

SPEECH RECOGNITION UNTUK KLASIFIKASI PENGUCAPAN NAMA HEWAN DALAM BAHASA SUNDA MENGUNAKAN METODE *LONG-SHORT TERM MEMORY*

Nur Aini Lailla Asri, Riza Ibnu Adam, Budi Arif Dermawan

Program Studi Informatika, Fakultas Ilmu Komputer Universitas Singaperbangsa Karawang
1910631170110@student.unsika.ac.id

ABSTRAK

Keberagaman suku, ras, agama, bahasa, dan adat istiadat membuat Indonesia memiliki kebudayaan yang dijalin erat dengan bahasa. Menurut data Badan Pusat Statistik (BPS) persentase penutur bahasa daerah semakin menurun pada generasi muda. Selain menurunnya penutur bahasa daerah juga terjadi kesalahan dalam pelafalan bahasa daerah. Kesalahan pelafalan ini dapat diatasi dengan meningkatkan latihan berbahasa, terutama membaca dan mendengarkan kosa kata bahasa daerah. Kosa kata yang dapat dilatih misalnya nama-nama hewan dalam bahasa daerah. Selain menambah kosa kata, dengan mengenal nama hewan dapat meningkatkan kecerdasan natural anak. Metode *deep learning* dapat digunakan untuk mengatasi pergeseran bahasa daerah salah satunya adalah menggunakan *speech recognition*. Metode *Long-Short Term Memory* dapat digunakan untuk klasifikasi suara pelafalan nama hewan dalam bahasa Sunda. Dataset yang digunakan adalah dataset baru dengan 16 cara yang berbeda dalam pengambilan data suara. Data terdiri dari 100 nama hewan yang digunakan sebagai kelas dalam proses klasifikasi. Hasil akurasi terbaik mencapai 97,50% didapat dengan menggunakan *epoch* 150 dan *batch size* 30. Pada proses *testing* menggunakan data dari dataset, dapat memprediksi semua nama-nama hewan yang digunakan sebagai data kelas. Namun program belum dapat memprediksi dengan tepat apabila menggunakan data dari luar dataset dikarenakan jumlah data suara yang digunakan sebagai dataset belum cukup banyak.

Kata kunci: Bahasa Sunda, Klasifikasi, LSTM

1. PENDAHULUAN

Indonesia memiliki keanekaragaman suku, ras, agama, bahasa, dan adat istiadat. Dengan keberagaman tersebut tumbuhlah berbagai kebudayaan yang dijalin erat dengan bahasa. Menurut data Badan Pusat Statistik (BPS) pada tahun 2011, suku bangsa di Indonesia mencapai 1.340 suku. Bangsa suku Sunda dengan 36,7 juta orang atau setara dengan 15,5% dari populasi penduduk Indonesia, menjadi penduduk kedua paling banyak setelah suku bangsa Jawa. Berdasarkan data BPS presentase penutur bahasa daerah semakin menurun pada generasi lebih muda. Faktor penyebab pergeseran bahasa menurut [1] dan [2] diantaranya karena faktor pendidikan, keluarga, status sosial, tempat tinggal, migrasi, ekonomi, dan multibahasa dalam keluarga. Menurut hasil penelitian yang dilakukan oleh Balai Bahasa Provinsi Jawa Barat (BBPJB) Kementerian dan Kebudayaan menyebutkan bahwa bahasa Sunda terancam punah dan di usia anak-anak hanya sekitar 40% yang memahami dan dapat berbahasa Sunda di Jawa Barat (Republika.co.id, 2013). Selain menurunnya penutur kata bahasa daerah, kesalahan dalam pelafalan bahasa juga terjadi. Faktor terjadinya kesalahan dalam pelafalan bahasa daerah antara lain karena faktor kebiasaan lingkungan, tidak adanya pelajaran bahasa daerah dan kurangnya keinginan untuk mempelajari bahasa daerah. Kesalahan pelafalan ini dapat diatasi dengan meningkatkan latihan berbahasa, terutama membaca dan mendengarkan kosa kata bahasa daerah kemudian melafalkannya [3].

Metode *deep learning* dapat digunakan untuk mengatasi pergeseran bahasa daerah salah satunya adalah *speech recognition*. *Speech recognition* adalah teknologi yang dapat mengonversikan perkataan yang berupa sinyal suara menjadi tulisan. Selain untuk mengidentifikasi ucapan, juga dapat dimanfaatkan untuk mempelajari bahasa baru [4]. Penelitian tentang pengenalan ucapan telah dilakukan sebelumnya, penelitian tersebut melakukan pendeteksian pengucapan bahasa Arab dengan metode *Convolutional Neural Network* (CNN) menggunakan 2.892 alfabet Arab. Model CNN memperoleh akurasi pengujian sebesar 100%. Namun terdapat permasalahan bahwa model mulai mengenali data suara yang spesifik dalam set *training* daripada *generalizing*, hal tersebut menyebabkan *validation loss* terus meningkat setiap *epoch* [5]. Pada penelitian pendeteksi kalimat *hoax* dari *Twitter* menggunakan metode LSTM-CNN dengan membandingkan hasil pengolahan data menggunakan *regularization* dan tanpa *regularization*, hasil pengolahan tanpa melakukan *regularization* mendapat nilai *validation loss* lebih kecil daripada yang melakukan *regularization*. Pada penelitian tersebut menunjukkan nilai *training loss* menurun tetapi nilai *validation loss* meningkat, hal tersebut menunjukkan bahwa model tanpa *regularization* mengalami *overfitting* [6]. Penelitian tentang teknik *regularization* pada *deep learning* menjelaskan bahwa model yang menggunakan teknik *regularization* menghasilkan kinerja lebih baik daripada model tanpa *regularization* dalam hal *training errors*. Hasil penelitian tersebut

menjelaskan bahwa metode *regularization* dapat menyelesaikan masalah *overfitting* dan *underfitting* secara efisien, tetapi pada kondisi penelitian tersebut *L1 regularization* dan skema *autoencoder* masih mengalami *overfitting* dan *underfitting* [7]. Penelitian tentang RNN yang sering terjadi *vanishing gradient* ketika proses training telah dilakukan, masalah tersebut dapat diselesaikan dengan menggunakan metode lanjutan yang dapat bekerja dengan baik pada masalah jeda waktu yang lama seperti metode *Long Short Term Memory* (LSTM) [8]. Kemampuan metode LSTM pernah dibandingkan dengan metode *Gated Recurrent Unit* (GRU). Hasil dari penelitian tersebut menyebutkan bahwa nilai akurasi dengan metode LSTM lebih tinggi, yakni sebesar 73% dan 64% jika menggunakan metode GRU [9].

Berdasarkan pembahasan tersebut, maka penelitian ini akan melakukan klasifikasi pelafalan fonem dalam bahasa daerah khususnya bahasa Sunda dengan menggunakan *speech recognition* yang menerapkan *deep learning* menggunakan metode *Long-Short Term Memory* (LSTM) dalam pemrosesan suara. Ada pun kata yang digunakan untuk *dataset* yakni menggunakan nama hewan dalam bahasa Sunda.

2. TINJAUAN PUSTAKA

2.1. Bahasa Sunda

Bahasa Sunda merupakan bahasa daerah yang digunakan di daerah Jawa Barat, Sebagian Jawa Tengah sebelah barat, dan Sebagian Daerah Ibu Kota Jakarta [4]. Bahasa Sunda merupakan suatu hasil kebiasaan masyarakat di tanah Sunda yang tercipta dari proses interaksi masyarakat dan menjadi ciri khas dari budaya Sunda itu sendiri [10].

2.2. Speech Recognition

Speech recognition atau *automatic speech recognition* (ASR) memungkinkan untuk membuat mesin paham dengan apa yang dilafalkan manusia. Input ucapan akan diproses menjadi sinyal dengan cara mengonversi gelombang suara menjadi sekumpulan angka yang dapat dipahami komputer dan diidentifikasi sebagai kata-kata [11].

2.3. Deep Learning

Deep Learning adalah cabang dari *machine learning* yang berfokus pada kecerdasan buatan di mana algoritma tersebut mempunyai sekumpulan algoritma yang berusaha memodelkan data pada abstraksi *high-level* pada data dengan menerapkan lapisan pemrosesan dengan struktur yang kompleks dan sebaliknya serta terdiri dari sejumlah transformasi non-linear. Algoritma ini didasari dari representasi *machine learning*. Perbedaan antara *deep learning* dengan *normal neural network* yaitu *deep learning* mempunyai *hidden layer*. *Deep learning* dapat melakukan *training* dengan diawasi atau tanpa diawasi. Ada banyak varian dari arsitektur *deep learning* ini, seperti *deep neural network*,

convolutional deep neural network, *recurrent neural network*, dan *deep belief network* [4].

2.4. Long Short Term Memory (LSTM)

Long Short-Term Memory (LSTM) merupakan salah satu jenis *Recurrent Neural Network* (RNN) yang divariasi dengan menerapkan *memory cell* yang dapat menyimpan informasi dalam waktu yang panjang. Dalam pemrosesan data sekuensial yang panjang, LSTM dianjurkan digunakan sebagai solusi mengatasi *vanishing gradient* pada RNN [12]. *Vanishing gradient* merupakan kondisi jaringan yang terlalu dalam, sehingga menyebabkan gradien dari *loss function* mengecil hingga mendekati nol setelah beberapa iterasi sebelum mencapai konvergen [13].

2.5. Dataset

Dataset adalah gabungan dari data yang menjelaskan suatu topik tertentu. Dataset merupakan salah satu faktor kesuksesan dalam pembuatan *machine learning* atau *deep learning*. Dataset digunakan untuk referensi pembelajaran dari model yang akan dibuat terhadap keputusan yang akan diambil [14].

2.6. File WAV

File WAV merupakan format audio standar Microsoft dan IBM untuk PC (*Personal Computer*). Biasanya dikodekan menggunakan PCM (*Pulse Code Modulation*). WAV merupakan data yang tidak terkompresi sehingga semua sampel disimpan di *harddisk*. File WAV jarang digunakan di internet karena ukurannya relatif besar dan maksimal file Wav 2GB [15].

2.7. Mel-Frequency Cepstral Coefficient (MFCC)

Mel-Frequency Cepstral Coefficients (MFCC) adalah salah satu metode yang sering digunakan dalam pengolahan suara yang memanfaatkan *speech recognition*. MFCC berfungsi sebagai *feature extraction*, yakni proses ekstraksi sinyal suara menjadi sejumlah parameter. Analisis suara yang dilakukan MFCC dilatari oleh persepsi pendengaran manusia, karena pendengaran manusia dapat memfilter frekuensi tertentu [11].

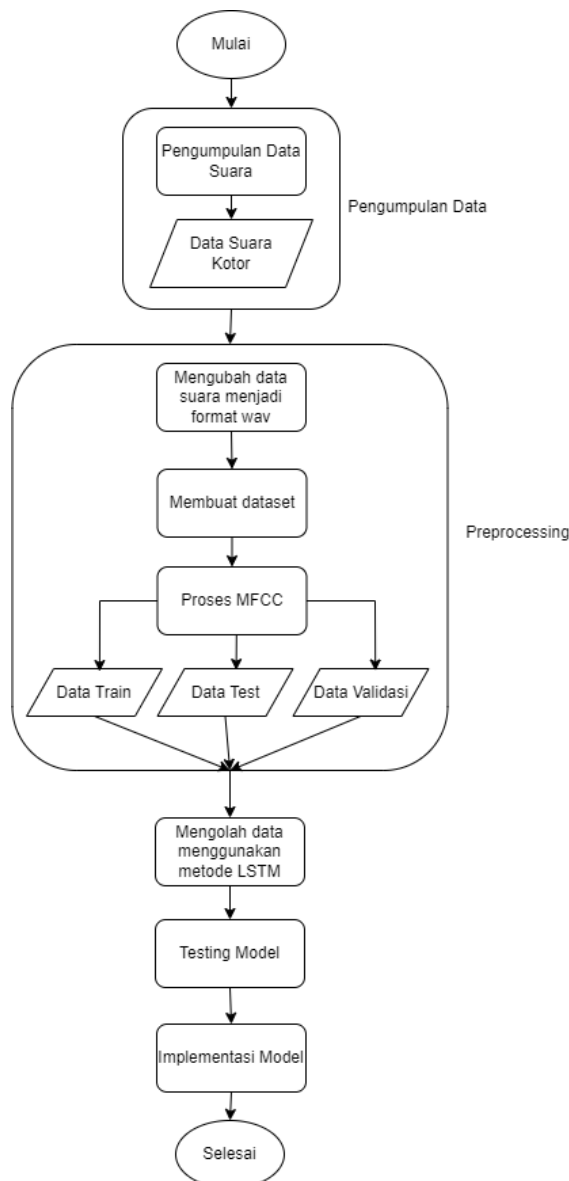
2.8. Python

Python merupakan bahasa pemrograman yang mempunyai kode sangat jelas, lengkap, dan mudah untuk dipahami. Python dianggap mampu menyelesaikan masalah *big data*, *data mining*, *deep learning*, *data science*, dan *machine learning*. Python disebut sebagai bahasa pemrograman multi-paradigma karena berbentuk pemrograman berorientasi objek, imperatif dan fungsional [16].

3. METODE PENELITIAN

Penelitian akan dilakukan melalui lima tahapan yakni pengumpulan data, *preprocessing*, *training model*, *testing model*, dan implementasi model.

Berikut *flowchart* proses klasifikasi suara yang digunakan dalam penelitian ini dapat ditunjukkan dalam Gambar 1.



Gambar 1. *Flowchart* proses klasifikasi suara

4. HASIL DAN PEMBAHASAN

4.1. Pengumpulan Data

Penelitian ini melakukan pembaruan dataset. Dataset diambil menggunakan suara anak laki-laki, anak perempuan, laki-laki dewasa, dan perempuan dewasa. Tahapan pengumpulan data untuk proses pembuatan dataset dimulai dengan mempersiapkan data nama-nama hewan yang akan digunakan dalam pelafalan yang nantinya digunakan untuk kelas dalam klasifikasi. Dalam penelitian ini menggunakan 100 nama hewan dalam bahasa Sunda, hal tersebut bertujuan untuk mengetahui performa LSTM dalam melakukan klasifikasi multi kelas. Dari 100 nama

hewan diambil data suara menggunakan 16 cara yang berbeda sehingga jumlah keseluruhan dataset adalah 1.600 data suara. Tujuan dari pengambilan suara dari *device* yang berbeda dan jarak yang berbeda adalah untuk mendapatkan *noise* yang berbeda dari masing-masing cara pengambilan.

4.2. Preprocessing

Pada proses ini data suara disiapkan sebelum digunakan untuk pemrosesan klasifikasi suara. Setelah pengambilan data suara, data diberi nama sesuai dengan nama hewan yang disebutkan, pelafalnya, *device* yang digunakan, dan jarak pengambilannya dengan format *nama_hewan-pelafal-device-jarak*. Pada tahap pengambilan data suara, format file rekaman yang dihasilkan setiap *device* berbeda-beda seperti format mp3, aac, dan M4A sehingga perlu untuk menyamakan format file suara tersebut. Dalam penelitian pengolahan data suara menggunakan format wav karena file wav tidak terkompres sehingga semua sampel dapat disimpan. Setelah format file disamakan, maka data dibagi sesuai kelas yang telah ditentukan yakni sesuai dengan nama-nama hewan yang digunakan. Hasil dataset yang telah dibagi dimasukkan ke Google Drive. Setelah data suara siap digunakan, data akan dimuat ke dalam program untuk proses pengolahan data suara. Pada penelitian ini melakukan pemrosesan data menggunakan Google Colaboratory. Dataset melalui proses ekstraksi MFCC sebelum dimasukkan ke data *array file_animal_data*. Ekstraksi MFCC bertujuan untuk mendapat nilai-nilai parameter. Label kelas hasil pembersihan nama file pada saat memuat dataset diubah dari yang berbentuk kategorial menjadi bentuk numerik. Pada proses split data, data dibagi untuk menjadi data *training*, validasi, dan *testing*.

4.3. Training Model

Pada tahap *training model*, dataset yang siap diolah akan digunakan dalam proses pelatihan model LSTM. Ada beberapa hal yang harus disesuaikan dengan jumlah kelas pada dataset yang digunakan dalam penelitian seperti jumlah *epoch* dan *batch size*. Pada penelitian ini menggunakan jumlah *epoch* 150 dan *batch size* sebesar 30. Hasil dari pelatihan program LSTM tersebut dapat di gambarkan dengan hasil *ploting training* dan *validation loss* serta *ploting training* dan *validation accuracy*. Hasil evaluasi disimpan dalam file dengan format .h5. Model tersebut akan disimpan apabila nilai *val_accuracy* pada *epoch* tersebut lebih besar dari nilai *val_accuracy* pada *epoch* sebelumnya.

Pada penelitian ini melakukan percobaan untuk menentukan nilai akurasi tertinggi dengan meakukan percobaan nilai unit, *epoch*, dan *batch size*. Untuk hasil percobaan menggunakan unit 128 dapat dilihat pada Tabel 1.

Tabel 1. Hasil percobaan menggunakan unit 128

Percobaan ke-	Unit	Epoch	Batch Size	Akurasi
1	128	50	32	49,38%
2	128	100	32	84,25%
3	128	150	32	93,25%
4	128	200	32	96,92%
5	128	250	32	97,33%

Dan hasil percobaan menggunakan unit 256 dapat dilihat pada Tabel 2.

Tabel 2. Hasil percobaan menggunakan unit 256

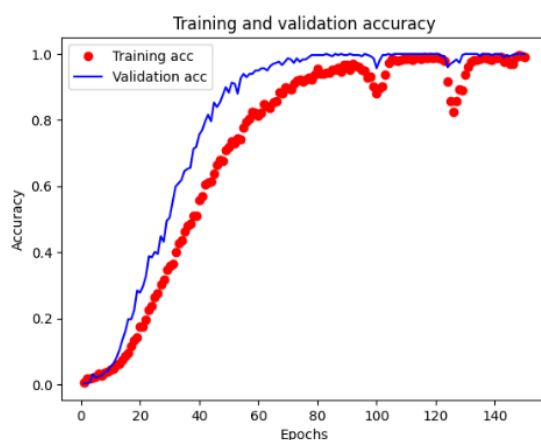
Percobaan ke-	Unit	Epoch	Batch Size	Akurasi
1	256	50	32	83,75%
2	256	100	32	97,42%
3	256	150	32	97,50%
4	256	200	32	97,42%
5	256	250	32	97%

Dari hasil percobaan unit dan *epoch* di atas dapat diketahui bahwa dengan unit 256 dan *epoch* sebesar 150 memiliki nilai akurasi paling tinggi. Sehingga untuk percobaan selanjutnya melakukan percobaan jumlah *batch size* dengan menggunakan unit 256 dan *epoch* 150.

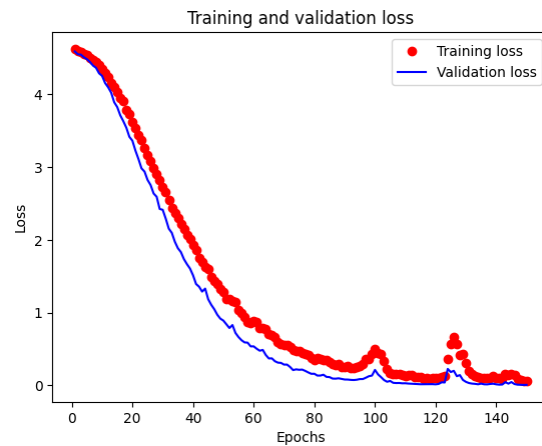
Tabel 3. Hasil percobaan batch size

Percobaan ke-	Unit	Epoch	Batch Size	Waktu Eksekusi	Akurasi
1	256	150	10	36 menit	97,50%
2	256	150	20	23 menit	97,50%
3	256	150	30	18 menit	97,50%
4	256	150	40	15 menit	91,58%
5	256	150	50	13 menit	96,08%

Dari hasil seluruh percobaan unit, *epoch*, dan *batch size*, pada penelitian ini menggunakan unit 256, *epoch* 150, dan *batch size* 30. Hasil training tersebut dapat digambarkan dalam plotting berikut



Gambar 2. Plotting training dan validation accuracy



Gambar 3. Plotting training dan validation loss

4.4. Testing Model

Pada proses ini, mengambil data suara dari dataset untuk *testing*, dimana dataset tersebut diambil secara acak dari dataset yang digunakan dan menggunakan data suara dari luar dataset. Pada proses ini menggunakan *function np.expand_dims()* untuk menambahkan dimensi baru pada posisi pertama (0), sehingga *array* berbentuk (1, jumlah data, jumlah fitur, 1). *Function* ini digunakan untuk memperhitungkan dimensi *batch* saat memprediksi dengan model. *Function argmax()* digunakan untuk mengambil nilai maksimum dari setiap prediksi dan menghasilkan prediksi kelas yang diklasifikasikan oleh model.

Selanjutnya kode yang disimpan di *model_A* disimpan dan di konversi menjadi format JSON dalam file 'speech_animal.json', *speech_animal.json* dimuat kembali untuk melanjutkan menggunakan model dari file tersebut untuk melakukan prediksi atau pelatihan. Dengan melakukan kompilasi model kembali dengan parameter yang serupa dengan model aslinya. Menggunakan fungsi kerugian atau *loss()* yang diatur menjadi *categorical_crossentropy*, *optimizer* menggunakan *RMSPProp*, metrik evaluasi diatur menjadi *categorical_accuracy*. Hal tersebut menyesuaikan untuk tugas klasifikasi multikelas. Hasil kompilasi tersebut dapat dilihat pada Gambar 4.

```

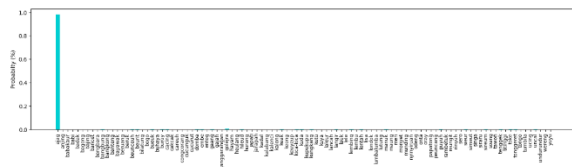
Model: "sequential"
-----
Layer (type)                Output Shape         Param #
-----
lstm (LSTM)                  (None, 256)          264192
dropout (Dropout)            (None, 256)           0
activation (Activation)       (None, 256)           0
dense (Dense)                (None, 100)          25700
activation_1 (Activation)     (None, 100)           0
-----
Total params: 289,892
Trainable params: 289,892
Non-trainable params: 0
None

```

Gambar 4. Hasil kompilasi

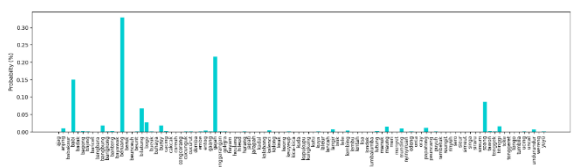
4.5. Implementasi Model

Hasil dari penelitian ini merupakan nilai akurasi dari klasifikasi pengucapan nama hewan dalam Bahasa Sunda. Hasil akurasi pada penelitian ini sebesar 97% dengan menggunakan 100 kelas. Proses *testing* pada Gambar 4.61 menggunakan data suara pelafalan nama hewan ajag untuk data prediksi, menghasilkan hasil prediksi sesuai dengan penyebutannya dengan diagram ditunjukkan dalam Gambar 5.



Gambar 5. Hasil prediksi menggunakan data dari dataset

Sedangkan untuk hasil testing dari luar dataset dapat ditunjukkan dengan Gambar 6.



Gambar 6. Hasil prediksi menggunakan data dari luar dataset

Kurangnya kemampuan model untuk memprediksi data suara dari luar dataset terjadi karena kurangnya keberagaman data pada setiap kelas. Hal tersebut pernah terjadi pada penelitian yang berjudul "Implementasi Deep Learning Menggunakan Metode CNN dan LSTM untuk Menentukan Berita Palsu dalam Bahasa Indonesia". Pada penelitian tersebut menyebutkan bahwa jumlah data yang digunakan belum cukup banyak, sehingga proses *training* dan pembentukan model kurang optimal [17].

5. KESIMPULAN DAN SARAN

Penelitian ini telah menerapkan metode LSTM untuk klasifikasi ucapan fonem nama hewan dalam bahasa Sunda. Penelitian ini melakukan pembaruan dataset pelafalan nama hewan dalam bahasa Sunda dan dataset tersebut dapat digunakan untuk klasifikasi menggunakan metode LSTM. Dengan menggunakan 100 nama hewan sebagai kelas pada dataset, program klasifikasi fonem bahasa Sunda menggunakan *epoch* 150 dan *batch size* 30 mendapat hasil cukup baik dalam klasifikasi suara dibandingkan dengan percobaan-percobaan lain yang dilakukan pada penelitian ini, dengan hasil nilai akurasi sebesar 97,50% dan waktu eksekusi *training* selama 18 menit. Pada penelitian selanjutnya disarankan untuk melakukan klasifikasi dengan metode LSTM yang lain, perpaduan metode, atau melakukan perbandingan metode untuk mengetahui manakah yang dapat melakukan klasifikasi pelafalan nama hewan dalam bahasa Sunda yang lebih baik. Walau nilai akurasi

yang didapat cukup tinggi, namun program belum dapat memprediksi suara dari luar dataset dengan baik sehingga masih diperlukan perbaikan dengan menambah dataset agar program belajar lebih banyak variasi suara.

DAFTAR PUSTAKA

- [1] W. P. Bhakti and I. Pekalongan, "Pergeseran Penggunaan Bahasa Jawa ke Bahasa Indonesia dalam Komunikasi Keluarga di Sleman," PBSI UPY, 2020.
- [2] R. Mustikasari and C. W. Astuti, "Pergeseran Penggunaan Bahasa Jawa pada Siswa TK dan KB di Kelurahan Beduri Ponorogo," 2020. [Online]. Available: <http://jurnal.unsur.ac.id/ajbsi>.
- [3] S. Ita Dhamina and L. Irma Wanti, "Kesalahan Pelafalan Fonem Bahasa Jawa Siswa Kelas Mengengah di Ponorogo," 2022.
- [4] D. N. Fathurrahman, A. B. Osmond, and R. E. Saputra, "Deep Neural Network untuk Pengenalan Ucapan pada Bahasa Sunda Dialek Tengah Timur (Majalengka)," 2018.
- [5] F. Alqadheeb, A. Asif, and H. F. Ahmad, "Correct Pronunciation Detection for Classical Arabic Phonemes Using Deep Learning," in *2021 International Conference of Women in Data Science at Taif University, WiDSTaif 2021*, Institute of Electrical and Electronics Engineers Inc., Mar. 2021. doi: 10.1109/WIDSTAIF52235.2021.9430236.
- [6] P. N. Anggreyani and W. Maharani, "Hoax Detection Tweets of the COVID-19 on Twitter Using LSTM-CNN with Word2Vec," *JURNAL MEDIA INFORMATIKA BUDIDARMA*, vol. 6, no. 4, p. 2432, Oct. 2022, doi: 10.30865/mib.v6i4.4564.
- [7] I. Nusrat and S. B. Jang, "A Comparison of Regularization Techniques in Deep Neural Networks," *Symmetry (Basel)*, vol. 10, no. 11, Nov. 2018, doi: 10.3390/sym10110648.
- [8] S. Hochreiter, "The Vanishing Gradient Problem During Learning Recurrent Neural Nets and Problem Solutions," *International Journal of Uncertainty, Fuzziness and Knowledge-Based System*, vol. 6, no. 2, pp. 107–116, Oct. 1998.
- [9] A. Hanifa, S. A. Fauzan, M. Hikal, and M. B. Ashfiya, "Perbandingan Metode LSTM dan GRU (RNN) untuk Klasifikasi Berita Palsu Berbahasa Indonesia," 2021. [Online]. Available: <https://covid19.go.id/p/hoax-buster>.
- [10] A. Fitriyani, K. Dalam, M. Nilai, B. Sunda, H. K. Suryadi, and S. Syam, "Peran Keluarga dalam Mengembangkan Nilai Budaya Sunda (Studi Deskriptif terhadap Keluarga Sunda di Komplek Perum Riung Bandung)."
- [11] A. Anggoro, A. B. Osmond, and R. E. Saputra, "Recurrent Neural Network untuk Pengenalan Ucapan pada Bahasa Sunda Dialek Utara," 2018.
- [12] R. Arief, N. A. Iriawan, D. A. Lawi, and E. Copresponder, "Klasifikasi Audio Ucapan

- Emosional Menggunakan Model LSTM,” *Konferensi Nasional Ilmu Komputer (KONIK)*, pp. 524–529, 2021.
- [13] G. Thiodorus, A. Prasetia, L. A. Ardhani, and N. Yudistira, “Klasifikasi Citra Makanan/Non Makanan Menggunakan Metode Transfer Learning dengan Model Residual Network,” *Teknologi*, vol. 11, no. 2, pp. 74–83, Jul. 2021, doi: 10.26594/teknologi.v11i2.2402.
- [14] M. Novela and T. Basaruddin, “Dataset Suara dan Teks Berbahasa Indonesia pada Rekaman Podcast dan Talk Show,” *Jurnal Fasikom*, vol. 11, no. 2, pp. 61–66, 2021.
- [15] H. Santoso and M. Fakhriza, “Perancangan Aplikasi Keamanan File Audio Format WAV (Waveform) Menggunakan Algoritma RSA,” 2018.
- [16] Jubile Enterprise, *Python untuk Programmer Pemula*. Alex Media Komputindo, 2019.
- [17] A. A. Kurniawan and M. Mustikasari, “Implementasi Deep Learning Menggunakan Metode CNN dan LSTM untuk Menentukan Berita Palsu dalam Bahasa Indonesia,” vol. 5, no. 4, pp. 2622–4615, 2020, doi: 10.32493/informatika.v5i4.7760.