

Corpus-Driven Analysis of Phonological and Morphemic Errors in the Tilawati Textbook: Evidence for Improving Qur'anic Literacy Materials

M. Baihaqi*

Universitas Islam Negeri Sunan Ampel Surabaya, Indonesia

Nadlir

Universitas Islam Negeri Sunan Ampel Surabaya, Indonesia

Candra Darmawan

Universitas Islam Negeri Raden Fatah Palembang, Indonesia

Ahmad Jundi Al Mubarak

Telkom University Surabaya, Indonesia

Nabilah Robbaniyah

Telkom University Surabaya, Indonesia

Abstract—This study critically examines the linguistic integrity of Tilawati, one of the most widely used Qur'an-based instructional textbooks for reading practice in Indonesia, with particular attention to the accuracy of its phonological and morphemic representations. The complete textbook was digitized and compiled into an annotated corpus, enabling the identification of two-dimensional error patterns across phonological (vowel and consonant) and morphemic (derivation, root–pattern, and affixation) levels. The analysis is theoretically grounded in error analysis theory, interlanguage theory, and core principles of corpus linguistics. The findings demonstrate that phonological and morphemic errors occur systematically rather than randomly, with a higher concentration of errors appearing in the initial stages of instruction. The most recurrent error categories involve vowel-related phonological deviations and root–pattern morphemic inaccuracies, predominantly manifested through substitution and omission processes. In addition, several errors exhibit cross-dimensional co-occurrence, indicating a compounded linguistic load at the level of individual tokens and instructional contexts. Such patterns suggest that inaccurate linguistic input at the level of basic linguistic units in the corpus may hinder learners' development of stable phonological and morphemic representations. From a pedagogical perspective, these results underscore the importance of providing precise and internally consistent linguistic input in Qur'anic reading materials. Methodologically, the study demonstrates the value of corpus-driven approaches for the systematic evaluation of instructional textbooks. Overall, the findings provide an empirical foundation for revising the Tilawati series and contribute to broader efforts to enhance the accuracy, pedagogical coherence, and linguistic reliability of Qur'anic literacy instruction (QLI) within formal educational contexts.

Index Terms—corpus-driven analysis, error analysis theory, phonological errors, morphemic errors, Qur'anic literacy instruction

I. INTRODUCTION

Reading the Qur'an constitutes an essential component of literacy within Islamic education in Indonesia (Supriadi et al., 2022). It occupies a central position in elementary madrasah education and community-based Qur'anic learning centers, where instruction integrates recitation with the application of tajwīd rules to ensure phonological accuracy (Basir, 2024). The Tilawati textbook has been widely distributed, with more than 700,000 users across Indonesia and several neighboring countries. Given its extensive use, the urgency for phonologically and linguistically accurate instructional materials becomes increasingly pressing. In early literacy development, the quality of phonological and morphemic input plays a crucial role in shaping learners' initial reading skills (Deacon et al., 2024; Share, 2021; Verhoeven & Perfetti, 2022). As reading provides a foundation for language learning (Damra et al., 2025), an empirical assessment of the linguistic accuracy of the Tilawati textbook is necessary to ensure effective Qur'anic reading instruction and to enhance the overall quality of Arabic literacy education in Indonesia.

Phonological and morphemic representations are particularly critical in Qur'anic reading instruction grounded in

Arabic. Empirical research has consistently demonstrated the impact of vowel marking, consonantal structure, and morphological patterning on Arabic reading acquisition (Saiegh-Haddad & Everatt, 2017; Tibi & Kirby, 2018). Inaccuracies in vowel diacritics or morphemic patterns within instructional materials may lead learners to develop unstable phonological and grammatical representations, which can negatively affect reading accuracy (Hussin et al., 2023). Although textbooks serve as cornerstone sources of linguistic input, systematic linguistic evaluations of Qur'anic qirā'ah materials remain limited, particularly in non-native Arabic contexts (Haris, 2022). Moreover, corpus-driven techniques for validating the linguistic quality of instructional materials have not yet been widely implemented in the Indonesian context (Alhawary, 2019).

Although error analysis has been extensively applied to learner-produced data, relatively few studies have examined the linguistic quality of teaching materials themselves. Curriculum and textbook reviews are often guided by intuitive judgments rather than evidence-based, systematic analyses, leaving important dimensions of instructional quality under-researched (Nadlir et al., 2025). The majority of textbook studies are not supported by corpus evidence and do not employ multifactorial annotation models (Biber et al., 2024; Biancardi et al., 2021). To date, no study has combined phonological–morphemic error classification, canonical error categories, and cross-dimensional co-occurrence analysis in an evaluation of the Tilawati textbook. This gap highlights the need for a corpus-driven linguistic audit of widely used Qur'anic reading materials.

From a theoretical perspective, interlanguage theory posits that systematic inaccuracies in instructional input may influence learners' developing linguistic systems and contribute to fossilization (Oliveira & Flasco, 2021). Such errors often exceed learners' ability to self-correct, particularly at early and intermediate stages of acquisition, and may pose persistent challenges even for more advanced learners. In the case of the Tilawati textbook, which is used by large numbers of learners, recurring input-related inaccuracies may therefore have a systemic impact. A corpus-driven approach enables the objective and quantifiable identification of linguistic irregularities across both micro-level units (e.g., tokens and morphemes) and macro-level instructional contexts (Klie et al., 2023). Accordingly, a systematic corpus-driven linguistic audit is required to ensure that Qur'anic reading instruction is grounded in accurate and reliable linguistic input.

Building on this theoretical and methodological foundation, the present study constructs an annotated corpus comprising the complete Tilawati textbook and applies error analysis to identify phonological and morphemic error patterns. The study addresses three main objectives: (1) to determine the extent and distribution of phonological and morphemic errors; (2) to analyze error patterns across instructional units; and (3) to examine cross-dimensional error co-occurrence in order to assess whether overlapping linguistic loads emerge within specific instructional contexts.

This study makes a methodological contribution by introducing a comprehensive manual multi-layer annotation system for phonological and morphemic features within a purpose-built two-dimensional corpus, combined with frequency, distributional, and co-occurrence analyses characteristic of contemporary corpus linguistics. Theoretically, it integrates error analysis theory and interlanguage theory to examine the pedagogical implications of linguistic input quality. Practically, the findings provide evidence-based recommendations for revising the Tilawati textbook, informing teacher training program development, and supporting curriculum design grounded in linguistic accuracy. This integration of corpus-driven methodology and textbook evaluation aligns with a growing trend in international research emphasizing empirical validation of instructional materials (Khojasteh & Mukundan, 2025).

II. LITERATURE REVIEW

Research on Qur'anic literacy and Arabic reading instruction consistently emphasizes that phonological and morphemic accuracy are crucial for early literacy acquisition. In Arabic, vowel fidelity and consonantal precision are not merely phonetic details; rather, they are central determinants of lexical access and decoding efficiency. Empirical studies in psycholinguistics and reading research demonstrate that inaccurate pronunciation of words or non-words, particularly involving short and long vowels (Tibi & Kirby, 2018), can interfere with phoneme–grapheme mapping and hinder the formation of stable orthographic representations. Textbooks constitute the primary source of linguistic input in instructional settings; as such, inaccuracies in learning materials may shape learners' developing phonological systems. The issue is particularly critical in Qur'anic reading, in which accuracy is both pedagogically and theologically mandated.

Beyond phonology, Arabic morphology—especially the root–pattern system—plays a fundamental role in lexical access and semantic processing. Morphological awareness has been shown to interact closely with phonological awareness in supporting reading development in Semitic languages (Deacon et al., 2024). Deviations in orthographic representation may therefore influence not only pronunciation but also morphological parsing and semantic interpretation. Nevertheless, the majority of research on Arabic literacy has focused on learner output rather than on the linguistic quality of instructional materials. Consequently, textbook-based sources of error remain underexplored, particularly in relation to early interlanguage development.

Historically, error analysis has been employed to examine learner language and to identify systematic developmental patterns and deviations. Both classical and contemporary perspectives within error analysis theory maintain that errors are not random mistakes but systematic manifestations of underlying linguistic processing mechanisms (Chuang, 2024). However, the application of error analysis to textbook evaluation remains limited. Assessments of instructional materials often rely on expert judgment or qualitative impressions without systematic empirical validation. Recent advancements in corpus linguistics have addressed this limitation by enabling frequency-based, replicable, and large-scale analyses of

linguistic features across complete texts (Biber et al., 2024).

Corpus-driven methods further allow researchers to examine not only individual errors but also their distributional patterns and cross-dimensional relationships. Co-occurrence analysis, a technique within corpus linguistics, can reveal compounded linguistic loads arising from the simultaneous presence of multiple deviations within the same lexical items or instructional segments. Such layered irregularities may increase cognitive processing demands in language learning contexts (Tabari et al., 2025). However, these analytical approaches have rarely been applied to Qur'anic reading materials and are even less common in the Indonesian educational context. This research gap underscores the need for a corpus-driven, multi-layered investigation of widely used textbooks such as the Tilawati series, which serve large populations of beginning learners and significantly shape foundational Arabic literacy development.

III. MATERIALS AND METHODS

Taking instructional materials as a source of linguistic input, this study employed a consistent and replicable methodological procedure to examine phonological and morphemic errors in the Tilawati textbook. Drawing on the tenets of error analysis and corpus linguistics, the study combined corpus compilation, multi-dimensional annotation, and statistical pattern analysis to ensure empirical validity and analytical transparency (Tognini-Bonelli, 2001). This section describes the research design, corpus construction, annotation framework, reliability procedures, and data analysis techniques applied in the study.

A. Research Design

The study implemented a corpus-driven error analysis framework in which analytical categories and generalizations were induced from observed linguistic data rather than tested against predetermined hypotheses. This approach was grounded in error analysis theory, which conceptualizes linguistic errors as systematic manifestations of underlying language systems rather than as arbitrary anomalies (Chuang, 2024). To support interpretive analysis, the study also drew on interlanguage theory, which views learner-related errors as developmental phenomena shaped by instructional input and exposure (Vasileanu & Vasile, 2024). Methodologically, the study was guided by core principles of corpus linguistics, which emphasize empirical generalization based on documented patterns across large bodies of linguistic data rather than on intuition or isolated examples (Tognini-Bonelli, 2001).

B. Materials and Corpus Construction

The corpus comprised the complete primary-level reading materials of the Tilawati series, one of the most widely used textbooks for Qur'anic reading instruction in Indonesia. All texts were digitized and compiled into a structured linguistic corpus following established guidelines for representativeness and accurate encoding of orthographic and diacritic forms (Biber et al., 2024). Word tokens were treated as the basic unit of analysis, allowing for systematic examination of phonological and morphemic accuracy across instructional units. This token-based corpus design is widely adopted in applied corpus research to facilitate reliable frequency and distributional analyses (Jablonkai et al., 2024).

C. Error Annotation Framework

The corpus was annotated using a two-tier error annotation model designed to capture cross-dimensional linguistic deviations. The first annotation layer addressed phonological errors, including vowel-related and consonant-related deviations. The second layer focused on morphemic errors, encompassing derivational inaccuracies, root-pattern errors, and affixation errors. In addition, each error instance was classified according to the four canonical categories of error analysis—substitution, omission, addition, and mis-selection—to enable systematic cross-category and cross-level comparison (Huidrom & Belz, 2023; Mellou & Sparos, 2005). Comparable multi-layer annotation approaches have been shown to enhance analytical precision in learner and textbook corpus research (Gessler et al., 2022).

D. Annotation Procedure and Reliability

To ensure consistency and objectivity, two trained annotators independently annotated a randomly selected subset of the corpus. Inter-annotator reliability was calculated using Cohen's kappa, a statistical measure widely employed in linguistic and psycholinguistic research to assess agreement on categorical annotation (Chamid et al., 2024). Following established methodological standards, a κ value of 0.80 or higher was set as the threshold for acceptable agreement, indicating high annotation reliability (Wong et al., 2021). All discrepancies were resolved through consensus discussion prior to full corpus annotation.

E. Data Analysis Techniques

Data analysis proceeded through three complementary corpus-driven procedures. First, frequency analysis was conducted to determine the overall occurrence rates of each error type across the corpus. Second, distributional analysis examined the dispersion of errors across instructional units to identify sections with heightened linguistic load. Finally, co-occurrence analysis using cross-tabulation matrices was applied to determine whether phonological and morphemic errors occurred jointly within the same lexical tokens or instructional contexts. Such co-occurrence patterns provide insight into compounded linguistic load that may increase cognitive processing demands for learners (Ruamsuk et al., 2025).

IV. RESULTS AND DISCUSSION

A. Results

To obtain an overview of the deviation patterns identified in the Tilawati textbook, all errors were first classified into phonological and morphemic categories. This distinction is essential for identifying the number of deficiencies at each linguistic level and for determining which dimensions are most affected by instructional inaccuracies. The major findings are summarized in Table 1.

TABLE 1
CLASSIFICATION OF LINGUISTIC ERRORS

Category	Subcategory	Frequency (n)	Percentage (%)	Incorrect	Correct	Description
Phonological	Vowel	128	42.52	كبيرة	كبيرة	The vowel (bi) is not long enough, so an additional vowel (i) should be added after (bi).
Phonological	Consonant	20	6.64	بيت	ثبت	The initial consonant (ba) is incorrect and should be replaced with another correct consonant, namely (tsa).
Morphemic	Derivational	23	7.64	خرج	خرج	The derivative of this word is a verb, not a noun, so it should be returned to the verb pattern (faala).
Morphemic	Root-Pattern	129	42.86	تزم	التزم	The origin of the word is (lazima-iltazama), so the origin of the word must be corrected.
Morphemic	Affixation	1	0.33	نعمة	نعم	The addition of ta marbutoh is inappropriate in the word (naima), so it should be removed to become (naima).
Total		301	100			

Table 1 indicates that phonological and morphemic errors occur in comparable proportions, although their internal distributions differ. Vowel-related phonological errors constitute the largest category (42.52%), followed closely by root-pattern morphemic errors (42.86%). By contrast, consonantal phonological errors and derivational morphemic errors occur less frequently, while affixation errors are rare. These findings suggest that the most vulnerable aspects of the Tilawati textbook concern vowel representation and its alignment with root-pattern morphology, both of which are essential for accurate decoding and meaning construction in Arabic reading.

To further classify the identified deviations, the analysis applied an error analysis framework that distinguishes four canonical error types: substitution, omission, addition, and mis-selection. This classification enables a more fine-grained examination of the linguistic processes underlying phonological and morphemic errors observed in the Tilawati textbook. The distribution of error types based on this framework is presented in Table 2.

TABLE 2
ERROR TYPES BASED ON AN ERROR ANALYSIS FRAMEWORK

Error Type	Phonological (n)	Morphemic (n)	Total	Incorrect	Correct	Description
Substitution	134	149	283	بتم	بتر	Error in consonant (ma) which should be (ra)
Omission	13	2	15	سهلة	سهولة	There is a missing vowel (u) so it should be returned to the long reading (hu)
Addition	0	1	1	نعمة	نعم	The addition of ta marbutoh is not in its place in the word (naima) so it should be removed to become (naima)
Mis-selection	1	1	2	بئس	بأس	Wrong choice of form available in the language system
Total	148	153	301			

As shown in Table 2, substitution is by far the most frequent error type, accounting for 283 out of 301 total errors across both phonological and morphemic domains. Omission follows with 15 instances, while addition and mis-selection occur only marginally, with one and two cases respectively. This distribution clearly indicates that the overwhelming majority of deviations in the Tilawati textbook stem from the replacement of linguistic elements rather than from overgeneration or incorrect form selection. The dominance of substitution suggests systematic alterations in phoneme or morpheme realization, which ultimately affect the structural integrity of the affected words.

To determine whether errors cluster in particular instructional units, a distributional analysis across units was conducted. This analysis allows assessment of whether linguistic load is evenly distributed throughout the instructional sequence. The results are presented in Table 3.

TABLE 3
DISTRIBUTION OF ERRORS ACROSS UNITS

Unit	Phonological (n)	Morphemic (n)	Total	Percentage (%)
Unit 1	38	39	77	25.58
Unit 2	110	114	224	74.42
Unit 3	0	0	0	0
Unit 4	0	0	0	0
Unit 5	0	0	0	0
Total	148	153	301	100

Table 3 demonstrates that all identified errors are concentrated in Unit 1 and Unit 2, with Unit 2 alone accounting for 74.42% of the total 301 errors. No errors were detected in Units 3–5. This highly uneven distribution indicates that inaccuracies are confined to the earliest instructional stages, which are crucial for establishing foundational phonological and morphemic representations. Errors introduced at this foundational level may become entrenched and potentially reinforced throughout subsequent stages of learning.

Beyond examining errors in isolation, the study also investigated whether phonological and morphemic errors co-occur within the same lexical items or instructional contexts. To this end, a co-occurrence analysis using cross-tabulation matrices was conducted. An overview of the results is provided in Table 4.

TABLE 4
CO-OCCURRENCE MATRIX

Category	Phonology Only	Morpheme Only	Co-occurring	Total
Substitution	24	30	229	283
Omission	13	2	0	15
Addition	0	1	0	1
Mis-selection	1	0	1	2
Total	38	33	230	301

Table 4 demonstrates that error co-occurrence is highly frequent, with 230 cases involving simultaneous phonological and morphemic deviations. The vast majority of these co-occurring cases stem from substitution errors (229 instances), while omission and addition do not exhibit co-occurrence, and mis-selection appears only marginally. This pattern indicates that when linguistic elements are replaced, the deviation often affects both phonological realization and morphemic structure simultaneously. Such intertwined disruptions suggest that certain lexical items carry concurrent phonological and morphological processing demands, potentially interfering with learners' ability to establish stable and accurate linguistic representations.

To illustrate these patterns more concretely, representative error examples were selected based on frequency, typicality, and pedagogical relevance. These examples are presented in Table 5.

TABLE 5
REPRESENTATIVE ERROR EXAMPLES

No	Incorrect Form	Correct Form	Category	Error Type	Explanation
1	كبيرة	كبيرة	Phonological – vowel	Omission	Phonological error in (bi) the omission of the vowel i which should be lengthened
2	بتم	بتر	Phonological – consonant	Substitution	There is a replacement of the consonant (ra) which does not correspond to the root form of the word
3	فججا	فججا	Phonological – vowel	Omission	Phonological error in (ja) the omission of the vowel a which should be lengthened

Table 5 provides illustrative instances of vowel omission and consonant substitution errors, such as كبيرة instead of كبيرة and بتم instead of بتر. These examples demonstrate that seemingly minor phonemic deviations may result in misalignment with well-formed root–pattern structures or even alter grammatical interpretation. From the perspective of error analysis theory, such deviations reflect systematic input-related inaccuracies rather than incidental mistakes, particularly when they recur across instructional units. Within an interlanguage framework, repeated exposure to these forms may shape learners' developing phonological and morphemic representations, increasing the likelihood that inaccurate patterns are internalized as part of the evolving linguistic system. In cases of vowel omission, the loss of vocalic information may obscure morphological boundaries and weaken phoneme–grapheme correspondence, while consonant substitution can distort the underlying trilateral root and disrupt form–meaning mapping. Given the centrality of vocalization and root–pattern morphology in Arabic, these micro-level deviations may accumulate and impose additional cognitive processing demands during decoding. Consequently, the examples in Table 5 highlight how inaccuracies in instructional materials can function as a source of systematic input error, with potential implications for the stabilization of learners' phonological and morphemic knowledge in Qur'anic literacy instruction.

B. Discussion

The findings of this study demonstrate the presence of systematic phonological and morphemic errors in the Tilawati textbook, concentrated particularly in its early instructional units. As shown in Table 1, vowel-related phonological errors

(128 cases) and root–pattern morphemic errors (129 cases) constitute the dominant categories. This pattern indicates that the most vulnerable aspects of the instructional input concern vowel representation and root–pattern morphology—two foundational components of accurate Arabic reading (Saiegh-Haddad & Everatt, 2017; Tibi & Kirby, 2018). Inaccurate grapheme representation at these levels may interfere with learners’ pronunciation, decoding processes, and comprehension from the outset of instruction.

The substantial proportion of vowel-related errors warrants particular attention. As illustrated in Tables 1 and 5, many of these cases involve vowel omission or inaccurate vowel realization, which may weaken phoneme–grapheme correspondence and disrupt the development of fluent decoding skills (Deacon et al., 2024; Share, 2021). Beginning learners, who rely heavily on textbooks as their primary source of linguistic input, may consequently develop unstable phonological representations when exposed repeatedly to inaccurate vocalization patterns. This interpretation aligns with research indicating that early exposure to incorrect phonological input can have long-term implications for reading development (Verhoeven & Perfetti, 2022).

Morphemic errors, especially those involving root–pattern structures, are similarly consequential. As indicated in Table 1, root–pattern errors represent the largest category of morphemic deviations (129 cases). Given the central role of the root–pattern system in Arabic morphology, inaccuracies at this level may obscure systematic relationships between form and meaning. Research on Arabic literacy has demonstrated that morphological awareness interacts closely with phonological awareness in supporting reading accuracy and vocabulary growth (Deacon et al., 2024; Saiegh-Haddad et al., 2023). From this perspective, repeated exposure to incorrect root–pattern forms may hinder the development of robust morphemic representations and compromise learners’ ability to infer word structure.

The classification of errors according to the error analysis framework (Table 2) further clarifies the underlying processes of deviation. Substitution overwhelmingly emerges as the most frequent error type (283 out of 301 cases), followed distantly by omission (15 cases), while addition (1 case) and mis-selection (2 cases) occur only marginally. This distribution indicates that the primary issue involves the replacement of linguistic elements rather than overgeneration or incorrect form selection. Within error analysis theory, substitution processes are often associated with unstable or incomplete linguistic representations (Chuang, 2024; Huidrom & Belz, 2023). When such patterns appear in instructional materials, they carry particular pedagogical significance, as textbooks function as authoritative models of linguistic form.

The distributional findings presented in Table 3 reveal a highly uneven concentration of errors across instructional units. All identified errors occur in Units 1 and 2, with Unit 2 alone accounting for 74.42% of the total deviations (224 out of 301). No errors were detected in Units 3–5. This concentration is pedagogically significant because the initial units are designed to establish foundational phonological and morphemic awareness. From the perspective of interlanguage theory, systematic inaccuracies encountered during early stages of learning may become entrenched and resistant to later correction if not addressed promptly (Oliveira & Flasco, 2021). Therefore, the high error density in the earliest units of the Tilawati textbook may exert a disproportionate influence on learners’ long-term reading development.

The co-occurrence analysis (Table 4) provides additional insight into the structural complexity of the identified deviations. A substantial majority of errors (230 out of 301) involve simultaneous phonological and morphemic deviations. Nearly all of these co-occurring cases stem from substitution processes (229 cases), while omission and addition do not display co-occurrence, and mis-selection appears only marginally. This pattern indicates that when substitution occurs, it frequently affects both phonological realization and morphemic structure at the same time. Such overlapping deviations may increase cognitive processing demands during decoding, thereby challenging learners’ ability to construct stable phonological and morphemic representations (Alnafisah et al., 2022). In instructional contexts, this compounded linguistic load may reduce the effectiveness of the input provided to beginning readers.

The examples presented in Table 5 further illustrate how even minor orthographic variations can produce substantial linguistic consequences. Vowel deletion and consonant substitution not only alter pronunciation but may also distort root–pattern alignment and grammatical interpretation. From the standpoint of error analysis theory, these deviations represent systematic input-level inaccuracies rather than isolated mistakes. Within an interlanguage framework, repeated exposure to such forms may influence the stabilization of learners’ developing linguistic systems. Consequently, inaccuracies in instructional materials should not be regarded as trivial, as they may cumulatively shape learners’ phonological and morphemic knowledge in Qur’anic literacy instruction.

Overall, the findings demonstrate that a corpus-driven approach provides a rigorous and empirically grounded method for evaluating the linguistic quality of instructional materials. By systematically analyzing frequency, distribution, and co-occurrence patterns, the approach reveals structural tendencies that may not be apparent through impressionistic textbook review (Biber et al., 2024). In the context of Qur’anic literacy instruction, such evidence-based evaluation is essential to ensure that teaching materials support, rather than hinder, the acquisition of accurate phonological and morphemic knowledge.

From a pedagogical perspective, these findings underscore the urgent need for systematic linguistic verification in the development and revision of Qur’anic literacy textbooks. Given that all identified errors are concentrated in the foundational units and that a substantial proportion involves substitution processes affecting both phonological and morphemic structures, editorial review should prioritize accuracy in vowel representation and root–pattern formation. Textbook validation procedures may benefit from incorporating corpus-based proofreading, multi-layer linguistic screening (phonological and morphological), and collaboration with specialists in Arabic linguistics before publication.

Additionally, teachers using the current edition should be encouraged to implement corrective scaffolding strategies, such as explicit phoneme–grapheme reinforcement and guided morphological analysis, to mitigate potential negative input effects. By strengthening both material design and classroom mediation, instructional contexts can better safeguard learners from internalizing unstable phonological and morphemic representations during the critical early stages of literacy acquisition.

V. CONCLUSIONS

This study demonstrates that phonological and morphemic errors in the early instructional units of the Tilawati textbook are systematic rather than random. Vowel representation and root–pattern morphology emerge as the primary sources of deviation, while substitution overwhelmingly constitutes the dominant error mechanism at the word level. These patterns indicate that inaccurate instructional input may hinder the development of stable phonological decoding and morphemic sensitivity among beginning learners.

The complete concentration of errors in the initial units, together with the high rate of phonological–morphemic co-occurrence, suggests that learners are exposed to structurally problematic input at a formative stage of Qur’anic literacy instruction. From the perspectives of error analysis theory and interlanguage theory, such recurring input-level inaccuracies may influence the stabilization of learners’ developing linguistic systems and increase cognitive processing demands during decoding. These findings highlight the value of corpus-driven evaluation for identifying structural weaknesses in instructional materials and underscore the need for textbook revision focused on accurate vowel marking and consistent root–pattern morphology. Future research may examine the effects of revised materials on learners’ reading development.

REFERENCES

- [1] Alhawary, M. T. (2019). *Arabic Second Language Learning and Effects of Input, Transfer, and Typology*. Washington, DC: Georgetown University Press. <https://doi.org/10.2307/j.ctvcb5c6q>
- [2] Alnafisah, M., Goodale, E., Rehman, I., Levis, J., & Kochem, T. (2022). The impact of functional load and cumulative errors on listeners’ judgments of comprehensibility and accentedness. *System*, 102906. <https://doi.org/10.1016/j.system.2022.102906>
- [3] Basir, A. (2024). Enhancing Qur’an Reading Proficiency in Madrasahs Through Teaching Strategies. *Nazhruna: Jurnal Pendidikan Islam*, 7(2), 373–389. <https://doi.org/10.31538/nzh.v7i2.4985>
- [4] Biancardi, B., Ceccaldi, E., Clavel, C., Chollet, M., & Dinkar, T. (2021). CATS2021: International Workshop on Corpora And Tools for Social skills annotation. In *Proceedings of the 2021 International Conference on Multimodal Interaction (ICMI 2021)* (pp. 857–859). Association for Computing Machinery. <https://doi.org/10.1145/3462244.3480977>
- [5] Biber, D., Larsson, T., & Hancock, G. (2024). The linguistic organization of grammatical text complexity: comparing the empirical adequacy of theory-based models. *Corpus Linguistics and Linguistic Theory*, 20, 347–373. <https://doi.org/10.1515/cllt-2023-0016>
- [6] Chamid, Ahmad, A., Widowati, & Kusumaningrum, R. (2024). Labeling Consistency Test of Multi-Label Data for Aspect and Sentiment Classification Using the Cohen Kappa Method. *Ingenierie Des Systemes d’Information*, 29(1), 161–167. <https://doi.org/10.18280/isi.290118>
- [7] Chuang, F.-Y. (2024). Automated grammatical error correction in researching written learner language. In *Routledge handbook of technological advances in researching language learning*. Routledge. <https://doi.org/10.4324/9781003459088-8>
- [8] Damra, H. M., Abu-Helu, S. Y., & Al Masadeh, A. (2025). The Influence of Audiobooks on the Development of English as a Foreign Language Learners’ Reading Fluency in Jordan. *Theory and Practice in Language Studies*, 15(3), 786–796. <https://doi.org/10.17507/tpls.1503.13>
- [9] Deacon, S. H., Robertson, E. K., Ryken, A., & Levesque, K. (2024). The magic in magician: Contributions of phonological dimensions of morphological awareness to children’s reading development. *Journal of Research in Reading*, 47, 12439. <https://doi.org/10.1111/1467-9817.12439>
- [10] Gessler, L., Levine, L., & Zeldes, A. (2022). Midas Loop: Prioritized Human-in-the-Loop Annotation for Large Scale Multilayer Data. In *Proceedings of the 16th Linguistic Annotation Workshop (LAW-XVI) within LREC 2022* (pp. 103–110). European Language Resources Association. <https://aclanthology.org/2022.law-1.13/>
- [11] Haris, A. (2022). Teaching Reading of Arabic Language in Indonesia: Reconstruction of the Contents and Scope of Nahwu science. *Eurasian Journal of Applied Linguistics*, 8(2), 122–136. <https://doi.org/10.32601/ejal.911547>
- [12] Huidrom, R., & Belz, A. (2023). Towards a Consensus Taxonomy for Annotating Errors in Automatically Generated Text. In *International Conference Recent Advances in Natural Language Processing, RANLP* (pp. 527–540). https://doi.org/10.26615/978-954-452-092-2_058
- [13] Hussin, M., Ismail, Z., & Naimah. (2023). Error Analysis of Form Four KSSM Arabic Language Text Book in Malaysia. *Theory and Practice in Language Studies*, 13(1), 175–185. <https://doi.org/10.17507/tpls.1301.20>
- [14] Jablonkai, R. R., Kim, J., & Yan, R. (2024). A corpus approach to systematic literature reviews. In *Routledge handbook of technological advances in researching language learning*. Routledge. <https://doi.org/10.4324/9781003459088-40>
- [15] Khojasteh, L., & Mukundan, J. (2025). A Systematic Review of Corpus-Based Methodologies in Textbook Analyses and Evaluation. *Language Teaching Research Quarterly*, 47, 175–195. <https://doi.org/10.32038/ltrq.2025.47.10>
- [16] Klie, J., Webber, B., & Gurevych, I. (2023). Annotation Error Detection: Analyzing the Past and Present for a More Coherent Future. *Computational Linguistics*, 49(1), 157–198. https://doi.org/10.1162/coli_a_00464
- [17] Mellou, K., & Sparos, L. (2005). Systematic errors in etiological epidemiological studies. *Archives of Hellenic Medicine*, 22(2), 199–207.

- [18] Nadlir, Mukhlisah, Baihaqi, M., & Huda, H. (2025). Humanizing teacher education for madrasah contexts: a curriculum model integrating ethical reflection on socio-scientific issues model integrating ethical reflection on socio-scientific issues. *Cogent Education*, 12(1). <https://doi.org/10.1080/2331186X.2025.2583513>
- [19] Oliveira, D., & Flasmio, P. (2021). How to promote the pragmatic awareness and avoid the fossilization phenomenon through a role play activity in English as a foreign language classes. *Brazilian English Language Teaching Journal*, 1–11. <https://doi.org/10.15448/2178-3640.2021.1.41723>
- [20] Ruamsuk, Y., Mingkhwan, A., & Unger, H. (2025). DIFFSTRACT: distinguishing the content of texts. In *Proceedings of the 2022 6th International Conference on Natural Language Processing and Information Retrieval (NLPPIR '22)* (pp. 31–35). Association for Computing Machinery. <https://doi.org/10.1145/3582768.3582787>
- [21] Saiegh-Haddad, E., & Everatt, J. (2017). Early literacy education in Arabic. In *The Routledge International Handbook of Early Literacy Education* (1st ed., pp. 185–199). Routledge. <https://doi.org/10.4324/9781315766027-17>
- [22] Saiegh-Haddad, E., Ghawi-Dakwar, O., Haj, L., Farraj-Bsharat, R., & Laks, L. (2023). SPELLING ARABIC: When Does Orthographic Knowledge End and Language Knowledge Start? In *Routledge International Handbook of Visualmotor Skills, Handwriting, and Spelling: Theory, Research, and Practice* (pp. 276–292). Routledge. <https://doi.org/10.4324/9781003284048-25>
- [23] Share, D. L. (2021). Is the Science of Reading Just the Science of Reading English? *Reading Research Quarterly*, 391–402. <https://doi.org/10.1002/rq.401>
- [24] Supriadi, U., Supriyadi, T., & Abdussalam, A. (2022). Al-Qur'an Literacy: A Strategy and Learning Steps in Improving Al-Qur'an Reading Skills through Action Research. *International Journal of Learning, Teaching and Educational Research*, 21(1), 323–339. <https://doi.org/10.26803/ijlter.21.1.18>
- [25] Tabari, M. A., Lu, X., & Wang, Y. (2025). Mapping the associations among task complexity, emotionality, and linguistic complexity in L2 writing. *Language Awareness*, 34(2), 345–372. <https://doi.org/10.1080/09658416.2024.2384754>
- [26] Tibi, S., & Kirby, J. R. (2018). Investigating Phonological Awareness and Naming Speed as Predictors of Reading in Arabic. *Scientific Studies of Reading*, 22(1), 70–84. <https://doi.org/10.1080/10888438.2017.1340948>
- [27] Tognini-Bonelli, E. (2001). *Corpus linguistics at work*. John Benjamins Publishing Company. <https://doi.org/10.1075/scl.6>
- [28] Vasileanu, M., & Vasile, C. M. (2024). An error analysis of possessor encodings in Romanian interlanguage: A corpus-based study. *Journal of Linguistic and Intercultural Education*, 17(3), 177–194. <https://doi.org/10.29302/jolie.2024.17.3.10>
- [29] Verhoeven, L., & Perfetti, C. (2022). Scientific Studies of Reading Universals in Learning to Read Across Languages and Writing Systems Universals in Learning to Read Across Languages and Writing. *Scientific Studies of Reading*, 26(2), 150–164. <https://doi.org/10.1080/10888438.2021.1938575>
- [30] Wong, K., Paritosh, P., & Aroyo, L. (2021). Cross-replication Reliability - An Empirical Approach to Interpreting Inter-rater Reliability. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)* (pp. 7053–7065). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2021.acl-long.548>

M. Baihaqi, M.A., Ph.D., is an Associate Professor in the Postgraduate Program at Universitas Islam Negeri Sunan Ampel Surabaya, Indonesia, specializing in Arabic Language Education. He obtained his doctoral degree from Sudan University with a concentration in Arabic Language Education.

His research and teaching interests include Arabic language pedagogy, corpus linguistics, and Qur'anic literacy instruction. His works focus on error analysis, instructional material development, and data-driven approaches to improving Arabic language learning. ORCID iD: <https://orcid.org/0009-0007-7171-8449>

Nadlir, M.Ag., is an Associate Professor in the Postgraduate Program at Universitas Islam Negeri Sunan Ampel Surabaya, Indonesia. His area of expertise is the development of Islamic Education curricula.

His academic work includes curriculum development, instructional innovation, and the integration of Islamic values into educational systems. He is also actively involved in developing value-based and character-oriented educational models.

Candra Darmawan, M.A., is an Associate Professor at the Faculty of Da'wah and Communication, Universitas Islam Negeri Raden Fatah, Indonesia. His expertise lies in Islamic preaching (da'wah) and Islamic civilization.

His research focuses on da'wah studies, Islamic communication, and the dynamics of civilization in contemporary society. He is also actively engaged in developing Islamic studies relevant to modern social challenges.

Ahmad Jundi Al Mubarak is an academic at Telkom University Surabaya, Indonesia, Faculty of Electrical Engineering, in the Computer Engineering Program.

His research interests include Arabic text analysis, natural language processing, and the application of technology in linguistic studies. His works contribute to the development of computational approaches to Arabic linguistic analysis in the Indonesian context. ORCID iD: <https://orcid.org/0009-0006-6736-5735>

Nabilah Robbaniyah is an academic at Telkom University Surabaya, Indonesia, in the Informatics Program.

Her research interests include information technology, data analysis, and computational system development. She is involved in research integrating technological approaches with language analysis and educational applications.