

IMPLEMENTASI DEEP LEARNING MODEL BIDIRECTIONAL LSTM DALAM ANALISIS SENTIMEN KEPUASAN PUBLIK TERHADAP GUBERNUR JAWA BARAT DEDI MULYADI DI INSTAGRAM

Adam Adnan¹, Bambang Irawan², Ahmad Faqih³, Ade Rizki Rinaldi⁴

Program Studi Teknik Informatika¹²³
Program Studi Rekayasa Perangkat Lunak³

STMIK IKMI Cirebon
<https://ikmi.ac.id/page/18/?lang=de>
adambong340@gmail.com

(*) Corresponding Author : adambong340@gmail.com
Published : 20 Juli 2026

Abstract—Social media has become a public interaction platform that allows people to express opinions, support, criticism, and evaluations of public figures and government policies, including through Instagram. Dedi Mulyadi, as a public figure, receives numerous public responses through the comment section of his official Instagram account. However, these comment data are still unstructured and have not been automatically analyzed to determine public sentiment tendencies objectively. In addition, sentiment data on social media are generally imbalanced datasets, where the number of comments in certain classes is significantly more dominant than others. This condition may cause classification models to focus more on learning patterns from the majority class and become less optimal in recognizing minority classes. This study aims to analyze Instagram comment sentiment using the Bidirectional Long Short-Term Memory (BiLSTM) method into three categories: positive, negative, and neutral, as well as to analyze the effect of applying oversampling on model performance. Data were collected from 20 posts on the Instagram account @dedimulyadi71 through a web scraping process. The research stages included data filtering, preprocessing, manual labeling, tokenization, padding, data transformation, and model evaluation using K-Fold Cross Validation ($k=5$). After the data selection and cleaning process, a total of 2,617 comments were obtained and used as the research dataset. The results showed that the sentiment distribution consisted of 1,763 positive comments, 672 neutral comments, and 182 negative comments, indicating that the dataset was imbalanced. To address this condition, the study used two experimental scenarios, namely without oversampling and with text augmentation-based oversampling. In the scenario without oversampling, the BiLSTM model achieved an average accuracy of 72.64%. Meanwhile, in the oversampling scenario, the model achieved an average accuracy of 74.24%. These results indicate that the application of oversampling was able to improve the model performance in sentiment classification compared to using the original data without class balancing. Overall, the BiLSTM model was able to identify sentiment patterns in public comments despite the informal, short, and varied nature of the data. This study is expected to serve as a reference for the development of Indonesian-language sentiment analysis on social media and support the application of artificial intelligence in data-driven public opinion analysis.

Keywords: Sentiment Analysis, BiLSTM, Deep Learning, Instagram, Oversampling, Public Comments.

Abstrak—Media sosial telah menjadi ruang interaksi publik yang memungkinkan masyarakat menyampaikan opini, dukungan, kritik, dan evaluasi terhadap figur publik maupun kebijakan pemerintah, termasuk melalui platform *Instagram*. Dedi Mulyadi sebagai tokoh publik memperoleh banyak respons masyarakat melalui kolom komentar pada akun *Instagram* resminya. Namun, data komentar tersebut masih bersifat tidak terstruktur dan belum dianalisis secara otomatis untuk mengetahui kecenderungan sentimen publik secara objektif. Selain itu, distribusi data sentimen pada media sosial umumnya tidak seimbang (*imbalanced dataset*), di mana jumlah komentar pada kelas tertentu jauh lebih dominan dibandingkan kelas lainnya. Kondisi ini berpotensi menyebabkan model klasifikasi lebih cenderung mempelajari pola dari kelas mayoritas dan kurang optimal dalam mengenali kelas minoritas. Penelitian ini bertujuan untuk menganalisis sentimen komentar *Instagram* menggunakan metode *Bidirectional Long Short-Term Memory (BiLSTM)* ke dalam tiga kategori, yaitu positif, negatif, dan netral, serta menganalisis pengaruh penerapan *oversampling* terhadap performa model. Data dikumpulkan dari 20 unggahan akun *Instagram* @dedimulyadi71 melalui proses *web scraping*. Tahapan penelitian meliputi penyaringan data,

preprocessing, pelabelan manual, *tokenisasi*, *padding*, transformasi data, serta evaluasi model menggunakan *K-Fold Cross Validation* ($k=5$). Setelah proses seleksi dan pembersihan data, diperoleh sebanyak 2.617 komentar yang digunakan sebagai dataset penelitian. Hasil penelitian menunjukkan bahwa distribusi sentimen terdiri atas 1.763 komentar positif, 672 komentar netral, dan 182 komentar negatif, sehingga dataset penelitian termasuk dalam kategori tidak seimbang. Untuk menangani kondisi tersebut, penelitian menggunakan dua skenario eksperimen, yaitu tanpa *oversampling* dan dengan *oversampling* berbasis *augmentasi teks*. Pada skenario tanpa *oversampling*, model *BiLSTM* memperoleh rata-rata akurasi sebesar 72,64%. Sementara itu, pada skenario dengan *oversampling* diperoleh rata-rata akurasi sebesar 74,24%. Hasil tersebut menunjukkan bahwa penerapan *oversampling* mampu meningkatkan performa model dalam klasifikasi sentimen dibandingkan penggunaan data asli tanpa penyeimbangan kelas. Secara keseluruhan, model *BiLSTM* mampu mengidentifikasi pola sentimen komentar publik meskipun data bersifat informal, singkat, dan bervariasi. Penelitian ini diharapkan dapat menjadi referensi dalam pengembangan analisis sentimen berbahasa Indonesia pada media sosial serta mendukung pemanfaatan kecerdasan buatan dalam analisis opini publik berbasis data.

Kata Kunci: Analisis Sentimen, *BiLSTM*, *Deep Learning*, *Instagram*, *Oversampling*, Komentar Publik.

INTRODUCTION

Perkembangan teknologi informasi telah mendorong perubahan signifikan dalam pola komunikasi masyarakat. Kehadiran media sosial menjadikan proses penyampaian informasi, opini, kritik, maupun dukungan berlangsung lebih cepat dan terbuka. Saat ini, media sosial tidak hanya berfungsi sebagai sarana interaksi personal, tetapi juga telah berkembang menjadi ruang publik digital tempat masyarakat mengekspresikan pandangan terhadap isu sosial, kebijakan pemerintah, maupun figur publik. Salah satu platform media sosial yang memiliki tingkat interaksi tinggi adalah *Instagram*. Platform ini memungkinkan pengguna berkomunikasi melalui unggahan foto, video, serta kolom komentar yang bersifat responsif dan *real-time*. Kolom komentar pada *Instagram* menjadi sumber data yang bernilai karena memuat berbagai bentuk tanggapan masyarakat secara spontan. Komentar-komentar tersebut dapat mencerminkan persepsi publik, tingkat penerimaan, kritik, maupun aspirasi terhadap suatu tokoh atau peristiwa tertentu.

Salah satu figur publik yang aktif memanfaatkan *Instagram* sebagai media komunikasi publik adalah Dedi Mulyadi. Melalui akun *Instagram* resminya, Dedi Mulyadi secara rutin membagikan kegiatan sosial, kebijakan, kunjungan lapangan, serta interaksi langsung dengan masyarakat. Aktivitas tersebut mendorong tingginya partisipasi publik dalam bentuk komentar yang beragam, mulai dari dukungan, pujian, saran, pertanyaan, hingga kritik.

Fenomena tingginya interaksi publik terhadap Dedi Mulyadi juga mendapat perhatian media nasional. Berdasarkan laporan CNN

Indonesia, popularitas Dedi Mulyadi di media sosial tergolong sangat tinggi dan menunjukkan dominasi interaksi publik yang signifikan dibandingkan akun resmi pemerintah daerah [1]. Selain itu, Kompas.com menyebutkan bahwa Dedi Mulyadi kerap dijuluki sebagai “gubernur konten” karena aktivitasnya yang intens dalam membagikan berbagai kegiatan melalui media sosial serta kemampuannya menarik perhatian masyarakat luas [2]. Fenomena tersebut menunjukkan bahwa media sosial kini tidak hanya berfungsi sebagai alat komunikasi, tetapi juga sebagai ruang pembentukan opini publik yang dapat memengaruhi citra dan legitimasi figur politik di era digital.

Meskipun demikian, ribuan komentar yang muncul pada media sosial masih berbentuk data teks tidak terstruktur. Data tersebut sulit dianalisis secara manual karena jumlahnya besar, terus bertambah, serta mengandung bahasa informal, singkatan, emoji, campuran bahasa daerah, hingga ungkapan *sarkastik*. Oleh karena itu, diperlukan pendekatan otomatis melalui analisis sentimen berbasis *Natural Language Processing (NLP)* untuk mengelompokkan opini masyarakat ke dalam kategori positif, netral, dan negatif.

Namun, salah satu permasalahan utama dalam analisis sentimen media sosial adalah ketidakseimbangan data (*imbalanced dataset*). Distribusi komentar pada media sosial sering kali didominasi oleh satu kelas sentimen tertentu, sedangkan kelas lainnya memiliki jumlah yang jauh lebih sedikit. Berdasarkan observasi awal pada dataset penelitian ini, diperoleh distribusi sentimen sebanyak 1.763 komentar positif, 672 komentar netral, dan 182 komentar negatif. Kondisi tersebut menunjukkan dominasi kelas positif yang cukup tinggi dibandingkan kelas lainnya.

Ketidakseimbangan data menjadi persoalan penting karena dapat menyebabkan

model klasifikasi lebih banyak mempelajari pola dari kelas mayoritas dan mengabaikan kelas minoritas. Akibatnya, model mungkin menghasilkan nilai akurasi yang terlihat tinggi, tetapi kurang mampu mengenali komentar netral maupun negatif secara optimal. Dalam konteks penelitian ini, komentar negatif memiliki nilai strategis karena sering kali berisi kritik, keluhan, dan evaluasi masyarakat terhadap figur publik. Jika kelas minoritas tidak terdeteksi dengan baik, maka hasil analisis sentimen berpotensi bias dan tidak merepresentasikan opini masyarakat secara menyeluruh.

Untuk mengatasi permasalahan tersebut, diperlukan strategi penanganan data tidak seimbang. Salah satu metode yang umum digunakan adalah *oversampling*, yaitu teknik menambah jumlah data pada kelas minoritas agar distribusi dataset menjadi lebih seimbang. Dalam penelitian ini, *oversampling* dilakukan melalui pendekatan *augmentasi* teks seperti *random swap*, *random deletion*, dan *synonym replacement*. Penerapan metode ini diharapkan dapat meningkatkan kemampuan model dalam mengenali pola sentimen pada kelas minoritas.

Dalam proses klasifikasi sentimen, pemilihan model juga menjadi faktor penting. Salah satu pendekatan *deep learning* yang banyak digunakan dalam analisis teks adalah *Bidirectional Long Short-Term Memory (BiLSTM)*. Model *BiLSTM* memiliki keunggulan karena mampu memahami konteks kata dari dua arah, yaitu dari urutan awal ke akhir dan dari akhir ke awal secara simultan. Kemampuan tersebut membuat *BiLSTM* lebih efektif dalam memproses teks media sosial yang cenderung pendek, informal, dan bergantung pada konteks kalimat.

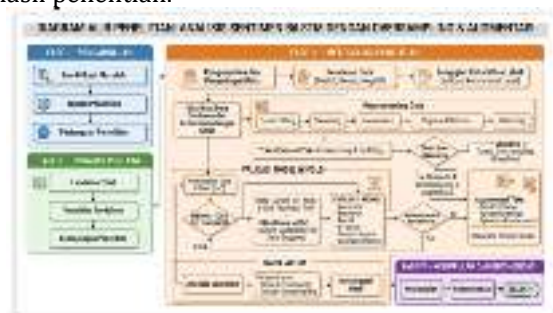
Sejumlah penelitian terdahulu menunjukkan bahwa model *BiLSTM* memiliki performa yang baik dalam klasifikasi sentimen pada data media sosial. Namun, sebagian besar penelitian masih berfokus pada isu layanan publik, ulasan produk, *cyberbullying*, maupun kebijakan nasional. Masih belum banyak penelitian yang secara khusus menganalisis sentimen publik terhadap figur politik Indonesia menggunakan data komentar *Instagram* dengan model *BiLSTM* serta menguji pengaruh ketidakseimbangan data melalui teknik *oversampling*.

Kondisi tersebut menunjukkan adanya *research gap* yang relevan untuk diteliti. Mengingat opini publik digital semakin sering dijadikan indikator penerimaan sosial terhadap figur pemerintahan, maka analisis sentimen

terhadap komentar masyarakat pada akun media sosial figur publik menjadi penting untuk dikaji secara ilmiah.

MATERIALS AND METHODS

Perlu dipahami bahwa penelitian analisis sentimen memerlukan alur kerja yang sistematis agar proses pengolahan data, pelatihan model, hingga evaluasi dapat dilakukan secara terstruktur dan menghasilkan model yang optimal. Pada penelitian ini digunakan pendekatan *Bidirectional Long Short-Term Memory (BiLSTM)* yang dikombinasikan dengan teknik *oversampling* dan *augmentasi* teks untuk mengatasi permasalahan ketidakseimbangan distribusi kelas sentimen. Diagram alir penelitian berikut menggambarkan tahapan penelitian mulai dari identifikasi masalah, pengumpulan data, preprocessing, pelatihan model, validasi menggunakan K-Fold, hingga interpretasi hasil penelitian.



Gambar 1. Alir Penelitian

Gambar tersebut menunjukkan keseluruhan alur penelitian analisis sentimen menggunakan model *BiLSTM* dengan penerapan *oversampling* dan *augmentasi* data teks. Tahapan penelitian dibagi menjadi empat fase utama. Fase pertama merupakan tahap pendahuluan yang meliputi identifikasi masalah, tujuan penelitian, dan penyusunan pertanyaan penelitian sebagai dasar pelaksanaan penelitian. Selanjutnya pada fase kedua dilakukan tinjauan pustaka yang mencakup landasan teori, penelitian terdahulu, serta identifikasi kesenjangan penelitian yang menjadi alasan dilakukannya penelitian ini.

Pada fase ketiga ditampilkan metodologi penelitian secara rinci. Tahap awal dimulai dari pengumpulan dan penyaringan data, kemudian dilakukan proses pelabelan sentimen ke dalam kategori positif, netral, dan negatif. Setelah itu dilakukan pengujian reliabilitas label untuk memastikan konsistensi hasil pelabelan data. Data yang telah siap kemudian diproses pada tahap preprocessing yang meliputi case folding, cleansing, normalisasi, stopword removal, dan stemming. Hasil preprocessing dilanjutkan ke tahap transformasi teks menggunakan tokenizing dan

padding agar data dapat diproses oleh model deep learning.

Diagram juga menunjukkan adanya dua skenario pengujian, yaitu skenario tanpa oversampling sebagai baseline dan skenario dengan oversampling serta augmentasi teks. Teknik augmentasi yang digunakan meliputi random swap, random deletion, dan synonym replacement untuk meningkatkan variasi data pada kelas minoritas. Selanjutnya dilakukan validasi model menggunakan metode K-Fold Cross Validation dengan nilai $k=5$. Pada setiap iterasi fold, data dibagi menjadi data training dan testing, kemudian model BiLSTM dilatih menggunakan mekanisme early stopping untuk mencegah overfitting.

Tahap evaluasi model dilakukan menggunakan beberapa metrik seperti accuracy, precision, recall, F1-score, Cohen's Kappa, dan confusion matrix. Setelah seluruh fold selesai diproses, dilakukan perhitungan rata-rata hasil evaluasi untuk membandingkan performa model pada kedua skenario. Tahap akhir penelitian menghasilkan analisis sentimen, interpretasi hasil, serta penarikan kesimpulan dan rekomendasi berdasarkan performa model yang diperoleh. Dengan adanya diagram alir ini, proses penelitian dapat dipahami secara lebih sistematis dan terstruktur mulai dari pengolahan data hingga evaluasi akhir model.

RESULTS AND DISCUSSION

Setelah data komentar melalui tahap *preprocessing*, langkah berikutnya adalah melakukan transformasi data teks. Tahap ini bertujuan untuk mengubah data teks hasil *preprocessing* menjadi representasi numerik agar dapat diproses oleh model *Bidirectional Long Short-Term Memory (BiLSTM)*. Model *deep learning* tidak dapat memproses data dalam bentuk teks mentah, sehingga setiap kata harus diubah menjadi bentuk angka yang terstruktur.

Data yang digunakan pada tahap ini merupakan hasil *preprocessing* sebanyak 2.617 komentar yang terdiri atas tiga kelas sentimen, yaitu 1.763 komentar positif, 672 komentar netral, dan 182 komentar negatif. Distribusi data sentimen ditunjukkan pada Tabel 1

Tabel 1 Distribusi Data Setelah Preprocessing

Kategori Sentimen	Jumlah Data
Positif	1.763
Netral	672
Negatif	182

Total	2.617
--------------	-------

Berdasarkan Tabel 1 terlihat bahwa kelas positif memiliki jumlah data paling dominan dibandingkan kelas netral dan negatif. Kondisi ini menunjukkan bahwa dataset bersifat tidak seimbang (*imbalanced dataset*), sehingga pada tahap eksperimen dilakukan pengujian dua skenario, yaitu tanpa *oversampling* dan dengan *oversampling*.

Transformasi data teks pada penelitian ini dilakukan melalui beberapa tahapan sebagai berikut:

Tokenizing merupakan proses pemecahan kalimat menjadi unit-unit kata (token). Setiap komentar yang telah dibersihkan akan dipisahkan menjadi daftar kata agar dapat dikenali oleh sistem. Contoh hasil *tokenizing* ditunjukkan pada Tabel 2

Tabel 2 Hasil Tokenizing

No	Teks Hasil Preprocessing	Hasil Tokenizing
1	mantap pak dedy	[mantap, pak, dedy]
2	kritik terima lapang dada keren	[kritik, terima, lapang, dada, keren]
3	jalan parung bogor rusak tolong baik	[jalan, parung, bogor, rusak, tolong, baik]
4	terus nyala bapak baik hati	[terus, nyala, bapak, baik, hati]
5	sehat sukses selalu pak	[sehat, sukses, selalu, pak]
6	sekolah daerah sangat bagus	[sekolah, daerah, sangat, bagus]
7	terima kasih bantu masyarakat	[terima, kasih, bantu, masyarakat]
8	kerennnn total bapakkkk salam jogja	[kerennnn, total, bapakkkk, salam, jogja]
9	orang komen luar jawa barat mulu bagus	[orang, komen, luar, jawa, barat, mulu, bagus]
10	gubernur baik indonesia	[gubernur, baik, indonesia]

Melalui proses *tokenizing*, teks yang semula berupa kalimat diubah menjadi susunan kata yang dapat diproses lebih lanjut.

Integer Encoding

Setelah *tokenizing*, setiap kata diubah menjadi bilangan bulat menggunakan Keras *Tokenizer*. Sistem akan membuat kamus kata (*vocabulary*), kemudian setiap kata diberi indeks numerik berdasarkan urutan kemunculan atau frekuensi.

Sebagai contoh:

Tabel 3 Contoh Vocabulary Index Tokenizer

Kata	Indeks
pak	1
keren	2
mantap	3
baik	4

Setelah proses *tokenizing* dilakukan, setiap kata yang terdapat pada dataset selanjutnya diubah menjadi representasi numerik melalui tahap *integer encoding* menggunakan Keras *Tokenizer*. Pada tahap ini, sistem membentuk kamus kata (*vocabulary*) dari seluruh data teks, kemudian memberikan indeks unik pada setiap kata berdasarkan urutan kemunculan atau frekuensi kata dalam dataset. Dengan demikian, setiap kata akan direpresentasikan dalam bentuk bilangan bulat sehingga dapat diproses oleh model *deep learning*. Contoh kamus kata dan indeks numerik yang dihasilkan *tokenizer* ditunjukkan pada Tabel 3

Tabel 3 Contoh Integer Encoding

No	Teks	Hasil Tokenizing	Hasil Integer Encoding
1	[mantap, dedy]	pak,	[3, 1, 25]
2	[kritik, lapang, keren]	terima, dada,	[8, 5, 17, 29, 2]
3	[gubernur, indonesia]	baik,	[12, 4, 10]
4	[jalan, bogor, tolong, baik]	parung, rusak,	[7, 33, 18, 14, 21, 4]
5	[terus, bapak, baik, hati]	nyala,	[15, 41, 11, 4, 26]
6	[sehat, selalu, pak]	sukses,	[13, 20, 32, 1]
7	[sekolah, sangat, bagus]	daerah,	[9, 24, 16, 6]
8	[terima, bantu, masyarakat]	kasih,	[5, 19, 22, 27]

9	[kerennnn, total, bapakkkk, joga]	[34, 30, 45, 28, 37]
10	[orang, komen, luar, jawa, barat, mulu, bagus]	[23, 31, 38, 35, 39, 44, 6]

Berdasarkan Tabel 3 setiap kata hasil *tokenizing* telah diubah menjadi representasi angka menggunakan Keras *Tokenizer*. Setiap kata memperoleh indeks numerik tertentu berdasarkan kamus kata (*vocabulary*) yang terbentuk dari seluruh dataset. Kata yang sama akan selalu memiliki indeks yang sama pada seluruh data. Proses *integer encoding* ini bertujuan agar data teks dapat dibaca oleh model *deep learning*, karena model hanya dapat memproses *input* dalam bentuk numerik. Hasil *encoding* selanjutnya digunakan pada tahap *padding* untuk menyeragamkan panjang setiap *sequence* data.

Padding

Panjang setiap komentar pada dataset tidak sama. Oleh karena itu, dilakukan *padding* untuk menyeragamkan panjang *input*. Pada penelitian ini digunakan panjang maksimum *sequence* sebesar 100 *token*.

Skenario Eksperimen Tanpa Oversampling

Pada skenario ini, model dilatih menggunakan data asli tanpa dilakukan penyeimbangan data. Dataset yang digunakan berjumlah 2.617 komentar dengan distribusi 1.763 data positif, 672 data netral, dan 182 data negatif. Hasil pengujian pada masing-masing *fold* ditunjukkan pada Tabel 4.

Tabel 4 Hasil Akurasi K-Fold Tanpa Oversampling

Fold	Akurasi
Fold 1	0,7233
Fold 2	0,7595
Fold 3	0,7476
Fold 4	0,7055
Fold 5	0,6960
Rata - rata	0,7264

Berdasarkan Tabel 4, akurasi tertinggi diperoleh pada *Fold 2* sebesar 0,7595, sedangkan akurasi terendah terdapat pada *Fold 5* sebesar 0,6960. Secara keseluruhan, model menghasilkan rata-rata akurasi sebesar 0,7264 atau 72,64%. Selain akurasi, evaluasi juga dilakukan menggunakan *precision*, *recall*, *F1-score*, dan *Cohen's Kappa*. Ringkasan hasil evaluasi rata-rata ditunjukkan pada Tabel 4.19.

Tabel 5 Rata-rata Metrik Evaluasi Tanpa Oversampling

Metrik	Nilai
--------	-------

Accuracy	0,7264
Precision	0,7014
Recall	0,7264
F1-Score	0,7066
Cohen's Kappa	0,3722

Nilai *precision* sebesar 0,7014 menunjukkan bahwa model cukup baik dalam memberikan prediksi yang relevan terhadap kelas sentimen. Nilai *recall* sebesar 0,7264 menunjukkan kemampuan model dalam mengenali data aktual dari masing-masing kelas. Nilai *F1-score* sebesar 0,7066 menunjukkan keseimbangan yang cukup baik antara *precision* dan *recall*. Sementara itu, nilai *Cohen's Kappa* sebesar 0,3722 menunjukkan tingkat kesepakatan model pada kategori *fair agreement* hingga *moderate agreement*. Berdasarkan *classification report* pada beberapa *fold*, model memiliki performa terbaik dalam mengenali kelas positif, sedangkan kesalahan klasifikasi masih cukup sering terjadi pada kelas netral dan negatif. Hal ini disebabkan jumlah data kelas positif yang jauh lebih dominan, sehingga model lebih banyak mempelajari pola dari kelas mayoritas.

Secara umum, hasil eksperimen tanpa *oversampling* menunjukkan bahwa model *BiLSTM* mampu melakukan klasifikasi sentimen dengan performa cukup baik meskipun dilatih menggunakan dataset tidak seimbang. Namun, masih terdapat kecenderungan bias terhadap kelas mayoritas, terutama dalam mengenali komentar negatif yang jumlah datanya paling sedikit.

Oleh karena itu, pada tahap selanjutnya dilakukan skenario eksperimen kedua dengan menerapkan teknik *oversampling* berbasis *augmentasi teks* untuk mengetahui apakah penyeimbangan data mampu meningkatkan performa model, terutama pada kelas minoritas.

CONCLUSION

Berdasarkan hasil penelitian mengenai analisis sentimen komentar publik terhadap figur Dedi Mulyadi menggunakan model *Bidirectional Long Short-Term Memory (BiLSTM)*, maka diperoleh kesimpulan sebagai berikut:

Kecenderungan sentimen publik pada komentar *Instagram* terhadap Dedi Mulyadi didominasi oleh sentimen positif. Dari total 2.617 komentar yang dianalisis, sebanyak 1.763 komentar atau 67,37% termasuk kategori

positif, sebanyak 672 komentar atau 25,68% termasuk kategori netral, dan sebanyak 182 komentar atau 6,95% termasuk kategori negatif. Hasil tersebut menunjukkan bahwa persepsi masyarakat terhadap figur Dedi Mulyadi pada dataset penelitian cenderung positif.

Model *Bidirectional Long Short-Term Memory (BiLSTM)* berhasil diterapkan untuk mengklasifikasikan komentar publik *Instagram* terhadap Dedi Mulyadi ke dalam tiga kategori sentimen, yaitu positif, netral, dan negatif. Penerapan model dilakukan melalui dua skenario eksperimen, yaitu skenario tanpa *oversampling* menggunakan distribusi data asli dan skenario dengan *oversampling* menggunakan teknik *augmentasi teks* berupa *random swap*, *random deletion*, dan *synonym replacement*. Model mampu memproses komentar yang bersifat singkat, informal, dan beragam sehingga dapat digunakan pada analisis sentimen media sosial berbahasa Indonesia.

Berdasarkan hasil evaluasi performa menggunakan metrik *accuracy*, *precision*, *recall*, *F1-score*, dan *Cohen's Kappa*, skenario dengan *oversampling* memberikan hasil terbaik dibandingkan skenario tanpa *oversampling*. Skenario dengan *oversampling* memperoleh *accuracy* sebesar 74,24%, *precision* sebesar 0,7432, *recall* sebesar 0,7424, *F1-score* sebesar 0,7415, dan *Cohen's Kappa* sebesar 0,4545. Sementara itu, skenario tanpa *oversampling* memperoleh *accuracy* sebesar 72,64%, *precision* sebesar 0,7014, *recall* sebesar 0,7264, *F1-score* sebesar 0,7066, dan *Cohen's Kappa* sebesar 0,3722. Hasil ini menunjukkan bahwa penerapan *oversampling* mampu meningkatkan performa model *BiLSTM* pada dataset yang memiliki ketidakseimbangan kelas.

REFERENCE

- [1] Litbang Kompas, "Litbang Kompas: KDM Kuat di Medsos Meski Warga Tak Puas Lapangan Kerja BacaLitbang Kompas: KDM Kuat di Medsos Meski Warga Tak Puas Lapangan Kerja," CNN Indonesia.
- [2] Kompas.tv, "Dedi Mulyadi soal Dijuluki 'Gubernur Konten': Itu Sebenarnya Pujian," Kompas. TV.