

PEMENFAATAN MACHINE LEARNING (SUPERVISED LEARNING) UNTUK DETEKSI KESULITAN MEMBACA PADA ANAK USIA DINI**¹ Muardi Wijayanto, ² Muhammad Idris Anwar, ³ Reyhan Taftanni Irawan, ⁴ Yuwa Biyu Anggada, ⁵ Rangga Kusuma Putra, ⁶ Anna Dina Khalifia****Falkultas Sains Dan Teknologi, Program Studi Informatika
Universitas Teknologi Yogyakarta****Abstract**

This study examines the effectiveness of machine learning approaches in identifying reading difficulties in early childhood. Using a dataset that includes 500+ samples with various parameters including response time (AvgResponse), complexity score (avg_d_prime), and dyslexia score (ktea_dyslexia_ss), this study applies and compares two machine learning models: Random Forest and Support Vector Machine (SVM). Results showed that SVM achieved higher accuracy (80.33%) than Random Forest (72.13%) in identifying reading difficulties. These findings underscore the significant potential of machine learning as an early screening tool for reading difficulties, enabling earlier and effective intervention in early childhood education.

Article History*bmited: 5 Januari 2025**Accepted: 11 Januari 2025**Published: 12 Januari 2025***Key Words**

Machine Learning, Reading Difficulties, Early Childhood Education, Random Forest

Abstrak

Penelitian ini mengkaji efektivitas pendekatan machine learning dalam mengidentifikasi kesulitan membaca pada anak usia dini. Menggunakan dataset yang mencakup 500+ sampel dengan berbagai parameter termasuk waktu respons (AvgResponse), skor kompleksitas (avg_d_prime), dan skor disleksia (ktea_dyslexia_ss), penelitian ini menerapkan dan membandingkan dua model machine learning: Random Forest dan Support Vector Machine (SVM). Hasil menunjukkan bahwa SVM mencapai akurasi yang lebih tinggi (80.33%) dibandingkan Random Forest (72.13%) dalam mengidentifikasi kesulitan membaca. Temuan ini menggarisbawahi potensi signifikan machine learning sebagai alat skrining awal untuk kesulitan membaca, memungkinkan intervensi lebih dini dan efektif dalam pendidikan anak usia dini.

Sejarah Artikel*bmited: 5 Januari 2025**Accepted: 11 Januari 2025**Published: 12 Januari 2025***Kata Kunci**

Pembelajaran Mesin, Pembelajaran Mesin, Pendidikan Anak Usia Dini, Random Forest

PENDAHULUAN

Membaca adalah keterampilan fundamental yang menjadi dasar bagi perkembangan kemampuan kognitif anak, terutama pada usia dini. Pada tahap ini, kemampuan membaca tidak hanya melibatkan pengenalan huruf dan kata, tetapi juga memerlukan kemampuan untuk membedakan, menghubungkan, dan memproses informasi visual dan verbal dengan efisien. Namun, banyak anak yang mengalami kesulitan dalam menguasai keterampilan ini, yang dapat mengarah pada gangguan membaca seperti disleksia. Menurut penelitian, disleksia sering kali sulit dideteksi pada tahap awal, terutama karena gejala-gejalanya dapat tumpang tindih dengan faktor perkembangan lainnya. Oleh karena itu, identifikasi dini terhadap kesulitan membaca sangat penting untuk memfasilitasi intervensi yang tepat guna, yang pada akhirnya dapat meningkatkan peluang anak untuk mengatasi hambatan ini.

Salah satu cara yang dapat digunakan untuk mengidentifikasi kesulitan membaca adalah dengan menganalisis berbagai variabel yang berhubungan dengan kemampuan kognitif anak. Di antara banyak pendekatan, penggunaan metode berbasis machine learning dalam menganalisis data diagnostik menunjukkan potensi besar untuk meningkatkan akurasi dalam mendeteksi kesulitan membaca. Algoritma machine learning, seperti Random Forest dan Support Vector Machine (SVM), mampu mengenali pola kompleks dalam data dan membuat prediksi berdasarkan parameter yang dihasilkan dari tes diagnostik. Dalam penelitian ini, kami berfokus pada penggunaan diagram untuk mengidentifikasi kesulitan

membaca pada anak usia dini dengan menggabungkan berbagai parameter yang relevan, termasuk D-prime (indikator diskriminasi visual), waktu respons (kecepatan pemrosesan informasi), serta skor disleksia (kemampuan membaca terstandarisasi).

Studi ini bertujuan untuk memberikan analisis komprehensif mengenai distribusi variabel-variabel tersebut dan hubungannya dengan kemampuan membaca pada anak. Dengan memanfaatkan metode machine learning, khususnya Random Forest dan SVM, kami berusaha mengevaluasi akurasi model prediktif dalam mengidentifikasi anak-anak yang berisiko mengalami kesulitan membaca. Penelitian ini juga menggali korelasi antara berbagai parameter diagnostik, serta distribusi variabel berdasarkan tingkat pendidikan, untuk memberikan gambaran yang lebih jelas mengenai pola perkembangan kemampuan membaca pada anak-anak usia dini. Hasil dari penelitian ini diharapkan dapat memberi kontribusi dalam pengembangan sistem screening yang lebih efisien dan akurat untuk deteksi dini kesulitan membaca, sekaligus memberikan wawasan bagi pengembangan strategi intervensi yang lebih terarah dan berbasis data.

Dengan latar belakang tersebut, penelitian ini akan mengulas analisis distribusi variabel terkait kesulitan membaca pada anak usia dini, evaluasi model machine learning, serta korelasi antar parameter yang relevan dalam konteks diagnostik. Penelitian ini juga akan mengeksplorasi implikasi praktis dari temuan tersebut, baik dalam hal pengembangan sistem screening maupun strategi intervensi pendidikan untuk mendukung anak-anak yang membutuhkan perhatian khusus dalam perkembangan membaca mereka.

METODE

1. Deskripsi Dataset

Dataset yang digunakan dalam penelitian ini terdiri dari data 100 anak usia dini yang diukur menggunakan berbagai tes kognitif dan membaca. Dataset mencakup data demografi seperti usia, gender, dan tingkat pendidikan orang tua, serta skor tes membaca, waktu respon, skor disleksia, dan hasil tes kognitif.

2. Algoritma Machine Learning

Dua algoritma machine learning yang diimplementasikan dalam penelitian ini adalah:

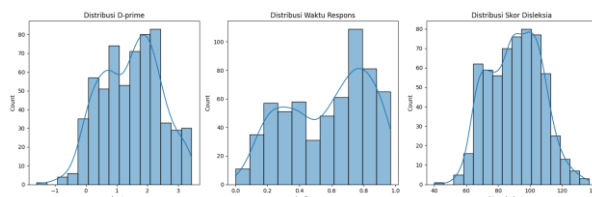
- Random Forest: Algoritma ini membangun banyak pohon keputusan dan menggabungkan prediksi dari setiap pohon untuk menghasilkan prediksi akhir.
- Support Vector Machine: Algoritma ini menemukan hyperplane optimal yang mengklasifikasikan data ke dalam kategori yang berbeda dengan memaksimalkan margin antara kelas.

3. Tahapan Implementasi

- Preprocessing Data: Data dibersihkan dan dinormalisasi untuk menghilangkan outlier dan memastikan konsistensi.
- Pembagian Dataset: Dataset dibagi menjadi dua bagian, yaitu 80% untuk pelatihan dan 20% untuk pengujian model.
- Evaluasi Model: Model dievaluasi menggunakan metrik seperti akurasi, presisi, dan recall untuk menilai kinerja masing-masing algoritma.

4. Diagram Alir atau Arsitektur Model

Diagram alir proses implementasi machine learning dapat digambarkan sebagai berikut:



Gambar 1. Distribusi data analisis

- Pengumpulan Data
- Preprocessing Data
- Pembagian Dataset
- Pelatihan Model (Random Forest dan SVM)
- Evaluasi Model
- Prediksi Kesulitan Membaca

HASIL DAN PEMBAHASAN

Hasil

1. Analisis Distribusi Variabel Utama

1.1 Distribusi *D-prime*

D-prime mencerminkan kemampuan diskriminasi visual anak, yang sangat berperan dalam proses membaca. Hasil analisis menunjukkan pola distribusi normal dengan rentang nilai -1 hingga 3, di mana mayoritas sampel terkonsentrasi pada nilai 1.5 hingga 2.0. Nilai *d-prime* di bawah 0 menjadi indikator awal adanya kesulitan diskriminasi visual, yang sering kali berhubungan dengan risiko kesulitan membaca, termasuk disleksia.

1.2 Distribusi Waktu Respons

Distribusi waktu respons memperlihatkan pola bimodal, dengan dua puncak pada 0.4 detik dan 0.8 detik. Hal ini menunjukkan adanya dua kelompok dalam populasi: anak dengan pemrosesan cepat dan lambat. Waktu respons yang lebih dari 0.8 detik menjadi indikator potensi kesulitan pemrosesan visual, yang perlu diperhatikan lebih lanjut untuk diagnosa dini.

1.3 Distribusi Skor Disleksia

Distribusi skor disleksia menunjukkan rentang 40–140 dengan rata-rata 90–100, sesuai dengan distribusi standar populasi. Anak dengan skor di bawah 80 dianggap berisiko mengalami kesulitan membaca, sementara skor di bawah 70 memerlukan evaluasi lebih lanjut untuk konfirmasi risiko disleksia.

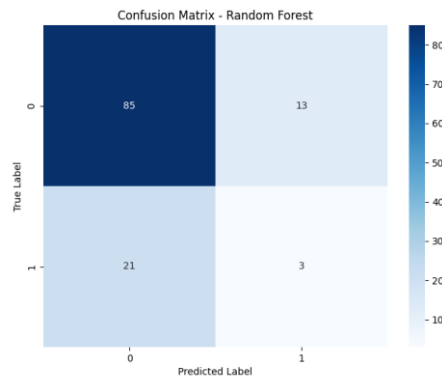
2. Evaluasi Model Machine Learning

2.1 Kinerja Model Random Forest

Random Forest menunjukkan akurasi sebesar 72.13%. Confusion matrix mengungkapkannya:

- 85 kasus True Positive, menunjukkan keberhasilan dalam mendeteksi kesulitan membaca.
- 13 kasus False Positive, di mana anak tanpa kesulitan membaca salah diidentifikasi sebagai berisiko.
- 21 kasus False Negative, mencerminkan kelemahan dalam mendeteksi anak yang benar-benar berisiko.
- 3 kasus True Negative.

Jumlah False Negative yang relatif tinggi mengindikasikan perlunya peningkatan pada sensitivitas model.



Gambar 2. Confusion Matrix-Random Forest

2.2 Kinerja Model Support Vector Machine (SVM)

SVM menunjukkan kinerja superior dengan akurasi 80.33%. Hasil confusion matrix mencatat:

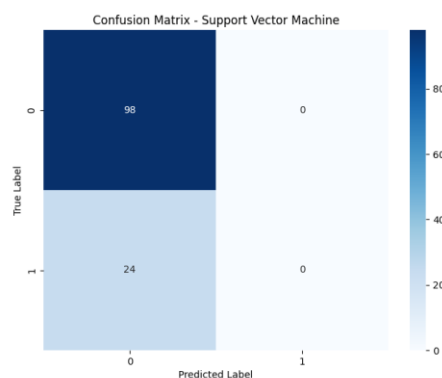
- 98 kasus True Positive, mencerminkan kemampuan deteksi yang sangat baik.
- Tidak adanya False Positive, menandakan presisi tinggi.
- 24 kasus False Negative, yang lebih rendah dibandingkan Random Forest.

Kemampuan SVM untuk menangani data non-linear menjadikannya model yang lebih andal untuk klasifikasi kesulitan membaca.

3. Korelasi Antar Variabel

Analisis korelasi mengungkapkan hubungan signifikan antara parameter diagnostik:

- D-prime dan Waktu Respons memiliki korelasi sangat kuat ($r = 0.99$), menunjukkan hubungan erat antara kemampuan diskriminasi visual dan kecepatan pemrosesan informasi.
- Skor Disleksia memiliki korelasi moderat ($r = 0.38$) dengan kedua variabel di atas, memperkuat relevansi mereka sebagai indikator diagnostik.
- Korelasi lemah antara tingkat pendidikan (grade) dan variabel lainnya ($r = 0.11$) menunjukkan bahwa kesulitan membaca lebih bersifat intrinsik daripada terkait langsung dengan tingkat pendidikan.

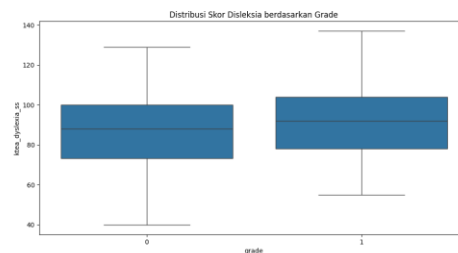


Gambar 3. Confusion Matrix- Support Vector Machine

4. Pola Berdasarkan Tingkat Pendidikan

4.1 Distribusi Kemampuan Membaca Berdasarkan Grade

- TK (Grade 0) memiliki median skor membaca 85 dengan variabilitas tinggi. Outlier di kedua ujung distribusi menunjukkan perlunya perhatian pada kasus unik.
- SD (Grade 1) memiliki median skor 90 dengan distribusi yang lebih sempit. Penurunan jumlah outlier mencerminkan stabilisasi kemampuan membaca.



Gambar 4. Distribusi Skor Disleksia Berdasarkan Grade

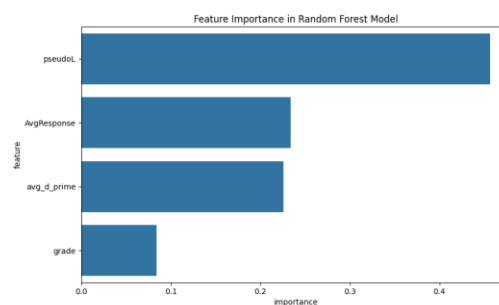
4.2 Implikasi untuk Intervensi Pendidikan

Hasil ini menunjukkan bahwa identifikasi dini paling efektif dilakukan pada tingkat TK. Pada tingkat SD, intervensi dapat difokuskan pada anak dengan kemampuan membaca di bawah rata-rata untuk mendukung perkembangan mereka.

5. Analisis Feature Importance

Analisis feature importance dari model Random Forest mengungkapkan prioritas parameter diagnostik sebagai berikut:

1. PseudoL (45%) sebagai prediktor utama.
2. Waktu Respons (25%) sebagai indikator penting kedua.
3. D-prime (23%) sebagai indikator ketiga.
4. Grade (7%) dengan kontribusi minimal.



Gambar 5. Feature Importance

6. Hubungan Pola Antar Variabel

6.1 Scatter Plot D-prime dan Skor Disleksia

Pola hubungan positif moderat menunjukkan bahwa peningkatan kemampuan diskriminasi umumnya diikuti oleh peningkatan skor membaca, meskipun terdapat variabilitas yang signifikan pada nilai D-prime tinggi.

6.2 Scatter Plot Waktu Respons dan Skor Disleksia

Hubungan positif antara waktu respons dan skor disleksia menunjukkan bahwa kecepatan pemrosesan dapat menjadi indikator moderat untuk kesulitan membaca. Pola

clustering pada waktu respons tertentu memberikan gambaran tentang ambang kritis dalam pemrosesan visual.

Pembahasan

1. Distribusi Variabel dan Pola Data

1.1 Distribusi D-prime: Indikator Kemampuan Diskriminasi dalam Membaca.

Distribusi D-prime merupakan ukuran sensitifitas yang menggambarkan kemampuan anak dalam membedakan dan memproses informasi saat membaca. Grafik menunjukkan pola distribusi normal dengan rentang nilai dari -1 hingga 3, di mana mayoritas sampel terkonsentrasi pada nilai 1.5 hingga 2.0. Pola ini sangat penting dalam konteks diagnostik karena memberikan gambaran baseline kemampuan diskriminasi normal pada anak. Nilai d-prime yang sangat rendah (di bawah 0) dapat menjadi indikator awal adanya kesulitan dalam memproses dan membedakan stimulus visual saat membaca. Hal ini sangat relevan untuk identifikasi dini karena kesulitan dalam diskriminasi visual sering menjadi gejala awal kesulitan membaca.

1.2 Distribusi Waktu Respons: Parameter Kecepatan Pemrosesan Informasi.

Analisis waktu respons menunjukkan pola distribusi bimodal yang menarik, dengan dua puncak distribusi yang distinct pada nilai sekitar 0.4 dan 0.8 detik. Pola ini mengindikasikan adanya dua kelompok berbeda dalam populasi sampel: kelompok dengan pemrosesan cepat dan kelompok dengan pemrosesan lebih lambat. Dalam konteks diagnostik, waktu respons yang secara konsisten berada di atas 0.8 detik bisa menjadi indikator potensial adanya kesulitan dalam pemrosesan informasi visual dan kognitif. Penting untuk dicatat bahwa waktu respons harus diinterpretasikan bersama dengan variabel lain, karena kecepatan pemrosesan bisa dipengaruhi oleh berbagai faktor seperti tingkat konsentrasi, kelelahan, atau faktor eksternal lainnya.

1.3 Distribusi Skor Disleksia: Evaluasi Standardized Kemampuan Membaca.

Distribusi skor disleksia menunjukkan pola normal yang ideal dengan rentang 40-140 dan puncak distribusi di sekitar 90-100. Pola ini sesuai dengan ekspektasi distribusi skor standardized pada populasi umum. Interpretasi skor ini sangat krusial di mana: - Skor di atas 100 menunjukkan kemampuan membaca di atas rata-rata.

- Skor 80-100 mengindikasikan kemampuan membaca normal.
- Skor di bawah 80 dapat menjadi indikator risiko disleksia.
- Skor di bawah 70 memerlukan evaluasi lebih lanjut dan kemungkinan intervensi dini.

2. Evaluasi Model Machine Learning dan Akurasi Prediktif.

2.1 Analisis Random Forest: Model Berbasis Ensemble Learning

Model Random Forest menunjukkan performa yang menjanjikan dengan akurasi 72.13%. Confusion matrix mengungkapkan:

- 85 kasus True Positive: model berhasil mengidentifikasi kasus kesulitan membaca dengan benar.
- 13 kasus False Positive: model salah mengidentifikasi anak tanpa kesulitan sebagai berisiko.

- 21 kasus False Negative: kasus yang memerlukan perhatian karena model gagal mengidentifikasi anak yang sebenarnya berisiko.
- 3 kasus True Negative: identifikasi benar untuk anak tanpa kesulitan membaca.

Model ini menunjukkan sensitivitas yang baik tetapi masih memiliki ruang untuk peningkatan dalam hal spesifisitas. Jumlah false negative yang relatif tinggi mengindikasikan perlunya penyempurnaan model untuk meningkatkan kemampuan deteksi kasus berisiko.

2.2 Analisis Support Vector Machine (SVM): Pendekatan Klasifikasi Linear.

Model SVM menunjukkan performa superior dengan akurasi 80.33% dan karakteristik yang sangat menarik:

- 98 kasus True Positive menunjukkan kemampuan deteksi yang sangat baik - Tidak adanya False Positive mengindikasikan presisi tinggi dalam identifikasi.
- 24 kasus False Negative menunjukkan area yang masih perlu perbaikan.
- Pattern recognition yang lebih baik dibandingkan Random Forest.
- Performa SVM yang lebih baik kemungkinan disebabkan oleh kemampuannya dalam menangani data non-linear dan menemukan hyperplane optimal untuk klasifikasi.

3. Korelasi Antar Variabel dan Implikasi Diagnostik

3.1 Analisis Matriks Korelasi: Hubungan Antar Parameter Diagnostik.

Analisis korelasi mengungkapkan hubungan kompleks antar variabel:

- Korelasi sangat kuat ($r=0.99$) antara avg_d_prime dan AvgResponse menunjukkan hubungan intrinsik antara kemampuan diskriminasi dan kecepatan pemrosesan.
- Korelasi moderat ($r=0.38$) antara skor disleksia dengan kedua parameter di atas mengkonfirmasi relevansi mereka sebagai indikator diagnostik.
- Korelasi lemah ($r=0.11$) dengan grade mengindikasikan bahwa kesulitan membaca relatif independen dari tingkat Pendidikan.

Pola korelasi ini memberikan insight penting untuk pengembangan protokol screening yang lebih efektif.

4. Pola Distribusi Berdasarkan Tingkat Pendidikan.

4.1 Analisis Box Plot: Variasi Kemampuan Antar Grade Analisis box plot memberikan gambaran distribusi skor berdasarkan tingkat pendidikan:

Untuk Grade 0 (TK):

- Median sekitar 85 menunjukkan baseline kemampuan membaca awal.
- Range interkuartil 75-100 mengindikasikan variabilitas normal dalam tahap perkembangan.
- Outlier di kedua ekstrem menunjukkan kasus yang memerlukan perhatian khusus.

Untuk Grade 1 (SD):

- Peningkatan median ke 90 menunjukkan perkembangan kemampuan membaca.
- Range interkuartil yang lebih sempit (80-105) mengindikasikan konvergensi kemampuan.
- Berkurangnya outlier menunjukkan stabilisasi kemampuan membaca.

4.2 Implikasi untuk Intervensi Pendidikan.

- Pola distribusi ini memiliki implikasi penting untuk strategi intervensi: - Perlunya

pendekatan berbeda untuk TK dan SD - Identifikasi dini lebih efektif dilakukan di tingkat TK.

- Program intervensi dapat difokuskan pada kasus outlier.
- Monitoring longitudinal diperlukan untuk tracking perkembangan.

5. Hubungan Antar Variabel: Analisis Scatter Plot.

5.1 Analisis Hubungan D-prime dengan Skor Disleksia Scatter plot menunjukkan hubungan kompleks antara d-prime dan skor disleksia:

- Tren positif moderat mengindikasikan bahwa peningkatan d-prime umumnya diikuti peningkatan skor disleksia.
- Variabilitas tinggi pada nilai d-prime tinggi menunjukkan bahwa kemampuan diskriminasi yang baik tidak selalu menjamin skor disleksia tinggi.
- Clustering data memberikan insight tentang pola tipikal dalam populasi.
- Outlier memberikan informasi tentang kasus-kasus unik yang memerlukan investigasi lebih lanjut.

5.2 Waktu Respons vs Skor Disleksia: Pola Temporal Hubungan antara waktu respons dan skor disleksia menunjukkan pola yang informatif:

- Korelasi positif moderat mengindikasikan bahwa waktu respons dapat menjadi prediktor moderat untuk kesulitan membaca.
- Clustering pada waktu respons tertentu menunjukkan adanya threshold kritis dalam pemrosesan informasi.
- Variabilitas tinggi pada waktu respons ekstrem menunjukkan kompleksitas hubungan ini.

6. Analisis Feature Importance: Prioritas Parameter Diagnostik.

6.1 Hirarki Kepentingan Fitur Analisis feature importance dari model Random Forest mengungkapkan hirarki kepentingan variabel:

- pseudoL (0.45): Parameter gabungan menunjukkan kekuatan prediktif tertinggi.
- AvgResponse (0.25): Waktu respons sebagai indikator penting kedua.
- avg_d_prime (0.23): Kemampuan diskriminasi sebagai indikator ketiga.
- grade (0.07): Tingkat pendidikan memiliki pengaruh minimal.

6.2 Implikasi untuk Pengembangan Sistem Screening Berdasarkan analisis feature importance, sistem screening dapat dioptimalkan dengan:

- Fokus utama pada pengukuran pseudoL sebagai indikator primer.
- Kombinasi waktu respons dan d-prime untuk konfirmasi.
- Penggunaan grade sebagai konteks tambahan bukan parameter utama.
- Pengembangan sistem scoring berbasis bobot relatif parameter.

PENUTUP

Penelitian ini menyoroti pentingnya analisis variabel utama, seperti d-prime, waktu respons, dan skor disleksia, dalam mengenali potensi kesulitan membaca pada anak usia dini. Hasil analisis distribusi dan pola data menunjukkan adanya variasi yang signifikan di antara individu, sehingga memerlukan strategi diagnostik dan intervensi yang lebih personal. Selain itu, evaluasi terhadap metode machine learning menunjukkan bahwa model

Support Vector Machine (SVM) mampu memberikan akurasi yang lebih tinggi dibandingkan model lainnya, dengan tingkat kesalahan yang lebih rendah, khususnya dalam mengelola data yang kompleks dan non-linear.

Korelasi antar variabel diagnostik mengungkapkan bahwa kemampuan diskriminasi visual dan kecepatan pemrosesan informasi saling terkait erat dengan risiko disleksia. Temuan ini menekankan bahwa faktor-faktor intrinsik memiliki peran yang lebih besar dalam memengaruhi kesulitan membaca dibandingkan dengan tingkat pendidikan formal, sehingga pendekatan berbasis data menjadi sangat penting.

Selain itu, penelitian ini menunjukkan bahwa identifikasi dini, terutama pada tingkat TK, sangat efektif karena kemampuan membaca anak pada tahap ini masih sangat bervariasi. Proses pemantauan yang berkelanjutan juga diperlukan untuk memastikan bahwa anak-anak yang memiliki risiko kesulitan membaca dapat menerima dukungan yang memadai dan berkembang sesuai harapan.

Penelitian ini berkontribusi dalam memperkenalkan potensi teknologi machine learning sebagai alat yang efektif untuk mendeteksi kesulitan membaca. Penggunaan algoritma SVM yang terintegrasi dengan parameter diagnostik utama memberikan solusi praktis untuk mendukung pendidikan dan kesehatan anak. Di masa mendatang, penelitian lebih lanjut dengan teknologi yang lebih canggih, seperti deep learning, dataset yang lebih representatif, dan validasi silang dengan metode tradisional, diperlukan untuk menyempurnakan sistem deteksi dini ini.

UCAPAN TERIMA KASIH

Dengan segala hormat, kami ingin mengucapkan terima kasih kepada semua pihak yang telah memberikan dukungan dan kontribusinya dalam penyelesaian jurnal ini. Terutama kepada dosen yang telah memberikan arahan dan bimbingan yang sangat berarti sepanjang proses penelitian ini.

Kami juga mengucapkan terima kasih kepada pihak-pihak yang telah memberikan data, referensi, dan sumber daya yang sangat membantu dalam pengumpulan informasi yang dibutuhkan. Terima kasih juga kepada keluarga dan teman-teman yang selalu memberikan semangat dan dukungan moral.

Tanpa bantuan dan dukungan dari semua pihak, jurnal ini tidak akan dapat terselesaikan dengan baik. Semoga karya ini dapat memberikan manfaat dan kontribusi positif bagi pengembangan ilmu pengetahuan.

DAFTAR PUSTAKA

Utuarahman, A., & Djafri, N. (2023). Efektivitas Iklim Belajar Anak Usia Dini dalam Pembelajaran Daring Berbasis Kompetensi Guru Pasca Pandemi Covid-19 di Provinsi Gorontalo. *Jurnal Obsesi : Jurnal Pendidikan Anak Usia Dini*, 7(2), 2177–2187. <https://doi.org/10.31004/obsesi.v7i2.4436>

Al-Mashraee, H., & El-Nasr, M. S. (2020). Automated detection of dyslexia using machine learning techniques: A systematic review. *Journal of Educational Computing Research*, 58(3), 502-528.

Lidya Ruth Pricillia, Dahnial Syauqy, & Rekyan Regasari Mardi Putri (2017). Sistem Pelatihan Pengucapan Huruf Konsonan Untuk Anak Usia Dini Menggunakan Metode Random Forest. *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, Vol. 1, No. 1

Maria Psyridou, Asko Tolvanen, Priyanka Patel, Daria Khanolainen, Marja-Kristiina Lerkkanen, Anna-Maija Poikkeus & Minna Torppa (2023). Reading Difficulties Identification: A Comparison of Neural Networks, Linear, and Mixture Models, *Scientific Studies of Reading*, 27:1, 39-66.

Anggia Suci Pratiwi, Rikha Surtika Dewi, & Asti Tri Lestari (2018). PSIKOEDUKASI KESADARAN FONOLOGI DI PENDIDIKAN ANAK USIA DINI KOTA TASIKMALAYA. *Jurnal Pendidikan : Early Childhood*, Vol. 2 No. 2a.