

Klasifikasi Tingkat Urgensi Pengaduan Masyarakat Menggunakan *Long Short-Term Memory* (Lstm) (Studi Kasus: Website Lapak Aduan Banyumas)

1st Kelvin Fauzian Setiawan
Fakultas Informatika
Telkom University Purwokerto,
Purwokerto, Indonesia
kelvinfauzian@students.telkomuniversity.ac.id

2nd Paradise, S.Kom., M.Kom.
Fakultas Informatika
Telkom University Purwokerto,
Purwokerto, Indonesia
paradise@telkomuniversity.ac.id

3rd Dedy Agung Prabowo, S.Kom.,
M.Kom.
Fakultas Informatika
Telkom University Purwokerto,
Purwokerto, Indonesia
dedyprabowo@telkomuniversity.ac.id

Abstrak — Penggunaan website pengaduan masyarakat sudah diterapkan di berbagai daerah, termasuk Lapak Aduan Banyumas milik Pemerintah Kabupaten Banyumas. Website ini memfasilitasi masyarakat untuk menyampaikan informasi, keluhan, pertanyaan, dan usulan terkait pelayanan daerah. Peningkatan jumlah aduan pada Desember 2023 (950 aduan) dan Januari 2024 (1.032 aduan) memperlambat proses penanganan, dan tanpa sistem skala prioritas penanganan menjadi tidak tepat sasaran. Penelitian ini bertujuan untuk mengimplementasikan sistem klasifikasi tingkat urgensi pengaduan masyarakat menggunakan metode *Long Short-Term Memory* (LSTM). Pada website Lapak Aduan Banyumas, pengaduan yang masuk diklasifikasikan ke dalam kategori “*Urgent*” berlaku jika kondisinya sangat memprihatinkan dan berdampak serius dan “*Not Urgent*” berlaku jika kondisinya serius tetapi tidak memerlukan penanganan segera. Penelitian ini menggunakan data yang diambil dari website Lapak Aduan Banyumas, terdiri dari 19.240 data pengaduan selama periode 2 Januari 2024 hingga 27 April 2025. Proses penelitian meliputi tahapan *Preprocessing Data*, Pelabelan Data, *Feature Extraction*, Pembuatan Model Klasifikasi, dan *Deployment*. Berdasarkan hasil penelitian, model berhasil mengklasifikasikan pengaduan ke dalam dua kategori dengan rata-rata akurasi sebesar 99.51%. Nilai *precision*, *recall*, dan *F1-score* juga tinggi dan seimbang di kedua kategori, yaitu 0.99 untuk *Urgent* maupun *Not Urgent*. *Macro* dan *weighted average* sebesar 0.99 menunjukkan bahwa model mampu menangani kedua kategori dengan sangat konsisten.

Kata kunci — Lapak Aduan Banyumas, Pengaduan, LSTM, Urgent, Not Urgent

I. PENDAHULUAN

Pelayanan publik mencakup berbagai aspek penting dalam kehidupan bernegara, di mana pemerintah berperan menyediakan layanan untuk kesejahteraan masyarakat [1]. Salah satu bentuknya adalah website pengaduan masyarakat, seperti Lapak Aduan Banyumas milik Pemerintah Kabupaten Banyumas, Jawa Tengah. Website ini memfasilitasi

partisipasi masyarakat dalam menyampaikan keluhan, informasi, pertanyaan, dan usulan, serta meningkatkan kepercayaan publik dan mendukung evaluasi kinerja pelayanan daerah [2].

Data pengaduan yang masuk melalui website Lapak Aduan Banyumas tercatat sebanyak 950 aduan pada bulan Desember 2023 dan 1.032 aduan pada bulan Januari 2024. Ketika jumlah aduan meningkat, penanganan menjadi tidak efisien, memakan waktu lebih lama, serta rentan terhadap keterlambatan dalam respons. Hal ini dapat mengakibatkan penundaan dalam penanganan aduan yang dilaporkan oleh masyarakat. Ditambah ketiadaan sistem yang mampu menentukan skala prioritas membuat semua aduan ditangani secara manual dan berurutan, tanpa mempertimbangkan tingkat urgensi serta potensi dampaknya. Akibatnya, aduan yang lebih mendesak bisa tertunda atau bahkan terabaikan, sementara aduan dengan tingkat kepentingan lebih rendah justru mendapat perhatian yang sama. Untuk mengatasi permasalahan ini salah satu pendekatan yang dapat digunakan adalah dengan menerapkan *natural language processing* atau sering disingkat sebagai NLP [3].

Natural Language Processing (NLP) adalah metode untuk membangun model komputasi yang memungkinkan interaksi antara manusia dan mesin melalui bahasa alami. NLP memberikan pendekatan khusus dalam pengolahan teks. Pada NLP terdapat salah satu metode yang dapat digunakan yaitu *Long Short-Term Memory* (LSTM) yang merupakan variasi dari permodelan *Recurrent Neural Network* atau sering disingkat sebagai RNN [3]. *Long Short-Term Memory* (LSTM) merupakan salah satu jenis arsitektur dalam *deep learning* yang digunakan untuk memproses dan menganalisis data berurutan, seperti teks. LSTM mampu menangani permasalahan kompleks dengan efektif. Dalam konteks klasifikasi teks, LSTM memerlukan proses pelatihan yang lebih intensif, baik dari segi waktu maupun sumber daya komputasi, tetapi menghasilkan tingkat akurasi yang lebih tinggi [4].

Penggunaan LSTM dalam klasifikasi teks ini didukung oleh beberapa penelitian sebelumnya yang berjudul "Penerapan Metode *Long Short Term Memory* untuk Klasifikasi pada *Hate Speech*," hasilnya menunjukkan tingkat akurasi yang cukup tinggi, yaitu 86,23% pada *training data* dan 87,10% *validation data*[5]. Penelitian selanjutnya yang berjudul "Implementasi Model *Long-Short Term Memory* (LSTM) pada Klasifikasi Teks Data SMS Spam Berbahasa Indonesia," hasilnya menunjukkan bahwa kinerja metode LSTM pada kasus klasifikasi teks jauh lebih baik daripada metode Naive Bayes dan KNN dengan masing-masing accuracy yang dihasilkan adalah LSTM 94%, Naive Bayes 67%, dan KNN 70% [6]. Dengan latar belakang tersebut, Penelitian ini bertujuan untuk mengimplementasikan sistem klasifikasi dengan judul "Klasifikasi Tingkat Urgensi Pengaduan Masyarakat Menggunakan metode *Long Short-Term Memory* (LSTM) (Studi Kasus: Website Lapak Aduan Banyumas)". Metode ini telah terbukti memiliki akurasi yang tinggi dalam penelitian sebelumnya. Model klasifikasi tersebut kemudian diimplementasikan ke dalam sebuah website yang dirancang mirip dengan website Lapak Aduan Banyumas.

II. KAJIAN TEORI

Pada penelitian klasifikasi tingkat urgensi pengaduan masyarakat menggunakan *Long Short-Term Memory* (LSTM) digunakan beberapa teori yang diperlukan untuk mendukung kegiatan yang dilakukan. Beberapa landasan teori yang dikemukakan mencakup konsep dasar dan definisi yang berkaitan dengan faktor-faktor pendukung dalam melaksanakan penelitian ini.

A. Pengaduan Masyarakat

Pengaduan masyarakat adalah bentuk pengawasan berupa saran, gagasan, atau keluhan konstruktif kepada aparat pemerintah. Pengaduan ini menjadi elemen penting dalam pelayanan publik untuk memperbaiki kekurangan yang ada [7].

B. Text Mining

Text mining bertujuan untuk mengekstrak informasi yang bermanfaat dari dokumen teks yang tidak terstruktur. Salah satu tugas utamanya adalah *text classification*, yaitu proses mengelompokkan dokumen ke dalam kelas tertentu berdasarkan label yang sudah ditentukan sebelumnya (*supervised learning*) [8]. Sebelum dilakukan analisis lebih lanjut, data teks harus melalui tahap *text pre-processing* agar menjadi lebih terstruktur dan siap digunakan. Tahapan *pre-processing* yang umum dilakukan di antaranya:

a. Case Folding

Case folding merupakan proses penyeragaman bentuk huruf atau mengubah semua karakter huruf besar (*uppercase*) menjadi karakter huruf kecil (*lowercase*) pada dokumen teks ataupun sebaliknya [9].

b. Filtering

Filtering merupakan proses membuang atau menghapus kata-kata serta tanda-tanda yang tidak bermakna secara signifikan, seperti URL, emoticon, dan tanda baca [9].

c. Tokenization

Tokenization merupakan proses pemotongan teks berdasarkan tiap kata menjadi potongan-potongan kecil yang disebut token [9].

d. Stopwords Removal

Stopwords Removal merupakan pemilihan sebagian kata dalam suatu dokumen yang dianggap tidak penting atau tidak menggambarkan isi dari sebuah kalimat [9].

e. Stemming

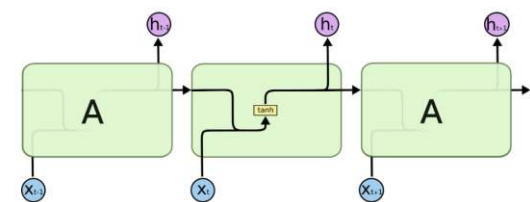
Stemming merupakan proses mentransformasikan kata-kata imbuhan dalam dokumen menjadi kata akarnya (*root word*) atau kata dasar [9].

C. Feature Extraction

Term Frequency-Inverse Document Frequency (TF-IDF) adalah metode pembobotan kata yang mengalikan frekuensi kemunculan kata dalam dokumen *Term Frequency* (TF) dengan kebalikannya di seluruh dokumen *Inverse Document Frequency* (IDF). Hasil bobot ini digunakan untuk proses analisis teks klasifikasi [18].

D. Long Short-Term Memory (LSTM)

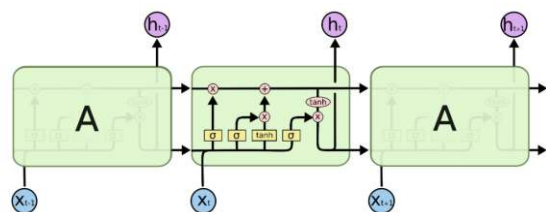
Long Short-Term Memory (LSTM) merupakan variasi dari pengembangan *Recurrent Neural Network* (RNN) yang dibuat untuk menghindari masalah ketergantungan jangka panjang (*long-term dependency*). Pada RNN sering mengalami masalah *vanishing gradient* yaitu gradien yang semakin mengecil hingga *layer* terakhir membuat nilai bobot tidak berubah sehingga menyebabkan proses hasil yang tidak optimal (konvergen) dan *exploding gradient* yaitu gradien semakin menjadi sangat besar, menyebabkan pembaruan bobot menjadi sangat besar sehingga model menjadi tidak stabil (divergen). Untuk mengatasi masalah ini, arsitektur LSTM dikembangkan. LSTM memiliki struktur yang lebih kompleks dengan adanya komponen-komponen khusus seperti *forget gate*, *input gate*, *cell state*, dan *output gate* [10].



GAMBAR 1

LAYER PADA RECURRENT NEURAL NETWORK [3]

Long Short-Term Memory (LSTM) juga mempunyai model seperti pada RNN, tetapi memiliki struktur yang lebih kompleks seperti pada Gambar 2 di bawah ini.



GAMBAR 2

LAYER PADA LONG SHORT-TERM MEMORY [3]

E. Confusion Matrix

Confusion matrix adalah tabel yang menyatakan klasifikasi jumlah data uji yang benar dan jumlah data uji yang salah. Contoh confusion matrix untuk klasifikasi biner [11] ditunjukkan pada Tabel 1 di bawah ini.

TABEL 1
CONFUSION MATRIX

		Kelas Prediksi	
		1	0
Kelas Sebenarnya	1	TP	FN
	0	FP	TN

Keterangan:

TP (*True Positive*): Jumlah dokumen dari kelas 1 yang benar diklasifikasikan sebagai kelas 1

TN (*True Negative*): Jumlah dokumen dari kelas 0 yang benar diklasifikasikan sebagai kelas 0

FP (*False Positive*) : Jumlah dokumen dari kelas 0 yang salah diklasifikasikan sebagai kelas 1

FN (*False Negative*) : jumlah dokumen dari kelas 1 yang salah diklasifikasikan sebagai kelas 0

F. Website

Website merupakan kumpulan halaman dalam satu domain yang saling terhubung dan diakses melalui URL. Website dapat menampilkan teks, gambar, maupun kombinasi konten statis dan dinamis. Hubungan antarmuka disebut *hyperlink*, sedangkan teks penghubungnya disebut *hypertext* [12]. Flask adalah *framework* yang menggunakan pemrograman Python, Flask bertindak sebagai *framework* aplikasi dan tampilan web. Pengembang dapat menggunakan Flask dan bahasa Python untuk membuat web terstruktur dan mengelola perilaku web dengan lebih mudah [13]. Python sendiri merupakan bahasa pemrograman interpretatif yang mudah dipelajari, mendukung paradigma objek, imperatif, dan fungsional, serta dilengkapi koleksi library yang banyak dan struktur kode yang jelas [14].

III. METODE

Penelitian ini bersifat kuantitatif, karena memanfaatkan data dari hasil pengumpulan data pada website Lapak Aduan Banyumas untuk proses klasifikasi menggunakan algoritma *Long Short-Term Memory* (LSTM). Tahapan penelitian meliputi:

1. Tahapan Identifikasi

Tahap identifikasi diawali dengan mengidentifikasi dan merumuskan masalah yang terjadi di website Lapak Aduan Banyumas. Dalam tahap ini, dilakukan juga studi literatur untuk memahami teori, penelitian terdahulu, dan konsep yang relevan untuk membangun pemahaman awal yang kuat sebagai dasar penelitian.

2. Tahapan Pengumpulan Data

Dalam penelitian ini, dilakukan pengambilan data pengaduan masyarakat dengan metode *web scraping* dari website Lapak Aduan Banyumas. Alamat website dapat diakses melalui URL: lapakaduan.banyumaskab.go.id. Data ini diperoleh pada tanggal 2 Januari 2024 sampai 27 April 2025 yang berjumlah 19.240 data. Proses pengumpulan data

dilakukan dengan cara *web scraping* pada halaman website Lapak Aduan Banyumas

3. Tahapan *Preprocessing Data*

Tahap *preprocessing data* mencakup pengumpulan data melalui web scraping dari website Lapak Aduan Banyumas, dilanjutkan dengan pembersihan data teks menggunakan beberapa tahapan seperti *case folding*, *filtering*, *tokenization*, *stopwords removal*, dan *stemming* agar data siap digunakan untuk pemodelan lebih lanjut.

4. Tahapan Pelabelan Data

Pada tahap pelabelan data, dilakukan dengan menggunakan *semantic labeling* pada dataset pengaduan berdasarkan kategori urgensi, dengan menentukan kata kunci yang mewakili kondisi “*Urgent*” dan “*Not Urgent*”

5. Tahapan *Feature Extraction*

Pada tahap ini dilakukan pembobotan kata menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk memberikan bobot penting pada kata-kata, sehingga sistem dapat mengenali makna dan relevansi masing-masing laporan.

6. Tahapan Pembuatan Model Klasifikasi

Tahap modeling mencakup proses *splitting dataset* menjadi data latih dan data uji, diikuti dengan pelatihan model LSTM menggunakan parameter yang telah ditentukan, pengujian model untuk mengukur performa, evaluasi menggunakan metrik akurasi, *precision*, *recall*, dan *f1-score*, hingga tahap *deployment* model ke dalam website Lapak Aduan Banyumas.

7. Tahapan *Deployment*

Peneliti mengembangkan antarmuka website yang menyerupai Lapak Aduan Banyumas, mencakup halaman-halaman utama seperti *register*, *login* masyarakat, *login* admin, beranda, alur pengaduan, daftar OPD, input pengaduan, daftar pengaduan, grafik pengaduan, statistik pengaduan, dan *dashboard* admin.

IV. HASIL DAN PEMBAHASAN

A. Pengumpulan Data

Dalam penelitian ini, dilakukan pengambilan data pengaduan masyarakat dengan metode *web scraping* dari website Lapak Aduan Banyumas. Alamat website dapat diakses melalui URL: lapakaduan.banyumaskab.go.id. Data ini diperoleh pada tanggal 2 Januari 2024 sampai 27 April 2025 yang berjumlah 19.240 data. Proses pengumpulan data dilakukan dengan cara *web scraping* pada halaman website Lapak Aduan Banyumas menggunakan beberapa *library* berbasis *Python* yaitu *DrissionPage*, *BeautifulSoup*, *requests*, *time* dan *csv*. Data yang diambil meliputi title (judul pengaduan), description (keterangan sumber dan tanggal), status (status penanganan), kategori (jenis permasalahan), kepada (instansi tujuan), content (isi aduan), dan page (halaman). Berikut adalah hasil pengumpulan data pengaduan.

	Tesis	Skripsi/tesis	Status	Isi	Aspek	Page
1	REVISI					
2	REVISI	REVISI				
3	REVISI	REVISI				
4	REVISI	REVISI				
5	REVISI	REVISI				
6	REVISI	REVISI				
7	REVISI	REVISI				
8	REVISI	REVISI				
9	REVISI	REVISI				
10	REVISI	REVISI				
11	REVISI	REVISI				
12	REVISI	REVISI				
13	REVISI	REVISI				
14	REVISI	REVISI				
15	REVISI	REVISI				
16	REVISI	REVISI				
17	REVISI	REVISI				
18	REVISI	REVISI				
19	REVISI	REVISI				
20	REVISI	REVISI				
21	REVISI	REVISI				
22	REVISI	REVISI				
23	REVISI	REVISI				
24	REVISI	REVISI				
25	REVISI	REVISI				
26	REVISI	REVISI				
27	REVISI	REVISI				
28	REVISI	REVISI				
29	REVISI	REVISI				
30	REVISI	REVISI				
31	REVISI	REVISI				
32	REVISI	REVISI				
33	REVISI	REVISI				
34	REVISI	REVISI				
35	REVISI	REVISI				
36	REVISI	REVISI				
37	REVISI	REVISI				
38	REVISI	REVISI				
39	REVISI	REVISI				
40	REVISI	REVISI				
41	REVISI	REVISI				
42	REVISI	REVISI				
43	REVISI	REVISI				
44	REVISI	REVISI				
45	REVISI	REVISI				
46	REVISI	REVISI				
47	REVISI	REVISI				
48	REVISI	REVISI				
49	REVISI	REVISI				
50	REVISI	REVISI				
51	REVISI	REVISI				
52	REVISI	REVISI				
53	REVISI	REVISI				
54	REVISI	REVISI				
55	REVISI	REVISI				
56	REVISI	REVISI				
57	REVISI	REVISI				
58	REVISI	REVISI				
59	REVISI	REVISI				
60	REVISI	REVISI				
61	REVISI	REVISI				
62	REVISI	REVISI				
63	REVISI	REVISI				
64	REVISI	REVISI				
65	REVISI	REVISI				
66	REVISI	REVISI				
67	REVISI	REVISI				
68	REVISI	REVISI				
69	REVISI	REVISI				
70	REVISI	REVISI				
71	REVISI	REVISI				
72	REVISI	REVISI				
73	REVISI	REVISI				
74	REVISI	REVISI				
75	REVISI	REVISI				
76	REVISI	REVISI				
77	REVISI	REVISI				
78	REVISI	REVISI				
79	REVISI	REVISI				
80	REVISI	REVISI				
81	REVISI	REVISI				
82	REVISI	REVISI				
83	REVISI	REVISI				
84	REVISI	REVISI				
85	REVISI	REVISI				
86	REVISI	REVISI				
87	REVISI	REVISI				
88	REVISI	REVISI				
89	REVISI	REVISI				
90	REVISI	REVISI				
91	REVISI	REVISI				
92	REVISI	REVISI				
93	REVISI	REVISI				
94	REVISI	REVISI				
95	REVISI	REVISI				
96	REVISI	REVISI				
97	REVISI	REVISI				
98	REVISI	REVISI				
99	REVISI	REVISI				
100	REVISI	REVISI				

GAMBAR 3
PENGUMPULAN DATA PENGADUAN

B. Preprocessing Data

1.. Case Folding

Tahapan *case folding* dilakukan untuk penyeragaman bentuk huruf besar (*uppercase*) menjadi huruf kecil (*lowercase*) pada teks pengaduan. Berikut adalah tabel hasil *case folding*.

TABEL 2
CASE FOLDING

Teks Input	Teks Output
ngarep unwiku tolong di aspal ulang atau di tambal, banyak lubang dan tiap hari ada yg jatuh pengendara motor 🙄	ngarep unwiku tolong di aspal ulang atau di tambal, banyak lubang dan tiap hari ada yg jatuh pengendara motor 🙄

b. Filtering

Tahapan *filtering* dilakukan untuk menghapus kata-kata dan tanda-tanda yang tidak bermakna secara signifikan, seperti URL, angka, simbol, dan tanda baca. Proses ini dilakukan menggunakan “import re (regular expression)” untuk menghapus URL yang diawali dengan http, https, atau www. Berikut adalah tabel hasil *filtering*.

TABEL 3
FILTERING

Teks Input	Teks Output
ngarep unwiku tolong di aspal ulang atau di tambal, banyak lubang dan tiap hari ada yg jatuh pengendara motor 🙄	ngarep unwiku tolong di aspal ulang atau di tambal banyak lubang dan tiap hari ada yg jatuh pengendara motor

c. Tokenization

Tahapan *tokenization* dilakukan untuk memotong teks berdasarkan tiap kata menjadi potongan-potongan kecil yang disebut token. Berikut adalah tabel hasil *tokenization*.

TABEL 4
TOKENIZATION

Teks Input	Teks Output
ngarep unwiku tolong di aspal ulang atau di tambal banyak lubang dan tiap hari ada yg jatuh pengendara motor	['ngarep', 'unwiku', 'tolong', 'di', 'aspal', 'ulang', 'atau', 'di', 'tambal', 'banyak', 'lubang', 'dan', 'tiap', 'hari', 'ada', 'yg', 'jatuh', 'pengendara', 'motor']

d. Stopwords Removal

Tahapan *stopwords removal* dilakukan untuk menghapus kata-kata umum yang dianggap tidak memiliki kontribusi signifikan terhadap makna utama teks, seperti kata sambung, kata depan, atau kata bantu. Proses ini dilakukan dengan mencocokkan setiap token dengan daftar *stopwords* bahasa Indonesia yang disediakan oleh *library* NLTK. Kata-kata yang cocok dengan daftar tersebut akan dihapus dari token, sehingga hanya menyisakan kata-kata penting yang lebih relevan untuk analisis lebih lanjut. Berikut adalah tabel hasil *stopwords removal*.

TABEL 5
STOPWORDS REMOVAL

Teks Input	Teks Output
['ngarep', 'unwiku', 'tolong', 'di', 'aspal', 'ulang', 'atau', 'di', 'tambal', 'banyak', 'lubang', 'dan', 'tiap', 'hari', 'ada', 'yg', 'jatuh', 'pengendara', 'motor']	['ngarep', 'unwiku', 'tolong', 'aspal', 'ulang', 'tambal', 'lubang', 'yg', 'jatuh', 'pengendara', 'motor']

e. Stemming

Tahapan *stemming* dilakukan untuk mengubah kata-kata berimbuhan menjadi bentuk kata dasar menggunakan *library* Sastrawi berbasis bahasa Indonesia. Tujuannya adalah untuk menyederhanakan variasi kata sehingga analisis teks dapat dilakukan secara konsisten terhadap bentuk dasar kata. Berikut adalah tabel hasil *stemming*.

TABEL 6
STEMMING

Teks Input	Teks Output
['ngarep', 'unwiku', 'tolong', 'aspal', 'ulang', 'tambal', 'lubang', 'yg', 'jatuh', 'pengendara', 'motor']	['ngarep', 'unwiku', 'tolong', 'aspal', 'ulang', 'tambal', 'lubang', 'yg', 'jatuh', 'kendara', 'motor']

C. Pelabelan Data menggunakan Semantic Labelling

Pada tahap pelabelan data, digunakan pendekatan *semantic labeling* untuk mengelompokkan laporan pengaduan masyarakat ke dalam dua kategori “Urgent” dan “Not Urgent”. Pelabelan dilakukan dengan mencocokkan kata atau frasa pada laporan terhadap daftar kata kunci yang telah ditentukan. Kategori “Urgent” memuat kata kunci yang mencerminkan kondisi darurat atau butuh penanganan segera, seperti “kebakaran”, “longsor”, atau “darurat”. Sebaliknya, kategori “Not Urgent” berisi kata kunci administratif atau keluhan ringan, seperti “ktp”, “informasi”, dan “usulan”. Jika laporan mengandung kata dari daftar “Urgent”, maka diberi label “Urgent”, sedangkan jika hanya memuat kata dari daftar “Not Urgent” atau tidak mengandung keduanya, diberi label “Not Urgent”. Hasilnya, dari total 19.240 laporan, sebanyak 12.904 masuk kategori “Urgent” dan 6.336 termasuk “Not Urgent”.

D. Term Frequency-Inverse Document Frequency (TF-IDF)

Pada tahap *feature extraction*, digunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk mengidentifikasi dan menentukan kata-kata kunci yang paling relevan dalam membedakan kategori *Urgent* dan *Not Urgent*. Hasil perhitungan menunjukkan kata-kata dengan skor TF-IDF tertinggi, seperti *jalan*, *min*, *mohon*, *desa*, *yg*, *banyumas*, dan *rusak*. Kata-kata tersebut memiliki bobot penting yang tinggi karena sering muncul dalam laporan pengaduan masyarakat. Berdasarkan hasil ini, kata-kata tersebut digunakan sebagai acuan dalam menyusun daftar kata kunci untuk mendukung proses klasifikasi tingkat urgensi pengaduan.

E. Pembuatan Model menggunakan Metode *Long Short-Term Memory* (LSTM)

Pada tahap pembuatan model klasifikasi, terdapat beberapa konfigurasi model *Long Short-Term Memory* (LSTM) yang akan digunakan, untuk lebih jelasnya disajikan pada Tabel 7 di bawah ini

TABEL 7
KONFIGURASI MODEL LSTM

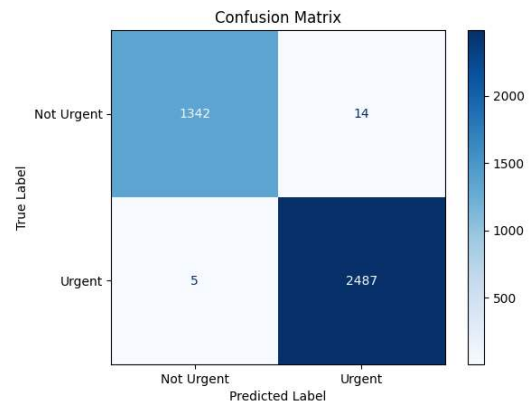
No	Komponen	Spesifikasi
1	Jumlah <i>neuron</i> pada <i>input layer</i>	100
2	Fungsi <i>input layer</i>	<i>Embedding layer</i>
3	Jumlah <i>hidden layer</i>	4 layer: [256,128,64,32] (Bidirectional LSTM)
4	Fungsi <i>output gate</i>	<i>Sigmoid activation</i>
5	<i>Batch size</i>	64
6	Dataset pemodelan	19.240 data. 80% data <i>training</i> dan 20% data <i>testing</i>
7	Optimizer	Adam, Nadam, RMSprop
8	<i>Loss Function</i>	Binary crossentropy
9	<i>Early Stopping</i>	<i>Patience: 7</i>
10	<i>Tracking</i>	val loss
11	<i>Apply Best Weights</i>	True
12	Jumlah <i>epoch</i>	50

Konfigurasi model LSTM berikut mencakup sejumlah komponen penting yang digunakan dalam proses klasifikasi urgensi aduan. Model ini menerima *input* berupa urutan teks sepanjang 100 token, yang direpresentasikan menggunakan *embedding layer* dengan dimensi vektor sebesar 50. Struktur jaringan terdiri dari empat *hidden layer* bertipe *Bidirectional LSTM* dengan jumlah unit berturut-turut sebesar 256, 128, 64, dan 32. Untuk klasifikasi, *output gate* menggunakan fungsi aktivasi *sigmoid*.

Pelatihan dilakukan dengan *batch size* sebesar 64, menggunakan *optimizer Adam, Nadam, RMSprop* (*Nadam* memberikan hasil terbaik) dan fungsi *loss* berupa *binary crossentropy*. Dataset terdiri dari 19.240 data, yang dibagi menjadi 80% data pelatihan dan 20% data pengujian. Untuk mencegah *overfitting*, diterapkan mekanisme *early stopping* dengan *patience* selama 7 *epoch*, serta pelacakan terhadap metrik *val loss*. Proses pelatihan berlangsung selama maksimal 50 *epoch* dengan penerapan bobot terbaik secara otomatis. Konfigurasi ini dirancang untuk memaksimalkan akurasi klasifikasi melalui pemanfaatan arsitektur LSTM yang dalam dan reguler.

F. Evaluasi Kinerja Model Menggunakan *Confusion Matrix*

Untuk mengevaluasi performa model dalam mengklasifikasikan laporan pengaduan, digunakan metode *confusion matrix* yang mampu menunjukkan jumlah prediksi benar dan salah berdasarkan kelas sebenarnya dan kelas yang diprediksi. Visualisasi *confusion matrix* membantu dalam memahami sejauh mana model berhasil mengenali laporan *urgent* maupun *not urgent* secara akurat.



GAMBAR 4
EVALUASI MENGGUNAKAN *CONFUSION MATRIX*

Gambar 4 *confusion matrix* menunjukkan performa model dalam mengklasifikasikan laporan menjadi dua kategori: *Not Urgent* (label 0) dan *Urgent* (label 1). Dari total 3848 data *testing*, hasil klasifikasi sebagai berikut:

- True Positive* (TP): 2.487 laporan *urgent* berhasil diklasifikasikan dengan benar sebagai *urgent*.
- True Negative* (TN): 1.342 laporan *not urgent* berhasil diklasifikasikan dengan benar sebagai *not urgent*.
- False Positive* (FP): 14 laporan tidak *urgent* salah diklasifikasikan sebagai *urgent*.
- False Negative* (FN): 5 laporan *urgent* salah diklasifikasikan sebagai *not urgent*.

Dari laporan klasifikasi (*classification report*), diperoleh metrik evaluasi melalui Tabel 8 di bawah ini:

TABEL 8
HASIL EVALUASI KLASIFIKASI

Label	Precision	Recall	F1-score	Support
Not Urgent (0)	0.99629	0.98968	0.99297	1356
Urgent (1)	0.99440	0.99799	0.99489	2492
Accuracy	-	-	0.99510	3848
Macro Avg	0.99535	0.99383	0.99393	3848
weighted Avg	0.99490	0.99494	0.99484	3848

Model berhasil mengklasifikasikan pengaduan ke dalam kategori *Urgent* dan *Not Urgent* dengan akurasi 99,51%. Nilai *precision*, *recall*, dan *F1-score* rata-rata 0,99 di kedua kategori menunjukkan model mampu membedakan dua kelas secara konsisten dan akurat. Nilai *macro average* dan *weighted average* juga mencapai 0,99, memperkuat konsistensi performa model.

G. Tampilan Antarmuka Website LAB

Peneliti telah mengembangkan antarmuka website yang menyerupai Lapak Aduan Banyumas, mencakup halaman-halaman utama seperti *register*, *login* masyarakat, *login* admin, beranda, alur pengaduan, daftar OPD, input pengaduan, daftar pengaduan, grafik pengaduan, statistik pengaduan, dan *dashboard* admin. Proses pengembangan website dilakukan menggunakan *software* Visual Studio Code, pada bagian *frontend* menggunakan bahasa

pemrograman HTML, CSS, dan JavaScript, sedangkan *backend* menggunakan bahasa pemrograman Python dengan *framework* Flask. Berikut adalah tampilan website Lapak Aduan Banyumas.



GAMBAR 5
TAMPILAN UTAMA WEBSITE LAB

H. Hasil Pengujian

1. Hasil Pengujian Sistem Menggunakan *Black Box Testing*

Pengujian sistem merupakan tahap untuk menilai kualitas aplikasi yang telah dirancang dan dikembangkan. Pada tahap ini digunakan metode *Black Box Testing* untuk menguji setiap fungsi dan alur proses tanpa memperhatikan struktur kode program, melainkan hanya menilai keluaran sistem berdasarkan input yang diberikan. Pengujian ini dilakukan secara langsung oleh Bapak Doni Prasetyo selaku Koordinator Sekretariat Lapak Aduan Banyumas. Hasil pengujian sistem menunjukkan bahwa aplikasi yang telah dirancang dan dikembangkan mampu menjalankan setiap fungsi sesuai dengan spesifikasi yang ditetapkan. Pengujian dilakukan terhadap dua jenis peran utama dalam sistem, yaitu Pengguna (Masyarakat) dan Admin.

2. Hasil Pengujian Tingkat Urgensi Pengaduan Masyarakat

Pengujian ini bertujuan untuk mengevaluasi ketepatan sistem dalam mengklasifikasikan tingkat urgensi dari berbagai jenis pengaduan masyarakat melalui 60 skenario uji yang bervariasi. Skenario yang digunakan mencakup kategori pengaduan *Urgent*, *Not Urgent*, *Formal Urgent*, *Unformal Urgent*, *Formal Not Urgent*, *Unformal Not Urgent*, serta masukan tidak valid seperti huruf acak, angka acak, simbol acak, dan kombinasi acak. Selain itu, pengujian juga melibatkan pengaduan dalam bahasa Jawa, baik yang bersifat *urgent* maupun *not urgent*. Untuk memverifikasi kebenaran klasifikasi tingkat urgensi yang dihasilkan oleh sistem, validasi dilakukan secara langsung oleh Bapak Doni Prasetyo selaku Koordinator Sekretariat Lapak Aduan Banyumas.

I. Analisis Hasil Pengujian

Pada tahap ini, pembahasan difokuskan pada analisis hasil pengujian yang telah dilakukan menggunakan metode *Black Box Testing* melalui *browser* Google Chrome, serta analisis terhadap sistem klasifikasi tingkat urgensi pengaduan masyarakat pada website Lapak Aduan Banyumas.

1. Analisis Hasil Pengujian *Black Box Testing*

Hasil pengujian *Black Box Testing* menunjukkan total 44 pola situasi. Hasil pengujian pada halaman pelanggan menunjukkan 33 skenario berhasil dan 0 tidak berhasil. Sementara itu, pada halaman admin terdapat 11 skenario berhasil dan 0 tidak berhasil. Rekapitulasi hasil pengujian

Blackbox Testing pada tabel 9 menunjukkan total 44 skenario yang berhasil dan 0 skenario yang tidak berhasil.

TABEL 9
ANALISIS HASIL PENGUJIAN *BLACK BOX TESTING*

No	Pola Situasi	Berhasil	Tidak Berhasil
1	Halaman Pengguna	33	0
2	Halaman Admin	11	0
Total Hasil Pengujian		44	0

Dari tabel rekapitulasi hasil pengujian di atas, dilakukan perhitungan menggunakan teknik analisis deskriptif untuk memperoleh nilai persentase kelayakan sistem. Rumus yang digunakan adalah sebagai berikut:

$$\text{Persentase Kelayakan} = \frac{\text{Skor Observasi}}{\text{Skor yang diharapkan}} \times 100\%$$

Dengan demikian, persentasenya adalah sebagai berikut:

$$\text{Persentase Kelayakan} = \frac{44}{44} \times 100\% = 100\%$$

Berdasarkan perhitungan *Black Box Testing* di atas, diperoleh total persentase kelayakan sebesar 100%. Dengan demikian, dapat disimpulkan bahwa sistem website Lapak Aduan Banyumas yang diuji sangat layak dan dapat digunakan dengan baik tanpa ditemukan kesalahan fungsional.

2. Analisis Hasil Pengujian Urgensi Pengaduan Masyarakat

Hasil pengujian menunjukkan bahwa telah dilakukan pengujian terhadap total 60 skenario untuk mengevaluasi kemampuan sistem dalam mengklasifikasikan tingkat urgensi pengaduan masyarakat. Skenario-skenario tersebut mencakup berbagai variasi bentuk dan gaya penulisan pengaduan, mulai dari yang bersifat formal hingga informal, serta pengaduan acak yang tidak mengandung makna urgensi. Tabel 10 di bawah ini menunjukkan hasil pengujian pada 60 skenario uji coba utama, masing-masing dengan 5 percobaan.

TABEL 10
ANALISIS HASIL PENGUJIAN URGENSI
PENGADUAN

No	Skenario Uji Coba	Berhasil	Tidak Berhasil
1	<i>Urgent</i>	5	0
2	<i>Not Urgent</i>	5	0
3	<i>Formal Urgent</i>	5	0
4	<i>Unformal Urgent</i>	5	0
5	<i>Formal Not Urgent</i>	5	0
6	<i>Unformal Not Urgent</i>	5	0
7	Acak – huruf (<i>Not Urgent</i>)	5	0
8	Acak – angka (<i>Not Urgent</i>)	5	0
9	Acak – simbol (<i>Not Urgent</i>)	5	0
10	Acak – kombinasi (<i>Not Urgent</i>)	5	0
11	Bahasa Jawa <i>Urgent</i>	5	0
12	Bahasa Jawa <i>Not Urgent</i>	5	0
Total Hasil Pengujian		60	0

Dari tabel rekapitulasi hasil pengujian di atas, dilakukan perhitungan menggunakan teknik analisis deskriptif untuk memperoleh nilai persentase kelayakan sistem. Rumus yang digunakan adalah sebagai berikut:

$$\text{Persentase Kelayakan} = \frac{\text{Skor Observasi}}{\text{Skor yang diharapkan}} \times 100\%$$

Berdasarkan hasil pengujian, diperoleh:

$$\text{Persentase Kelayakan} = \frac{60}{60} \times 100\% = 100\%$$

Dengan demikian, total persentase kelayakan sistem adalah sebesar 100%. Artinya, seluruh skenario uji coba yang dilakukan berhasil diklasifikasikan dengan tepat oleh sistem, tanpa adanya kegagalan dalam penentuan tingkat urgensi pengaduan. Berdasarkan hasil ini, dapat disimpulkan bahwa sistem klasifikasi tingkat urgensi pengaduan pada website *Lapak Aduan Banyumas* memiliki performa yang sangat baik dan layak digunakan.

V. KESIMPULAN

Berdasarkan hasil penelitian, sistem klasifikasi tingkat urgensi pengaduan masyarakat dengan metode *Long Short-Term Memory* (LSTM) menunjukkan performa sangat baik. Model berhasil mengklasifikasikan pengaduan ke dalam kategori *Urgent* dan *Not Urgent* dengan akurasi 99,51%. Nilai *precision*, *recall*, dan *F1-score* rata-rata 0,99 di kedua kategori menunjukkan model mampu membedakan dua kelas secara konsisten dan akurat. Nilai *macro average* dan *weighted average* juga mencapai 0,99, memperkuat konsistensi performa model. Selain itu, hasil pengujian sistem di website *Lapak Aduan Banyumas* pada 60 skenario uji menunjukkan keberhasilan 100%, tanpa kesalahan klasifikasi. Validasi dilakukan langsung oleh Bapak Doni Prasetyo selaku Koordinator Sekretariat *Lapak Aduan Banyumas*. Dengan demikian, sistem klasifikasi berbasis LSTM ini dinyatakan layak digunakan karena mampu mengatasi ketidakakuratan dalam menentukan tingkat urgensi pengaduan yang masuk.

REFERENSI

- [1] E. Sahadi. *Urgensi Peningkatan Kualitas Pelayanan Publik Dalam Mewujudkan Tata Kelola Pemerintahan Yang Baik Perspektif Siyasa Idariyah (Studi Kasus Di Desa Sukaraja Kecamatan Kedurang Ilir Kabupaten Bengkulu Selatan)*. Thesis (Diploma), UIN Fatmawati Sukarno, Bengkulu, 2021. [Online]. Available: <http://repository.iainbengkulu.ac.id/id/eprint/6840>.
- [2] I. Amirrosyad. *Efektivitas Aplikasi Lapak Aduan Dalam Pelayanan Pengaduan Masyarakat Di Kabupaten Banyumas (Studi Kasus Pada Dinas Komunikasi Dan Informatika Kabupaten Banyumas)*. Thesis, Institut Pemerintahan Dalam Negeri, 2023. [Online]. Available: <http://eprints.ipdn.ac.id/id/eprint/13531>.
- [3] F. Rozi, V. N. Wijayaningrum, and N. Khozin. "Klasifikasi Teks Laporan Masyarakat Pada Situs *Lapor!* Menggunakan Recurrent Neural Network," *SISTEMASI: Jurnal Sistem Informasi*, vol. 9, no. 3, pp. 633–645, 2020. doi: <https://doi.org/10.32520/stmsi.v9i3.977>.
- [4] S. Ramadhanti, D. P. Rini, and D. Rodiah. "Klasifikasi Emosi Pada Kalimat Di Twitter Menggunakan Algoritma Long Short-Term Memory". Thesis, Universitas Sriwijaya, 2023. [Online]. Available: <http://repository.unsri.ac.id/id/eprint/137933>.
- [5] B. Arief, H. Kholifatullah, and A. Prihanto. "Penerapan Metode Long Short Term Memory Untuk Klasifikasi Pada Hate Speech," *Journal of Informatics and Computer Science*, vol. 4, no. 3, pp. 292–297, 2023. doi: <https://doi.org/10.26740/jinacs.v4n03.p292-297>.
- [6] E. Dwi Pratama. "Implementasi Model Long-Short Term Memory (LSTM) pada Klasifikasi Teks Data SMS Spam Berbahasa Indonesia," *The Journal on Machine Learning and Computational Intelligence (JMLCI)*, vol. 1, no. 2, pp. 38–42, Jul. 2022. doi: <https://doi.org/10.26740/vol1iss2y2022id12>.
- [7] H. Sabeni and E. D. Setiamandani. "Pengelolaan Pengaduan Masyarakat Dalam Upaya Meningkatkan Kualitas Pelayanan Publik," *Jurnal Ilmu Sosial dan Ilmu Politik*, vol. 9, no. 1, pp. 43–52, 2020. doi: <https://doi.org/10.33366/jisip.v9i1.2214>.
- [8] Ananto, M. Ihsan, and W. S. Winanhju. "Klasifikasi Kategori Pengaduan Masyarakat Melalui Kanal LAPOR! Menggunakan Artificial Neural Network," *INFERENSI*, vol. 2, no. 2, pp. 71–79, 2019. doi: <http://dx.doi.org/10.12962/j27213862.v2i2.6821>.
- [9] M. S. Fani, R. Santoso, and Suparti. "Penerapan Text Mining untuk Melakukan Clustering Data Tweet Akun Bliibli Pada Media Sosial Twitter Menggunakan K-Means Clustering," *JURNAL GAUSSIAN*, vol. 10, no. 4, pp. 583–593, 2021. doi: <https://doi.org/10.14710/j.gauss.10.4.583-593>.
- [10] A. Hanifa, S. A. Fauzan, M. Hikail, and M. B. Ashfiya. "Perbandingan Metode LSTM dan GRU (RNN) untuk Klasifikasi Berita Palsu Berbahasa Indonesia," *DINAMIKA REKAYASA*, vol. 17, no. 1, pp. 33–39, 2021. doi: <http://dx.doi.org/10.20884/1.dr.2021.17.1.436>.
- [11] D. Normawati and S. A. Prayogi. "Implementasi Naïve Bayes Classifier Dan Confusion Matrix Pada Analisis Sentimen Berbasis Teks Pada Twitter," *Jurnal Sains Komputer & Informatika (J-SAKTI)*, vol. 5, no. 2, pp. 697–711, 2021. doi: <http://dx.doi.org/10.30645/j-sakti.v5i2.369>.
- [12] D. Febri Kuncoro, U. Juniarti, J. Syahputra, R. Bagus, B. Sumantri, and R. Suryani. "Rancang Bangun Sistem Pengaduan Masyarakat Berbasis Web Dengan Metode Waterfall," *Jurnal Sistem Informasi dan Teknologi Peradaban (JSITP)*, vol. 3, no. 2, pp. 14–19, 2022. doi: <https://doi.org/10.58436/jsitp.v3i2.1259>.
- [13] P. T. L. Dewi, F. Dewata, and M. A. Nugroho. *Implementasi Machine Learning Model Deployment Pada Website Pemantauan Kondisi Sungai Citarum Menggunakan Platform As-A-Service*. Thesis, Telkom University, 2022.

- [14] M. Romzi and B. Kurniawan. “Implementasi Pemrograman Python Menggunakan Visual Studio Code,” JIK, vol. 11, no. 2, pp. 1–9, 2020. [Online]. Available:

<https://journal.unmaha.ac.id/index.php/jik/article/view/198>.

