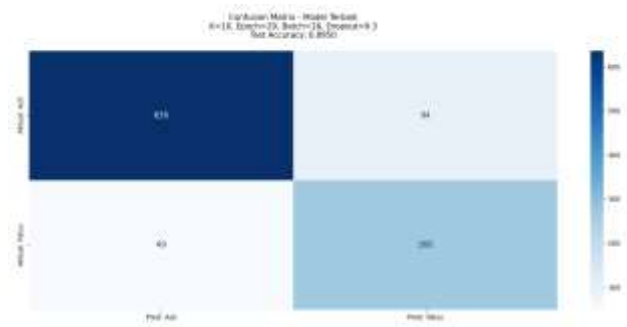
Berdasarkan hasil pengujian yang telah dilakukan, tidak seluruh kombinasi parameter ditampilkan dalam jurnal ini. Untuk menjaga keringkas dan fokus analisis, hanya kombinasi parameter yang bersifat representatif dan relevan yang disajikan, yaitu kombinasi 21, kombinasi 47, dan kombinasi 65, ditampilkan pada TABEL 5. Ketiga kombinasi tersebut dipilih karena mewakili hasil terbaik pada masing-masing variasi nilai *K* dan menunjukkan perbedaan performa model yang signifikan.

Hasil pada TABEL 5 menunjukkan bahwa peningkatan nilai *K* pada metode *cross-validation* memberikan pengaruh terhadap stabilitas performa model. Kombinasi 21 dan kombinasi 47 menghasilkan performa yang cukup baik, namun kombinasi 65 menunjukkan peningkatan akurasi dan konsistensi metrik evaluasi yang lebih tinggi. Hal ini mengindikasikan bahwa konfigurasi parameter pada kombinasi 65 mampu meningkatkan kemampuan generalisasi model dalam mengklasifikasikan ulasan konsumen.

C. Analisis Hasil Terbaik

Analisis hasil terbaik difokuskan pada kombinasi parameter yang menghasilkan performa paling optimal berdasarkan seluruh rangkaian pengujian. Berdasarkan hasil evaluasi, kombinasi 65 dipilih sebagai model terbaik karena menghasilkan nilai performa tertinggi dibandingkan kombinasi lainnya. Ringkasan hasil evaluasi model terbaik ditampilkan pada TABEL 5.

Berdasarkan TABEL 5, model dengan kombinasi 65 menghasilkan nilai akurasi sebesar 89.50%, presisi 87.15%, *recall* 86.34%, dan *F1-macro* 84.41%. Nilai tersebut menunjukkan bahwa model mampu melakukan klasifikasi ulasan konsumen dengan tingkat ketepatan yang tinggi serta keseimbangan performa antar kelas. Untuk melihat distribusi hasil klasifikasi secara lebih rinci, dilakukan analisis menggunakan *confusion matrix* yang ditampilkan pada GAMBAR 4.



GAMBAR 4
(CONFUSION MATRIX)

Hasil *confusion matrix* menunjukkan *confusion matrix* dari model terbaik pada sistem rekomendasi keaslian barang berdasarkan klasifikasi ulasan Tokopedia menggunakan metode *LSTM*. Model terbaik diperoleh melalui kombinasi ke-65 pada *fold* kesepuluh dengan konfigurasi parameter $K = 10$, *epoch* 20, *batch size* 16, dan *dropout* 0.3, menghasilkan nilai akurasi pengujian sebesar 0.8950 atau 89.50%. Berdasarkan hasil tersebut, model berhasil memprediksi 636 data aktual dengan label “Asli” secara benar dan hanya mengalami kesalahan pada 94 data yang secara keliru diklasifikasikan sebagai “Palsu”. Sementara itu, pada data dengan label aktual “Palsu”, model mampu mengklasifikasikan 268 data secara tepat, dengan 49 data lainnya terdeteksi sebagai “Asli”.

V. KESIMPULAN

Simpulan harus diuraikan dalam bentuk paragraf yang berisi poin utama pembahasan hasil penelitian, berupa uraian dan tidak boleh menggunakan pointer. Penelitian ini berhasil menerapkan metode *Long Short-Term Memory (LSTM)* untuk membangun sistem klasifikasi otomatis yang merekomendasikan keaslian barang berdasarkan ulasan pengguna di Tokopedia. Sistem dikembangkan melalui tahapan pengumpulan data, *preprocessing* teks, pembentukan *word embedding*, pelabelan data, serta pelatihan model menggunakan *K-Fold Cross Validation*.

Hasil pengujian menunjukkan bahwa model *LSTM* memiliki performa klasifikasi yang baik dengan akurasi sebesar 89.50%. Model terbaik diperoleh pada kombinasi ke-65, yang menghasilkan nilai *precision* 87.15%, *recall* 86.34%, dan *F1-score macro* 84.41%. Meskipun demikian, model masih mengalami keterbatasan dalam mengklasifikasikan ulasan yang bersifat ambigu atau memiliki sentimen campuran.

Untuk pengembangan penelitian selanjutnya, disarankan dilakukan optimalisasi pada tahap *preprocessing* teks, khususnya dalam menangani bahasa tidak baku dan ulasan ambigu. Selain itu, penerapan teknik penanganan ketidakseimbangan data serta eksplorasi model berbasis

pretrained language models seperti *IndoBERT* atau arsitektur *transformer* diharapkan dapat meningkatkan pemahaman konteks bahasa dan performa klasifikasi. Pengujian pada skala data dan platform e-commerce yang lebih luas juga perlu dilakukan untuk mengukur kemampuan generalisasi model.

REFERENSI

- [1] R. F. Abdillah and A. N. Pramesti, “Dampak Rating dan Ulasan Konsumen Terhadap Keputusan Pembelian di e-Commerce,” *SEMINAR NASIONAL AMIKOM SURAKARTA (SEMNAS) 2024*, 2024.
- [2] R. Fahrozi, D. Rahmawati, V. Muldani, and M. Saddam, “The influence of online customer review on trust and its implications for purchasing decisions on the Tokopedia marketplace,” *Jurnal Administrare: Jurnal Pemikiran Ilmiah dan Pendidikan Administrasi Perkantoran*, 2022, [Online]. Available: <http://creativecommons.org/licenses/by/4.0/>
- [3] A. A. Saputri, “Rational Choice Jual-Beli Barang Counterfeit,” *Journal of Economic Business Ethic and Science Histories*, vol. I, 2023.
- [4] F. Rumaisa, Y. Puspitarani, A. Rosita, A. Zakiah, and S. Violina, “Penerapan Natural Language Processing (NLP) Di Bidang Pendidikan,” *Jurnal Inovasi Masyarakat*, 2021.
- [5] R. A. Pramunendar, D. P. Prabowo, and R. A. Megantara, “Metode Recurrent Neural Network (RNN) Dengan Arsitektur LSTM Untuk Analisis Sentimen Opini Publik Terkait Vaksin Covid-19,” *Jurnal Informatika Upgris*, vol. 8, no. 1, 2022.
- [6] O. V. Putra, A. Musthafa, and K. P. Wibowo, “Klasifikasi Ekspresi Teks Berbahasa Jawa Menggunakan Algoritma Long Short Term Memory,” *Komputika : Jurnal Sistem Komputer*, vol. 10, no. 2, pp. 137–143, Aug. 2021, doi: 10.34010/komputika.v11i1.4616.
- [7] T. Arwindarti, E. I. Setiawan, and S. Imron, “Klasifikasi Sentimen Opini Publik Pada Instagram Pemerintah Kabupaten Bojonegoro Menggunakan LSTM,” *Teknika*, vol. 13, no. 1, pp. 1–9, Nov. 2023, doi: 10.34148/teknika.v13i1.699.
- [8] N. F. Basri and E. Utami, “Sistemasi: Jurnal Sistem Informasi Penerapan Model Word2Vec dan LSTM dalam Analisis Sentimen Ulasan Pengguna Mobile legends Application of Word2Vec and LSTM Models in Sentiment Analysis of Mobile Legends User Reviews,” *Sistemasi: Jurnal Sistem Informasi*, vol. 14, no. 2, pp. 2540–9719, 2025, [Online]. Available: <http://sistemasi.ftik.unisi.ac.id>
- [9] D. Tri Hermanto, A. Setyanto, and T. L. Emha, “Algoritma LSTM-CNN untuk Sentimen Klasifikasi dengan Word2vec pada Media Online,” *Citec Journal*, 2021.
- [10] R. Amelia and D. B. Santoso, “Prediksi Genre Film Dengan Klasifikasi Multi Kelas Sinopsis Menggunakan Jaringan LSTM,” *Journal of Information Technology and Computer Science (INTECOMS)*, 2023, [Online]. Available: <https://www.researchgate.net/publication/375187872>

- [11] D. Roy and M. Dutta, "A systematic review and research perspective on recommender systems," *J Big Data*, vol. 9, no. 1, Dec. 2022, doi: 10.1186/s40537-022-00592-5.
- [12] R. Huang, "Improved content recommendation algorithm integrating semantic information," *J Big Data*, vol. 10, no. 1, Dec. 2023, doi: 10.1186/s40537-023-00776-7.
- [13] R. A. Ichwanto, K. Q. Fredlina, and I. G. J. E. Putra, "Klasifikasi Kategori Feedback EDOM Primakara University dengan Algoritma RNN LSTM," 2025.

