

発達性ディスレクシア児のための漢字学習支援システムの開発と評価

梶谷 星^{†1} 川合 康央^{†1}

概要：本研究は、発達性ディスレクシアを持つ子どもたちのための学習支援システムの開発に関するものです。ディスレクシアは読み書きに困難をもたらす学習障害であり、学業成績や自尊心に深刻な影響を及ぼす可能性があり、過去の研究や事例研究を行い、特に聴覚法を用いた漢字学習の効果に注目しました。これらの知見に基づき、小学校 1～3 年生で学習する 440 文字の漢字に対応する学習支援システムを開発しました。システム開発にあたり、音声合成、音声認識、画像認識といった技術の検証を行い、可読性に影響を与えるフォントの効果についても調査しました。本研究では、ディスレクシアを持つ子どもたちの学習を支援する効果的な ICT ベースのアプローチの可能性を示唆しています。今後は、実際のディスレクシア児を対象とした評価実験を通じてシステムの有効性を検証する予定です。

1. はじめに

ディスレクシアは学習障害の一種であり、特に文字の読み書きに困難を抱える点が特徴である。この障害は、国語などの言語に関する教科に限らず、算数や理科といった幅広い分野に影響を与える。中でも発達性ディスレクシアは、児童の成長過程において早期に現れることが多く、学業全般における多様な問題を引き起こす可能性がある。当事者は、読み書きの困難さに直面することで、学業成績が伸び悩むだけでなく、自己評価の低下や学校生活における孤立感を抱くことがある。このような心理的および社会的な影響は、学業生活にとどまらず、将来の進路や職業選択にも悪影響を及ぼす可能性がある。本研究は、このような課題を克服するため、発達性ディスレクシアを抱える当事者を支援する学習支援アプリケーションの開発を目的とする。この取り組みにより、当事者が直面する学習の障壁を低減し、自信の向上および学業生活の質的改善を目指す。

2. 研究の背景と目的

2.1 ディスレクシアの概要と学校教育

ディスレクシアは、「発達性ディスレクシア」と「後天性ディスレクシア」の 2 種類に分類される。発達性ディスレクシアは、脳損傷の既往がないにもかかわらず生じるものであり、読みの問題に加え、書字にも困難を伴う。一方、後天性ディスレクシアは後天的な脳損傷によって発生し、主に読みのみに障害が現れる[1]。

文部科学省の調査（令和 4 年）によれば、「読む」または「書く」に顕著な困難を示す児童生徒の割合は 3.5%に達する[2]。この調査は令和 4 年 1 月から 2 月にかけて全国の公立学校に在籍する約 74,919 名を対象に実施された。

現行の学校教育システムにおけるディスレクシア児への対応には、多くの課題が存在する。第一に、診断の問題が挙げられる。発達性ディスレクシアの正確な評価には、統一された診断テストが必要であるが、全国規模での導入は

限定的である。この状況は早期発見と適切な支援開始を妨げる一因となっている。第二に、教材不足が課題である。ディスレクシア児の特性に配慮した教材、特に漢字学習を重視したドリルはほとんど存在していない。第三に、指導法の適切性が挙げられる。我が国では従来型のトップダウン式指導法が多用されているが、これはディスレクシア児の学習特性に適合していない。最後に、支援アプローチの問題がある。現行の支援体制では、本人の要望がある場合にのみサポートが提供される状況にとどまっている。ディスレクシア児に対する教育支援の改善には、具体的な解決策の開発と実装が求められる。現在の支援方法としては、通常の学習方法に加えて音声入力やキーボード入力などの ICT 技術を利用した学習支援が導入されている[3]。しかし、これらの多くは読み書きを可能にするのではなく、代替手段の提供にとどまっている。

2.2 ディスレクシアの概要と学校教育

ディスレクシア児への支援方法は、言語圏によって異なるアプローチが採用されている。

アルファベット圏では、視覚的な工夫を中心に支援が行われている。具体例として、フォントの選定では太さが均一で文字間の空白が広い書体が用いられ、文字の認識を容易にしている。また、レイアウトの工夫として、両端揃えを避けることで行の移行を分かりやすくし、段分けや行間隔の拡大により文字の視認性を向上させている。さらに、視覚的補助として、下線の代わりに語全体を色で覆う方法や、各行に番号を振ることで読み返しを減らす方法が採用されている。これらの取り組みは視覚的情報処理の負担を軽減し、読解効率を高めることを目的としている。

一方、漢字圏では、文字の複雑さを考慮した独自のアプローチが必要とされる。特にディスレクシア児は、視覚的同時処理障害（視覚刺激を並列処理する能力）を抱える可能性が高い。この障害は漢字の認識に直接的な影響を与え

^{†1} 文教大学大学院 情報学研究科

るため、ルビ振りが有効な支援方法とされている。しかし、ルビ振りを煩雑に感じて利用を避けるケースも多く、効果的な活用方法の模索が求められている。アルファベット圏と漢字圏の支援方法を比較すると、視覚的情報処理の負担軽減が共通の目標となっている。しかし、文字体系の違いにより具体的なアプローチには相違がある。アルファベット圏ではテキスト提示方法に重点が置かれるのに対し、漢字圏では文字認識と意味理解の支援が重要視される。特に、漢字圏におけるルビ振りは音読と意味理解を同時に支援する有力な手法であり、その効果的な実施方法が課題となっている。

3. 先行事例

ディスレクシア児への漢字学習支援アプリケーションの開発に先立ち、既存の支援方法や関連研究の事例について調査を行った。

3.1 支援方法

まず、DAISY というマルチメディア教科書の活用が挙げられる[4]。DAISY は音声や映像を用いた教科書を提供し、「分かち書き」による単語や文節の区切りを導入することで、読みやすさの向上を図っている。この手法はディスレクシア児の読解を支援し、学習効率を向上させることができる。支援方法の有用性については、2 段階方式の音読指導が効果的であることが確認されている[5]。具体的には、文字の解読を容易にする解読指導が誤読数の減少に寄与し、単語形態の認識を高める語彙指導が音読時間の短縮に効果を示している。さらに、直感的なデザインの活用も有効であるとされている。音楽教育の分野においては、フィギュアノートという色と形を用いた直感的なデザインの楽譜がある[6]。この手法では、楽器と楽譜に同じ色と形のシール

を貼り、五線譜の理解が困難な生徒でも楽器演奏に参加できるようにした事例がある。これにより、直感的なデザインが学習補助において有効であることが示唆されている。

3.2 関連研究

関連する研究事例として、『かんじダス』というアプリケーションがある[7]。このアプリケーションでは、漢字をイラストにドラッグして意味を合わせる「絵合わせ」形式を採用しており、Web アプリケーションとしてオンライン・オフラインの両方で利用可能である。24 種類の問題が収録され、1 回のゲームで 4 組 1 セットのランダム出題形式を採用している。Web アプリケーションの特性上、多端末対応や管理の容易性が特徴であり、学習支援ツールとして有用である。また、書字学習に困難を抱える生徒を対象とした ICT 活用型自己学習支援システムの開発事例も挙げられる[8]。このシステムは iPad をハードウェアとして使用し、FileMaker を用いて教材を作成したものである。学習後にテストを実施し、読みと意味を確認、さらに「答えを確認する」ボタンで正解を表示する仕組みを導入している。このアプローチは、失敗を最小限に抑えながら成功体験を積み重ねる学習方法を提案しており、ICT を活用した学習支援の有効性を示している。さらに、発達性読み書き障害を持つ児童生徒を対象とした漢字書字訓練の研究がある[9]。この研究では、聴覚法と視覚法の 2 種類の漢字書字訓練が比較されている。聴覚法では、漢字の構成要素を文章化して覚える方法が用いられている。例えば、「給」という漢字を「給食の時間に糸を合わせる」と覚える方法である。この研究の結果、14 名中 12 名で聴覚法が視覚法より有効であることが示され、聴覚法の有用性が示唆された。

これらの先行事例の分析から、効果的な支援には以下の

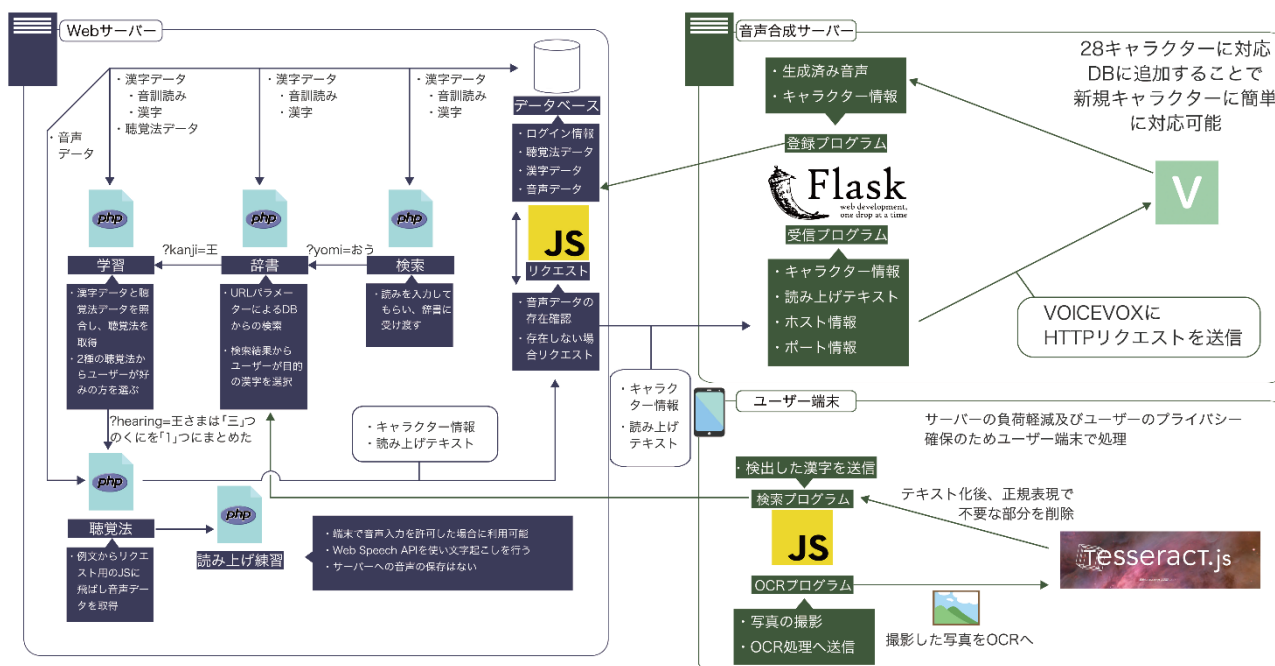


図 1 システム構成図

要素が重要であることが明らかになった。ICT 技術を活用したマルチメディアの利用、視覚的および聴覚的アプローチの組み合わせ、直感的なデザインと操作性、個別学習と即時フィードバックの提供、そして聴覚法を用いた記憶支援である。本研究では、これらの知見を基盤とし、より効果的な漢字学習支援アプリケーションの開発を目指す。

4. システム開発

4.1 環境構築

本研究では、先行研究で効果が示された「聴覚法」を採用し、ディスレクシア児の読みの学習を支援する Web システムを開発した。本システムの対象とする漢字は、教育漢字 1,026 字のうち、小学 1 年生から小学 3 年生までに学習する 440 字に設定した。この選定理由は、小学校で学ぶ漢字 1,026 字の 90%、さらに中学校で学ぶ漢字 1,110 字の 72% がこの 440 字の組み合わせで構成されているためである。

システムの開発には、2 台のサーバーを使用した分散型アーキテクチャを採用した。Web サーバー（CPU: Core i3-4130, RAM: 16GB）では、Web サイト公開用に Apache2 を使用し、MariaDB を用いて聴覚法の例文やユーザーの学習情報、生成した音声データを格納している。一方、音声合成サーバー（CPU: Ryzen 7-1700, RAM: 32GB, GPU: GTX1060）では、音声合成エンジンとして VOICEVOX を採用している。さらに、Web サーバーからのリクエストを受け取り、VOICEVOX に指示を出すための Flask が稼働している。この 2 台のサーバーは連携して動作し、ユーザーに対してシームレスな学習体験を提供する。システム全体の構成図は図 1 に示す通りである。

本システムの特徴として、高い拡張性が挙げられる。URL パラメーターを活用することで、データベースに新しい漢字、その音訓読み、学年情報、聴覚法の例文を追加することで容易に新しい漢字を追加できる。この機能により、システムの対象とする漢字や学年範囲を将来的に拡張する際の柔軟性が確保されている。

4.2 技術検証

システム開発に際して、主要な要素に関する技術検証を実施した。特に音声合成、音声入力、画像認識、LLM (Large Language Model) の 4 つの技術に焦点を当て、それぞれの性能と適合性を評価した。

音声合成技術として、本システムでは VOICEVOX を採用した。VOICEVOX は、エディタ、エンジン、コアの 3 つのコンポーネントから構成される Web サーバー型のソフトウェアである。[10]標準設定では、50021 番ポートへの HTTP リクエストにより音声を生成する。HTTP リクエストに必要な情報は、ホスト IP、ポート番号、読み上げテキスト、およびキャラクターID である。キャラクターID を指定することで、多様な音声キャラクターを選択可能である。

音声生成プロセスは 2 段階で構成される。最初に、音声

生成に必要な JSON データを取得し、次にその JSON データを使用して実際の音声を生成する。この方法により、音声合成に要する時間は約 2 秒と短く抑えられており、さらに辞書登録機能が充実している点から、VOICEVOX の採用を決定した。

次に、音声入力には OpenAI が開発したオープンソースの音声認識ソフトウェアである Whisper を検討した。Whisper を動作させるために、Git, Python, PyTorch, FFmpeg を用いた環境を構築し、venv モジュールで軽量な仮想環境を作成した。検証には、Windows 10 Pro, Ryzen 7 1700 (3.7GHz OC), RTX 4070 (VRAM 12GB), DDR4 32GB のシステムを使用した。ファイルサイズ 1.2MB, 長さ 3 秒の WAV 形式ファイル 3 種類を用いてテストを行った結果、base モデル以上であれば十分な認識精度が得られることが確認された (表 1)。一方で、複数リクエストを同時処理する際に発生する時間的コストが課題となった。そのため、通常の使用においては Speech Recognition の高精度版を採用し、Whisper は特定の条件下でのみ利用する方針とした。

表 1 Whisper 検証結果

モデルサイズ	要求 VRAM	時間	精度
tiny	1GB	8s	×
base	1GB	9s	○
small	2GB	11s	○
medium	5GB	17s	○
large	10GB	34s	○

画像認識技術としては、Google が提供する Tesseract OCR を採用した。Tesseract は Apache License 2.0 のもとで公開されており、無償かつ無制限で商用利用が可能である。

しかし、漢字ドリルなどで一般的に使用される黄色い背景に黒字で表記された漢字の認識に課題があることが判明した。この問題に対応するため、ページセグメンテーションモード (PSM) の調整を行った。その結果、PSM5 または PSM6 を使用することで、正確な文字認識が可能であることが確認された。

LLM の検証は 2 種類実施した。1 つ目は、LLM を活用したルビ振りの正確性の検証であり、2 つ目は LLM を活用した聴覚法の生成である。

4.2.1 LLM を活用したルビ振り

LLM を活用したルビ振りの正確性検証では、以下のようなプロンプトを用いた。

=プロンプト=

あなたは文章中の漢字に正しい読みを振る AI です。こちらから日本語の文章が入力として与えられたら、以下の形式で必ずすべての語に読み仮名 (ルビ) を埋め込んで変換してください。返事は不要ですので入力に対する出力のみ行うようにしてください。

出力形式 (入力: 吾輩は猫である)
出力: <!吾輩=わがはい>は<!猫=ねこ>である

注意点

文脈を考慮し、不自然な読み方にならないようにしてください。
(例: 「生きる」を「<!生=なま>きる」としない)

活用語尾をルビに含めないようにしてください。
(例: 「生きる」を「<!生きる=いきる>」とせず「<!生=い>」き

とすること)
出力は入力文を漢字単位で以下の書式「<!漢字=読み>」に置き換え、ひらがなの部分は変えずに残してください。
このプロンプトで、<!吾輩=わがはい>のように短い形式で出力させ、HTML 形式のルビ (例: <ruby>吾輩<rp></rp><rt>わがはい</rt><rp></rp></ruby>) に比べて使用されるトークン数を削減した。また、HTML 以外の言語で処理する際にも簡潔なコード変換を可能とした。検証には ChatGPT o1 を用い、10 件の例文を生成した。例文には「読み方が変わる語」を複数含ませ、次のような文章を使用した。(例文:実家の近くの坂を上り、高台の道の上にある店へ向かいながら、地元の名産品が生のまま販売されていることや、一部は加工されて生きがよくなる工夫がされていることを思い出していた。) 生成された出力について、<!単語=ルビ>の 1 セットを 1 件とみなし、誤ったルビ振りが行われている件数をカウントした。その結果を表 2 に示す。

表 2 LLM を用いたルビ振り検証結果

モデル	誤認識件数
GPT4o_mini	54/143
GPT4o	00/143
Gemini_1.5_Flash	42/143
Gemini_2.0_Flash	09/143
Grok2	06/143

GPT4o_mini は 20 字程度の短文では正確なルビ振りが可能であったが、文が長くなると変換を行わなかったり、「高台」を「こうだい」と出力するなど、文脈を無視した結果が増加した。Gemini_1.5_Flash では、<!冷たい=つめたい>のように「活用語尾をルビに含めないように」との指示を無視した出力が大半を占めており、Gemini_2.0_Flash でも同様の傾向が見られた。Grok2 は、「名産品」を<!名産=めいさん><!品=しな>と細かく区切りすぎる特徴があった。これらの結果から、最終的に、ルビ振り性能が最も高かったのは GPT4o であり、このモデルを本システムで採用することとした。

4.2.2 LLM を活用した聴覚法の生成

LLM を活用した聴覚法生成実験では、以下のようなプロンプトを用いた。

=プロンプト=

あなたは「聴覚法」を用いて漢字を分解し、短い例文を作成する AI です。以下の手順に従って出力してください。

【手順】

1.漢字入力：一文字の漢字を入力。

2.分解：

- 漢字を最小単位まで分解（ひらがな、カタカナ、漢字、絵文字可）。

- 構成要素を簡条書きで示す。

3.例文作成：

- 分解した要素全てと入力漢字を含む自然な日本語の短文を作成。

【出力形式】

1.構成要素：- 列挙。

2.例文：- 漢字と全要素を含む文。

【注意点】

- できる限り細かく分解。

- 例文に全ての要素と漢字を必ず含む。

- 柔軟に他の文字種や絵文字を使用。

【例】

- 入力：給

- 構成要素：糸・合

- 例文：給食の時間に糸を合わせる

聴覚法生成の検証では、先のルビ振り実験結果から、GPT4o の日本語読解能力の高さを評価し、GPT4o を利用して実験を行った。検証対象として、以下の文字を選定した。

「明」「晴」「銅」「線」各文字の構成要素を基に生成された例文の一部を以下に示す。

「明」構成要素：日、月

1. 明日は日が昇る前に月を観察しよう。
2. 夜明け前の空には日と月が並ぶことがある。
3. 彼の明るい笑顔が日と月のように輝いていた。

「晴」構成要素：日、月、青

1. 晴れた空には日と月が青く輝いている。
2. 天気が晴れると日差しが青空を包み込む。
3. 朝晴れると、日が昇るにつれて青が濃くなる。

「銅」構成要素：金、同、口、一

1. 銅の輝きは、金属の中でも同じ口の形をした装飾品で一際目立つ。
2. 職人は銅の表面を磨き、金属の同じ口模様を一行に並べた。
3. 銅は金属として、同じ口型の部品を一つ一つ作るのに適している。

「線」構成要素：糸、水（氵）、一、川

1. 線路のそばには糸のような一筋の川が流れている。
2. 線香の煙が糸のように水面の川へと漂っていった。
3. 線の先には川があり、水が糸のように一つの流れを作っている。

「明」「晴」に関しては聴覚法として十分に活用可能なレベルの例文が生成された。しかし、「銅」「線」では分解しすぎてしまっていることで、聴覚法として分かりづらい例文が多く生成された。このため、分解のレベルを下げることで実用的な例文が得られる可能性が示唆された。

また、プロンプトを改良するか、こちらで構成要素を指定することで、十分に実用可能な例文を作成できることが分かった。さらに、より高性能なモデルでの検証を行うことで、より優れた聴覚法を生成できる可能性がある。

4.3 書体の検証

先行研究では、アルファベット圏において「太さが均一で文字中の隙間が大きい字体を使用する」ことが読解力向上に効果的であると示唆されている。この知見を基に、日本語フォントにおいても同様の効果が見られるかを検証するため、フォントの違いが読解力や読みやすさに与える影響を調査した。比較対象としたフォントは、「HG 教科書体」「UD デジタル教科書体 N-R Regular」「ヒラギノ角ゴ ProN W3」の 3 種類である。それぞれのフォントの特徴は以下の通りである。

HG 教科書体は、小学校学習指導要領の学年別漢字配当表に基づき作られた教科書体である。手書き文字に近いデザインであり、小学校児童が書写する際の手本となるよう工夫されている。Office 製品に同梱される伝統的な教科書体として選出した。UD デジタル教科書体 N-R Regular は、ユニバーサルデザインを採用した教科書体である。筆書きの楷書ではなく、硬筆やサインペンの手書きを意識したデザインとなっており、書き方の方向や点、払いの形状を保ちつつ太さの強弱を抑えている。このフォントは、視認性と

学習用途を考慮して選定した。ヒラギノ角ゴ ProN W3 は、高速道路標識やカーナビなどで広く採用される角ゴシック体である。文字の隙間を均等に配置し、画線両端のアクセントによって可読性と存在感を両立している。日常生活での視認性の良さを評価し、比較対象として選出した。

フォントの見やすさを評価するため、読解力問題を用いた調査を実施した。問題文には注意深く読まなければ誤答する可能性がある問題を9問用意した。問題はリーディングスキルテスト[11]を基に、生成AI(ChatGPT)を活用して作成したものである。本実験には34名の被験者が参加した。被験者には、各フォントで書かれたテストを一度だけ読んだ後、回答する形式とした。実験では、フォントごとに作成した問題を画像データ化し、3問ずつフォントを切り替えてGoogleフォーム上で提示した。被験者は3つのグループに分かれ、それぞれ異なる順序でフォントを提示した。画像データ作成時には行間の調整を行わず、フォントサイズを統一することで、フォント以外の要因による影響を最小限に抑えるよう配慮した。

表3 フォント別正答率

	合計	合計
HG 教科書体	73/108	6.08/9.00
UD デジタル教科書体	84/108	7.00/9.00
ヒラギノ角ゴ W3	82/108	6.83/9.00

この結果から、視認性の順位は「UD 教科書体 > ヒラギノ角ゴ > HG 教科書体」となった。UD 教科書体が最も高い正答率を示し、ユニバーサルデザインを採用した書体の効果が確認された。ヒラギノ角ゴも比較的高い正答率を示し、日常的な視認性の高さが裏付けられる結果となった。一方、伝統的なHG教科書体は他の2つのフォントと比較して正答率が低い結果となった。これらの結果は、フォントの違いが読解力や読みやすさに影響を与える可能性を示唆している。特に、ユニバーサルデザインを採用したUD教科書体は、学習環境における有効な選択肢であることが示された。ただし、本調査のサンプル数は34件と比較的少ないため、結果の一般化には慎重を期する必要がある。また、本調査はディスレクシアの症状を持つ被験者を対象としていないため、ディスレクシア児への効果を直接的に示すものではない。この点については、今後さらなる検証が求められる。

5. 開発システムの詳細

本研究で開発したシステムは、ディスレクシア児の漢字学習を支援するためのWebアプリケーションである。以下に開発したシステムの詳細について記載する。

5.1 主要機能

このシステムは、辞書機能、検索機能、聴覚法学習機能、テスト機能、LLMを活用したルビ振り機能の5つの主要機能を持っている。辞書機能では、漢字の読み方、意味、例

文を表示する。検索機能では、文字入力およびOCR(光学文字認識)による漢字検索をサポートしている。聴覚法学習機能では、VOICEVOXを使用して、合成音声による漢字の読み上げを行い、読み上げた文章を表示する。テスト機能では、漢字の理解度を確認するためのテストを実施する。LLMを活用したルビ振り機能では、ユーザーが入力した文章を4.2.1のプロンプトでGPT4oのAPIに送信し、帰ってきた文章をHTMLで表示できる形に変換することでルビ振りした文章を表示することができる。これらの機能の組み合わせにより、ディスレクシア児の多様な学習ニーズに対応することが可能となっている。

5.2 データベース構造

本システムのデータベース構造は、漢字テーブル、読みテーブル、聴覚法テーブル、ユーザーテーブル、ユーザー漢字テーブル、およびユーザー漢字例文テーブルの6つの主要テーブルで構成されている。漢字テーブルには、各漢字そのものが格納されており、KanjiIDを主キーとして一意に識別している。読みテーブルでは、各漢字に対応する音読みおよび訓読みを複数保持可能としている。ReadingTypeカラムによって読みの種類を区別し、一つの漢字に対して複数の読みを適切に管理できる設計となっている。聴覚法テーブルは、ユーザーおよびAIによって追加された聴覚法の例文を無制限に保持できる構造であり、OverallUsageCountによって全体の使用頻度を管理している。この構造により、頻繁に使用される例文を把握し、より効果的な学習支援が可能となる。ユーザーテーブルには、ユーザーの基本情報(ユーザー名やメールアドレスなど)を保持しており、UserIDを主キーとして各ユーザーを一意に識別している。ユーザー漢字テーブルでは、各ユーザーが学習した漢字の情報を記録している。DateLearnedカラムにより学習日を管理しており、必要に応じて学習時間の記録を追加することも可能である。ユーザー漢字例文テーブルは、各ユーザーの例文使用回数を管理するためのものであり、どの例文をどの程度使用しているかを追跡可能である。このデータベース構造により例文の使用頻度を管理することで、人気の聴覚法を優先的に提示し、選択を容易にすることができる。また、お気に入りの例文を登録できる仕組みにより、カスタマイズ性の高い学習環境を提供できる。

5.3 動作フロー

システムの動作フローは、図2に示されている通りである。メイン画面から、検索、辞書、テストの各機能にアクセスできる構造となっている。検索画面では、漢字を検索した結果を表示する。辞書画面では、選択された漢字の詳細情報を表示する。学習画面では、聴覚法を用いた学習を実施する。テスト画面では、漢字の理解度を確認するテストを実施する。この構造により、ユーザーは自身の学習ニーズに応じて柔軟に各機能を利用することが可能である。

画面遷移図

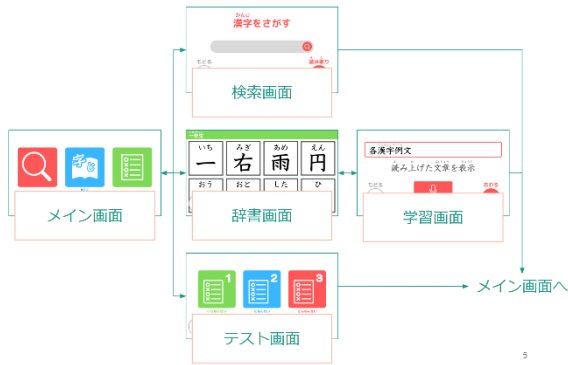


図2 画面遷移図

5.4 セキュリティ対策

セキュリティ面では、ユーザーデータの保護と安全な運用を確保するため、以下のセキュリティ対策を実施している。パスワードの暗号化、SSH/TLS 通信の採用、SQL インジェクション対策を施している。これらの対策により、ユーザーの個人情報や学習データの保護を図っている。さらに、ユーザーが撮影した OCR 用の画像はユーザー端末内で処理を行う仕様とし、本システムのサーバーには送信しない設計としている。これにより、サーバー側で保持するデータ量を削減し、安全性を向上させている。

5.5 今後の課題と展望

本研究で開発したシステムは、ディスレクシア児の漢字学習支援に一定の効果を示すことが期待されるが、いくつかの課題も明らかになっている。

現在、システムにおける聴覚法の例文作成は人手で行われており、高学年の漢字への対応やコンテンツの拡張が困難となっている。この課題に対する改善策として、LLM (Large Language Model) の活用を検証した。しかし、LLM を用いた例文生成では、漢字を過度に分解した結果、分かりづらい例文が生成されるという問題がある。この問題に対処するため、より高性能なモデルを活用した検証やプロンプトの最適化を行うことにより、例文作成の時間的コストを抑えつつ、高学年の漢字などへの拡張も可能になると考えられる。また、本研究の大きな課題として、ディスレクシア児の被験者を対象とした評価が実施できていない点が挙げられる。システムの有効性を正しく評価し、さらなる改善につなげるためには、実際のターゲットユーザーを対象とした詳細な調査が不可欠である。

6. まとめ

本研究では、発達性ディスレクシア児を対象とした漢字学習支援アプリケーションの開発を行った。本システムは、ディスレクシア児の特性に配慮しつつ、効果的な漢字学習を支援することを目的としている。開発したシステムには、以下の5つの主要機能が備わっている。

辞書機能では、漢字の読み方、意味、例文を学習でき、漢字検索機能では、文字入力および OCR (光学文字認識) を利用した漢字検索を実現する。聴覚法学習機能では、

では、VOICEVOX を使用し、音声による漢字学習を支援する。テスト機能では、漢字の理解度を確認するためのテストを実施する。ルビ振り機能では、LLM (Large Language Model) を活用し、ユーザーが提示した文章に対してルビを振る機能を提供する。また、フォントの影響調査において、UD 教科書体が最も高い読解性を示した。この知見を基に、システムデザインに UD 教科書体を反映させた。データベース設計においては、漢字や学年情報、ユーザーの学習データを効率的に管理できる構造を採用した。これにより、個々のユーザーの学習進捗に応じたカスタマイズされた学習体験を提供できるようになった。また、将来的な機能拡張や対象学年の拡大にも柔軟に対応できる拡張性を確保している。今後の課題としては、高学年への対応拡大、LLM を活用した例文生成の改善、さらにディスレクシア児を対象とした調査の実施が挙げられる。これらの課題に取り組むことで、システムの有効性をさらに高め、より広範な支援を提供することが可能になると考えられる。今後も継続的な改善を重ね、より多くの人々に寄り添う支援システムの実現を目指していく。

謝辞

本研究は JSPS 科研費 JP 23K11728 及び文教大学大学院情報学研究科共同研究費の助成を受けたものです。

参考文献

- [1] 宇野彰. 発達性読み書き障害. 高次脳機能研究, 2016, vol. 36, no. 2, p. 170-176.
- [2] 文部科学省初等中等教育局特別支援教育課. 通常の学級に在籍する特別な教育的支援を必要とする児童生徒に関する調査結果について. 2022, p. 36.
- [3] 小澤亘, 蓮尾和美, 神鞠子, 中塚愛里. ディスレクシア児童に対する「テストの音声化」支援. 日本デジタル教科書学会年次大会発表原稿集, 2016, vol. 5, p. 25-26.
- [4] 稲田健実. 特別支援教育における ICT 活用—学習障害における DAISY と音声付き教科書—. 知能と情報, 2020, vol. 32, no. 4, p. 135-140.
- [5] 小枝達也, 内山仁志, 関あゆみ. 小学 1 年生へのスクリーニングによって発見されたディスレクシア児に対する音読指導の効果に関する研究. 脳と発達, 2011, vol. 43, no. 5, p. 384-388.
- [6] 阪井恵. Universal Design for Learning (UDL) の視点でつくる音楽授業: その可能性と課題. 明星大学大学院教育学研究科年報, 2023, no. 8, p. 1-13.
- [7] 阿子島茂美, 吉野中, 漆澤恭子, 関口洋美. 漢字に親しむアプリ「かんじダス」の開発 (1): ディスレクシアのための教材アプリの技術的検討. 日本教育心理学会総会発表論文集第 57 回総会発表論文集, 2015, p. 316.
- [8] 樋熊一夫, 大庭重治. 漢字の書字学習に困難を示す生徒を対象とした ICT 活用による自己学習支援システムの開発に関する試行的研究. 特殊教育学研究, 2022, vol. 60, no. 1, p. 33-44.
- [9] 栗屋徳子, 春原則子, 宇野彰, 金子真人, 後藤多可志, 狐塚順子, 孫入里英. 発達性読み書き障害児における聴覚法を用いた漢字書字訓練方法の適用について. 高次脳機能研究, 2012, vol. 32, no. 2, p. 294-301.
- [10] VOICEVOX の全体構造について
",<https://github.com/VOICEVOX/voicevox/blob/main/docs/%E5%85%A8%E4%BD%93%E6%A7%8B%E6%88%90.md>, (参照 2024/12/20)
- [11] 新井紀子. AI vs. 教科書が読めない子どもたち. 東洋経済新報社, 2018.