



A Morphological Mismatch? Comparing Turkish High School English Textbooks and University Entrance Exams

形態不一致？比較土耳其高中英語教科書與大學入學英語考試

Tan Arda Gedik^{1,2}

Received: 30 March 2025 / Revised: 4 June 2026 / Accepted: 21 July 2026
© The Author(s) under exclusive licence to National Taiwan Normal University 2026

Abstract

This study investigates the morphological complexity of high school English textbooks and English university entrance exams in Turkey, focusing on indices of inflectional and derivational morphology, morpheme family size, frequency, and morphological density. Using automated corpus-based analyses, the study compares eight textbooks and ten exams across multiple morphological measures. The results indicate systematic differences between the two corpora, with the most robust effects emerging in derivational and prefix-related indices. Exam materials consistently exhibited higher derivational complexity, greater prefix use, and larger prefix family sizes, whereas many inflectional and suffix-related differences did not remain robust after correction for multiple comparisons. These findings suggest that university entrance exams rely more heavily on morphologically productive word-formation processes than those emphasized in official high school textbooks. The results point to a structural misalignment between instructional materials and assessment demands. The study highlights the importance of incorporating morphologically rich input and explicit attention to word-formation processes in English curricula to better support students' preparation for high-stakes examinations and the transition to tertiary education.

摘要

本研究探討土耳其高中英語教科書與大學入學英語考試在形態複雜度上的差異，重點分析屈折形態、派生形態、詞素家族大小、詞頻以及形態密度等指標。研究採用自動化語料庫分析方法，比較八套高中英語教科書與十份大學入學英語考試，並從多項形態學指標進行分析。研究結果顯示，兩類語料在多項形態指標上存在系統性的差異，其中以派生形態及前綴相關指標的差異最為明顯。與高中英語教科書相比，大學入學英語考試持續呈現較高的派生形態複雜度、較頻繁的前綴使用，以及較大的前綴家族大小；然而，多數屈折形態及後

Extended author information available on the last page of the article

Published online: 10 August 2026

Springer

綴相關指標在多重比較校正後未能維持顯著差異。研究結果顯示，大學入學英語考試比官方高中英語教科書更依賴具有高度構詞生產力的詞彙形成歷程。本研究指出，教學教材與評量要求之間存在結構性的落差。研究結果強調，在英語課程中納入更豐富的形態輸入，並加強對構詞歷程的明確教學，有助於提升學生面對高風險考試的準備，並促進其順利銜接高等教育。

Keywords Corpus linguistics · Morphological complexity · Turkey · Textbooks · University entrance exams

關鍵詞 語料庫語言學 · 形態複雜度 · 土耳其教科書 · 大學入學考試

Introduction

English language education in Turkey has long been a subject of intense discussion, reflecting the importance placed on English as a global language and its critical role in academic and professional success. Over the years, numerous reforms have been introduced to improve the quality of English instruction in schools, driven by the growing recognition of the need to equip students with competitive language skills. One of the most significant changes in recent years has been the integration of English instruction into earlier grades, beginning from the 2nd grade, a shift from its earlier introduction in the 4th grade. Accompanying this structural change was the reorganization of the educational system into a 4 + 4 + 4 model, requiring adjustments in curriculum design and teaching strategies (Milli Eğitim Bakanlığı, 2018). However, Hatipoğlu (2016) argues that these changes have not always followed empirical research.

In response to these changes, the national English language curriculum has undergone revisions to better align with contemporary language teaching methodologies and international standards such as the Common European Framework of Reference for Languages (CEFR). The curriculum emphasizes the importance of creating authentic learning environments and integrating communicative language teaching principles, encouraging learners to develop their listening, speaking, reading, and writing skills in meaningful contexts. The stated goal is to foster independent, proficient, and confident language users, capable of navigating real-world communicative situations effectively (Milli Eğitim Bakanlığı, 2018).

Despite these curricular reforms and the emphasis on communicative competence, assessment practices in Turkey continue to be dominated by high-stakes university entrance exams. These exams not only determine access to higher education but also exert a powerful backwash effect on classroom practices. Hatipoğlu (2016) mentions that Turkey has “one important high-stakes exam, which determines whether students gain entry to prestigious colleges or tertiary institutions” (p. 2). These high-stakes English university entrance exams are important because they determine what programs language major students can enroll in, at which university in Turkey on a ranking and score-based system. The English university entrance exam consists of 80 multiple-choice questions

to be completed within 120 min. The exam is designed in a classical multiple-choice format and includes 12 distinct question types. Specifically, the distribution of questions is as follows: vocabulary knowledge (5 questions, including 1 on phrasal verbs), grammar (10 questions, with 4 on tenses, 2 on prepositions, 3 on conjunctions, and 1 on quantifiers), cloze tests (5 questions), sentence completion (8 questions), finding the Turkish equivalent of an English sentence (6 questions), finding the English equivalent of a Turkish sentence (6 questions), reading comprehension (15 questions), identifying the sentence closest in meaning (5 questions), identifying the sentence that maintains coherence in a paragraph (5 questions), choosing the appropriate expression for a given situation (5 questions), dialogue completion (5 questions), and detecting the sentence that disrupts coherence in a paragraph (5 questions).

Central to achieving these objectives of the ministry and helping students succeed in the English university entrance exam is the role of English high school textbooks, which remain the most ubiquitous and influential teaching materials in Turkish English as a foreign language (EFL) classrooms. Textbooks serve as both a guide for instructional practice and a key source of language exposure, especially because many students may lack the means to purchase additional materials for socioeconomic or other reasons.

Despite the significance textbooks hold, concerns have been raised regarding the extent to which these materials align with the linguistic demands of high-stakes English university entrance exams from a linguistic perspective. Research in various contexts has highlighted discrepancies between English textbooks and standardized assessments, noting a lack of alignment in linguistic focus and skill development (Nur & Islam, 2018; Tai & Chen, 2015; Underwood, 2010).

In Turkey, this gap is particularly evident, although few studies have examined this issue directly. Previous studies suggest that English language education in Turkey is characterized by a sharp divide between the linguistic expectations of high school curricula and English university entrance exams (Gedik, 2021; Gedik & Kolsal, 2022). High school English instruction, guided by national curricula prepared by the Ministry of National Education (MoNE), often prioritizes basic grammar and vocabulary, catering to general communicative competence as Gedik and Kolsal (2022) have demonstrated. However, university entrance exams are designed to assess higher-order linguistic skills, requiring students to navigate more complex lexical and grammatical structures. Previous studies indicate that there is a linguistic disconnect between the English high school textbooks and the English university entrance exams in terms of lexical and syntactic demands (Gedik & Kolsal, 2022), and that primary, middle, and high school English textbooks do not follow an incremental increase of the quality (diversity) and quantity (frequency) of input of selected constructions that are taught in these books (Gedik, 2021).

This disconnect may create significant challenges for students transitioning from secondary to tertiary education, particularly those from under-resourced schools who may lack access to supplemental learning materials. One way to conceptualize the consequences of this misalignment is through the notion of the negative backwash effect (Prodromou, 1995). Backwash describes how assessment practices influence teaching and learning. When there is a strong alignment, testing can promote

positive learning outcomes; however, when instruction and assessment diverge, a negative backwash effect emerges. In the Turkish context, where high school textbooks emphasize relatively simple and communicative language while entrance exams demand complex and academically oriented input, this misalignment risks discouraging students and reinforcing inequities. This linguistic disconnect between instructional materials and assessment formats presents a challenge for students and educators alike, particularly in light of the government's provision of standardized textbooks to promote educational equity (Gençoğlu, 2017).

In this study, we focus on morphological complexity in two corpora (high school textbooks and exams) that Gedik and Kolsal (2022) compiled, mainly because of the availability of a research tool to analyze this and to expand the discussion of the disconnect between the two corpora (see Gedik, 2021; Gedik & Kolsal, 2022). The role of morphological complexity in second language acquisition (SLA) and language assessment has garnered significant attention in linguistic and educational research. Morphology is foundational to language proficiency as it enables learners to process and produce both simple and complex linguistic forms. Morphological knowledge and morphological awareness are a critical skill that underpins vocabulary acquisition, grammatical accuracy, and reading comprehension (Kieffer & Lesaux, 2012; Nation, 2001), all of which are arguably relevant in high-stakes examination conditions. However, the alignment between language instruction materials and the morphological demands of high-stakes assessments, such as university entrance exams, remains underexplored, especially from a morphological perspective. Therefore, this study addresses this gap by examining the morphological complexity of English textbooks used in Turkish high schools compared to the language found in university entrance exams, with a focus on indices related to inflectional and derivational morphology, morpheme family size, and morphological density.

Morphological indices, such as inflectional and derivational measures, provide a detailed lens through which to evaluate the linguistic complexity of texts. Inflectional morphology, which involves grammatical modifications such as tense, number, and aspect, is essential for sentence-level comprehension and production (Brezina & Pallotti, 2019) and helps to understand who did what to whom. In contrast, derivational morphology, which creates new words through the addition of prefixes and suffixes, is a key driver of vocabulary expansion and academic language proficiency (Schmitt & Zimmerman, 2002). Additional indices, including morpheme family size, log frequency scores, and morphological density, offer insights into the richness and productivity of the linguistic environment. Together, these measures allow for a comprehensive analysis of the extent to which educational materials prepare learners for the linguistic demands of university entrance exams.

Automated measures have become increasingly popular for analyzing learner language production and educational materials, as they offer efficient, objective, and reproducible methods for assessing linguistic complexity. For instance, automated tools like text analysis of morphology in monolingual input (TAMMI, Tywoniw & Crossley, 2020) and other computational linguistics tools allow researchers to measure a wide range of morphological and syntactic features in learner production and

textbooks, such as morpheme frequency, inflectional and derivational complexity, and lexical diversity. These automated measures are particularly valuable in educational settings because they provide consistent and comprehensive data across large corpora, which would be challenging to achieve manually. Moreover, automated measures allow for the analysis of subtle variations in language use, which are often critical for understanding the developmental stages of language learners (Lu, 2012). This approach contributes to a more comprehensive understanding of the linguistic features that shape second language acquisition.

Therefore, this study addresses the following research question: To what extent does the morphological complexity of English high school textbooks in Turkey differ from that of English university entrance exams, and what implications do these differences have for English language education in Turkey?

Morphological Complexity in L2

Morphological complexity is central to understanding second language acquisition (SLA). Morphemes, the smallest units of meaning, can be measured in various ways to reveal differences in linguistic input and output. Our discussion here will revolve around the categories of indices selected for the analysis—which we explain in the section. There are several morphological complexity measures that can be used, which we categorize as basic morpheme counts, frequency measures, family size, and morphological complexity indices.

Basic morpheme counts provide insights into the linguistic density of texts. Measures like the average number of inflected or derivational tokens per word indicate the frequency of morphemes that modify meaning or grammatical function (Tywoniw & Crossley, 2020). For instance, in SLA, exposure to higher counts of derivational tokens, such as words with affixes like “un-“ or “-ness,” can significantly expand a learner’s lexicon and improve comprehension and production of complex vocabulary (Brezina & Pallotti, 2019). Previous studies have shown that increased exposure to inflected forms contributes to learners’ ability to internalize grammatical structures and to produce morphologically accurate sentences (Ellis, 2006).

Frequency measures, such as the average frequency or log frequency of affixes, provide an additional layer of understanding. Frequency is an important driver of language acquisition, broadly speaking (Bybee, 2010). Frequent morphemes are processed more quickly and efficiently, aiding vocabulary acquisition and comprehension (Sánchez-Gutiérrez et al., 2017). Studies have demonstrated that L2 learners often benefit from encountering high-frequency affixes, as these forms facilitate lexical access and recognition (Jiang, 2000). For example, Ellis (2013) argues frequency and frequency effects will facilitate tense-aspect learning, whereby affixes can play a role, in a second language.

Family size indices reflect the productivity of morphemes in forming word families. Larger morpheme family sizes, such as those associated with high-frequency roots or affixes, have been shown to facilitate lexical retrieval and processing (Tywoniw & Crossley, 2020). Research by Sánchez-Gutiérrez et al.

(2017) demonstrated that L2 learners exposed to morphologically productive roots and affixes acquired vocabulary more efficiently and formed stronger lexical networks. For instance, learners exposed to one root in other words were better able to decode novel words using similar morphemes. Additionally, larger family sizes have been linked to faster word recognition, which supports reading fluency and comprehension in L2 learners (Nation, 2001).

Finally, morphological complexity indices (MCIs), such as those measuring inflectional or derivational variety and type-token ratios (TTRs), capture the diversity and density of morphemes within a text. High MCI values indicate texts with rich morphological structures that challenge learners to engage deeply with the material. Previous studies have demonstrated that texts with greater morphological complexity are particularly effective at promoting advanced linguistic skills. Brezina and Pallotti (2019) found that L2 learners exposed to morphologically diverse texts demonstrated significant gains in grammatical accuracy and lexical diversity over time.

Looking at these families of indices, it is safe to assume that, if English textbooks are the only source of input that Turkish students are exposed to (especially students from disadvantaged backgrounds), then the textbooks would need to have enough morphological complexity that can adequately prepare them for real-life communicative demands, or in the case of this study, English university entrance exams. If the two do not match in terms of complexity, one could argue that students would be at a severe disadvantage, as it has been discussed previously (Gedik, 2021; Gedik & Kolsal, 2022). To our knowledge, no other studies so far have compared the morphological complexity overlap of English exams and English textbook materials in any countries.

Interestingly, only a handful of studies have systematically analyzed the morphological complexity of English textbooks, largely because such analyses are time-intensive and require substantial manual effort. Traditionally, linguistic features like affix frequency, morphological diversity, and word families were studied through manual corpus annotation, which limited the scope of such research. However, advances in computational tools like TAMMI (Tywoniw & Crossley, 2020) have made these investigations significantly more accessible. TAMMI enables researchers to extract and calculate measures related to morpheme counts, morphological variety, and family size indices with minimal manual intervention (see for a more detailed explanation).

One study exemplifying the application of corpus-based morphological analysis was conducted by Morita et al. (2019), which examined the frequency and variety of affixes in Japanese junior high school and senior high school English textbooks. Their findings revealed that senior high school textbooks provided significantly greater exposure to affixes compared to junior high school textbooks. Specifically, prefixes like “un-” and “re-” and suffixes like “-ly” and “-ment” appeared more frequently in senior high school materials, suggesting that these textbooks aim to gradually introduce learners to more complex morphological patterns. However, the study also found that many affixes, particularly those with lower frequencies, were insufficiently represented in the textbooks to foster incidental learning. The authors concluded that explicit instruction in morphological patterns is necessary to supplement the input provided by textbooks.

The present study builds on this literature by providing a detailed comparison of morphological indices across high school and exam corpora in Turkey. In light of the foregoing discussion in the introduction, the present study aims to examine if there are statistically significant differences in the selected indices of morphological complexity between the textbook and exam corpus.

Methodology

Corpora

Gedik and Kolsal (2022) compiled the two corpora. English textbooks were obtained from eba.gov.tr and English university entrance exams from ösym.gov.tr. ÖSYM, the Measurement, Selection, and Placement Center, is responsible for preparing and administering Turkey's nationwide university entrance exams and for student placements. EBA, on the other hand, is the online platform where both teachers and students can access educational content, including textbooks for K–12. The English textbooks, along with supplementary materials such as workbooks and listening transcripts, for grades 9 through 12 were identified and downloaded in PDF format. Simultaneously, English university entrance exams from the years 2010 to 2019 were also gathered and saved in PDF format. In total, eight textbooks and ten exams were analyzed.

The textbooks cover each year of high school (9th–12th grades) and were published by various publishing houses: MEB's *Relearn*, *Teenwise*, and *Progress* for the 9th grade; *Count Me In* and *Gizem* for the 10th grade; *Sunshine* and *Silverlining* for the 11th grade; and *Count Me In* for the 12th grade, each with corresponding workbooks. Despite being produced by different publishers, the textbooks follow the same guidelines set by the Ministry of National Education (MEB) and are subject to a review process by MEB-appointed committees. The CEFR levels assigned to the textbooks for each grade are as follows: A1–A2 for 9th grade, A2+–B1 for 10th grade, B1+–B2 for 11th grade, and B2+ for 12th grade. The total number of tokens in the textbook corpus was 301,255, while the ten university entrance exams, released by ÖSYM between 2010 and 2019, contained a total of 66,913 tokens.

The following steps were undertaken to compile the corpora: (a) the PDF files were converted into DOCX format using an online document converter; (b) both corpora were cleaned to remove any errors, typos, or unnecessary symbols or images resulting from the conversion process that might interfere with the analysis; (c) the cleaned DOCX files were then converted into compatible TXT files for use with the analysis tools.

Automated Text Analysis and Data Analysis

To answer our research question, we used the corpus from the previous study (freely available to download at https://osf.io/82wzx/overview?view_only=c73c10e5b5a8456ea657d21719d1e6fc), and we employed TAMMI (Tool for Automatic

Measurement of Morphological Information, Tywoniw & Crossley, 2020), to automatically analyze the morphological complexity of the high school English textbook corpus and the university entrance exam corpus in Turkey. TAMMI is specifically designed to compute a variety of morphological indices, offering a detailed analysis of both basic and complex morphemes in textual data. It processes plain text files to calculate a wide range of measures related to morpheme counts, morphological variety, and complexity. These measures include morpheme type-token ratios (TTR), morphological complexity indices (MCI), and variables drawn from the MorphoLex database (Sánchez-Gutiérrez et al., 2017), which includes morpheme frequency/length, family size counts, frequency, and hapax counts.

TAMMI can analyze multiple text files simultaneously, which allows for the efficient processing of large corpora. When given a folder of plain text files as input, TAMMI generates a comprehensive comma-separated values (.csv) file that compiles the results into a format compatible with spreadsheet software, making it easy to manipulate and visualize the data.

We analyzed the CSV output in SPSS, version 27. All variables were first assessed for the assumptions of independent-sample *t*-tests. Normality was tested using Shapiro–Wilk tests separately for the high school and exam corpora, and homogeneity of variances was assessed using Levene’s tests. Several variables violated the assumption of normality in one or both groups (Shapiro–Wilk $p < .05$), though Levene’s tests indicated that variances were generally comparable across groups. Because some variables did not meet the normality assumption, we used Welch’s *t*-tests for all comparisons, which are robust to unequal variances and sample sizes. For each variable, we report the group means and standard deviations, Welch’s *t* statistic with the associated degrees of freedom, the *p*-value, and Cohen’s *d* as a measure of effect size. We corrected for multiple comparisons ($p = .05/26$ comparisons, $\alpha = .00192$).

We selected 26 of the 42 indices TAMMI can produce (see the spreadsheet for the full list of indices here <https://docs.google.com/spreadsheets/d/1ciHNNH-oKEtXDyPxWqSA7vWPkeW8-nys8/edit?gid=1834676141#gid=1834676141>). Table 1 presents the names of the selected indices and their descriptions taken from the spreadsheet. A slightly more detailed description of each index tested here can be found in the Appendix with examples.

The indices chosen were specifically selected to capture key aspects of inflectional and derivational morphology, morpheme frequency, family size, and overall morphological complexity, which are central to understanding the linguistic demands placed upon students. For instance, indices such as Inflected_Tokens, Derivational_Tokens, and num_morphemes_per_cw were prioritized for their ability to reflect the grammatical and lexical sophistication of the texts. Indices that were more detailed or focused on specific, less central aspects of morphology, such as those measuring very specific or infrequent morphemes, were excluded as they did not significantly contribute to the core comparisons of the study.

Table 1 Indices and their descriptions

Index	Description
Inflected_Tokens	Average number of inflected tokens per content word
Derivational_Tokens	Average number of tokens with derivational affixes per content word
Tokens_w_Prefixes	Average number of tokens with prefixes per content word
Tokens_w_Affixes	Average number of tokens with affixes per content word
Compounds	Average number of compound words per content word
number_prefixes_per_cw	Average number of prefixes per content word
number_roots_per_cw	Average number of roots per content word
number_suffixes_per_cw	Average number of suffixes per content word
number_affixes_per_cw	Average number of affixes per content word
num_morphemes_per_lemma	Average number of morphemes per word lemma
num_morphemes_per_cw	Average number of morphemes per content word
affix_freq_per_cw	Average frequency of affixes per content word
affix_log_freq_per_cw	Average log frequency of affixes per content word
prefix_freq_per_cw	Average frequency of prefixes per content word
prefix_log_freq_per_cw	Average log frequency of prefixes per content word
suffix_freq_per_cw	Average frequency of suffixes per content word
suffix_log_freq_per_cw	Average log frequency of suffixes per content word
affix_family_size_per_cw	Average family size for affixes per content word
prefix_family_size_per_cw	Average family size for prefixes per content word
suffix_family_size_per_cw	Average family size for suffixes per content word
mean_subset_inflectional_variety_10	Average inflectional variety
inflectional_TTR_10	Inflectional type-token ratio (TTR)
inflectional_MCI_10	Inflectional Mean Complexity Index (MCI)
mean_subset_derivational_variety_10	Average derivational variety
derivational_TTR_10	Derivational type-token ratio (TTR)
derivational_MCI_10	Derivational Mean Complexity Index (MCI)

Results

We ran Welch's independent-sample *t*-tests in SPSS, version 27, to investigate whether the two corpora differed across the selected morphological indices and obtained Cohen's *d* values (see Table 2 below). The Cohen's *d* values indicate the magnitude of the differences between the two corpora.

After applying a Bonferroni correction ($\alpha=.00192$), a subset of the observed effects remained statistically robust (see Table 2 below, 12 of the 26 remained significant). Specifically, Derivational_Tokens, Tokens_w_Prefixes, number_prefixes_per_cw, number_roots_per_cw, number_affixes_per_cw, num_morphemes_per_lemma, prefix_family_size_per_cw, prefix_freq_per_cw, prefix_log_freq_per_cw, affix_log_freq_per_cw, mean_subset_derivational_variety_10, and derivational_MCI_10 all remained significant. In contrast, the remaining measures, including

Table 2 Means, standard deviations, group comparisons, and Cohen's *D* for the selected morphological indices

Variable	Exam M	Exam SD	High school M	High school SD	t	df	p	Cohen's <i>d</i>	Significance after correction
Inflected_Tokens	0.33	0.02	0.30	0.03	-2.04	11.79	0.064	1.01	
Derivational_Tokens	0.27	0.04	0.21	0.03	-4.04	15.85	0.00097	1.88	*
Tokens_w_Prefixes	0.08	0.02	0.04	0.02	-5.20	15.08	0.00011	2.47	*
Tokens_w_Affixes	0.22	0.03	0.18	0.03	-3.09	15.81	0.00711	1.44	
Compounds	0.02	0.01	0.03	0.00	1.64	15.13	0.122	0.74	
number_prefixes_per_cw	0.08	0.02	0.04	0.02	-5.23	15.05	0.0001	2.49	*
number_roots_per_cw	0.97	0.01	0.94	0.01	-4.13	10.85	0.00173	2.07	*
number_suffixes_per_cw	0.26	0.03	0.21	0.03	-3.05	14.92	0.0081	1.45	
number_affixes_per_cw	0.34	0.04	0.25	0.05	-3.90	14.55	0.00149	1.87	*
num_morphemes_per_lemma	1.30	0.05	1.19	0.06	-4.29	13.74	0.00077	2.08	*
num_morphemes_per_cw	1.63	0.07	1.50	0.08	-3.68	13.26	0.00268	1.79	
prefix_family_size_per_cw	23.15	4.69	11.77	4.66	-5.13	15.19	0.00012	2.43	*
suffix_family_size_per_cw	305.84	41.26	245.10	37.61	-3.26	15.67	0.00502	1.53	
affix_family_size_per_cw	328.99	44.60	256.87	41.44	-3.55	15.58	0.00279	1.67	
prefix_freq_per_cw	83,120.82	18,169.89	40,278.59	15,548.29	-5.39	15.89	0.000062	2.51	*
prefix_log_freq_per_cw	0.46	0.09	0.24	0.09	-5.23	15.05	0.0001	2.49	*
suffix_freq_per_cw	710,857.28	103,064.29	564,348.32	108,380.79	-2.91	14.78	0.01085	1.39	
suffix_log_freq_per_cw	1.59	0.20	1.29	0.21	-3.04	14.83	0.0084	1.45	
affix_freq_per_cw	793,978.10	117,893.80	604,626.91	120,630.46	-3.34	14.98	0.00446	1.59	
affix_log_freq_per_cw	2.05	0.27	1.52	0.30	-3.84	14.55	0.00168	1.84	*
mean_subset_inflectional_variety_10	2.87	0.10	2.68	0.24	-2.18	8.82	0.05754	1.13	
inflectional_TTR_10	0.19	0.01	0.17	0.02	-2.26	8.74	0.05069	1.17	
inflectional_MCI_10	-0.96	0.02	-0.97	0.01	-1.90	15.85	0.07633	0.87	

Table 2 (continued)

Variable	Exam M	Exam SD	High school M	High school SD	t	df	p	Cohen's <i>d</i>	Significance after correction
mean_subset_derivational_variety_10	3.02	0.40	2.26	0.43	-3.85	14.55	0.00167	1.84	*
derivational_TTR_10	0.85	0.05	0.79	0.06	-2.56	14.15	0.02257	1.23	
derivational_MCI_10	-0.90	0.03	-0.94	0.02	-4.07	15.02	0.00101	1.82	*

Note. Asterisks indicate $p < .00192$; exact p -values are reported because several corrected effects are close to the adjusted threshold.

Inflected_Tokens, Tokens_w_Affixes, Compounds, number_suffixes_per_cw, num_morphemes_per_cw, suffix_family_size_per_cw, affix_family_size_per_cw, affix_freq_per_cw, suffix_freq_per_cw, suffix_log_freq_per_cw, and all inflectional complexity indices, as well as derivational_TTR_10, did not survive correction, despite many showing large effect sizes. This pattern indicates that the most robust differences between corpora are concentrated in derivational morphology and prefix-related measures, whereas other effects were less robust or more variable under stringent statistical control.

Basic Morpheme Counts

The exam corpus demonstrated greater use of morphology than the high school corpus across several basic morpheme indices, although not all differences survived correction for multiple comparisons. The Inflected_Tokens index was higher in the exam corpus ($M=0.33$, $SD=0.02$) compared to the high school corpus ($M=0.30$, $SD=0.03$), $t(11.79)=-2.04$, $p=.064$, $d=1.01$, but this difference was not statistically significant.

By contrast, Derivational_Tokens were significantly more frequent in the exam corpus ($M=0.27$, $SD=0.04$) than in the high school corpus ($M=0.21$, $SD=0.03$), $t(15.85)=-4.04$, $p<.001$, $d=1.88$, and this effect survived correction. Similarly, Tokens_w_Prefixes showed a robust difference, with the exam corpus ($M=0.08$, $SD=0.02$) containing more prefixed tokens than the high school corpus ($M=0.04$, $SD=0.02$), $t(15.08)=-5.20$, $p<.001$, $d=2.47$, also surviving correction. Tokens_w_Affixes followed the same trend ($M=0.22$, $SD=0.03$ vs. $M=0.18$, $SD=0.03$), $t(15.81)=-3.09$, $p=.007$, $d=1.44$, but did not remain significant after correction.

The Compounds index did not differ significantly between corpora, $t(15.13)=1.64$, $p=.122$, $d=0.74$. By contrast, number_prefixes_per_cw (exam: $M=0.08$, $SD=0.02$; HS: $M=0.04$, $SD=0.02$) showed a robust difference, $t(15.05)=-5.23$, $p<.001$, $d=2.49$, which survived correction. The number_roots_per_cw was also higher in the exam corpus ($M=0.97$, $SD=0.01$) compared to the high school corpus ($M=0.94$, $SD=0.01$), $t(10.85)=-4.13$, $p=.001$, $d=2.07$, and remained significant after correction.

Similarly, number_affixes_per_cw ($t(14.55)=-3.90$, $p=.001$, $d=1.87$) and num_morphemes_per_lemma ($t(13.74)=-4.29$, $p<.001$, $d=2.08$) remained significant after correction. In contrast, number_suffixes_per_cw ($p=.008$) and num_morphemes_per_cw ($p=.003$) did not remain significant after correction despite large effect sizes.

Frequency Measures

Differences across frequency-based indices were also observed, though only a subset survived correction. Prefix_freq_per_cw showed a robust effect, with higher values in the exam corpus, $t(15.90)=-5.39$, $p<.001$, $d=2.51$, which remained

significant after correction. Similarly, *prefix_log_freq_per_cw* also survived correction, $t(15.05) = -5.23$, $p < .001$, $d = 2.49$. *Affix_log_freq_per_cw* showed a significant difference ($t(14.55) = -3.84$, $p = .001$, $d = 1.84$), which also survived correction. However, other frequency measures, including *affix_freq_per_cw* ($p = .004$), *suffix_freq_per_cw* ($p = .010$), and *suffix_log_freq_per_cw* ($p = .008$), did not remain significant after correction despite large effect sizes. These findings suggest that the most robust frequency differences are driven by prefix-related distributions and overall affix log frequency.

Family Size Indices

Indices reflecting morphological productivity also revealed strong group differences. *Prefix_family_size_per_cw* showed a robust effect, $t(15.19) = -5.13$, $p < .001$, $d = 2.43$, which survived correction. Similarly, *suffix_family_size_per_cw* ($t(15.67) = -3.26$, $p = .005$, $d = 1.53$) did not survive correction. *Affix_family_size_per_cw* ($t(15.58) = -3.55$, $p = .003$, $d = 1.67$) also did not remain significant after correction. Thus, only prefix family size showed a fully robust effect, indicating that productivity differences are particularly pronounced in prefixal morphology.

Morphological Complexity

Indices of morphological complexity showed a mixed pattern. *Inflectional_TTR_10* ($p = .051$) and *mean_subset_inflectional_variety_10* ($p = .057$) did not reach significance. Similarly, *inflectional_MCI_10* ($p = .076$) was not significant. By contrast, derivational measures showed stronger effects. *Derivational_MCI_10* ($t(15.02) = -4.07$, $p = .001$, $d = 1.82$) and *mean_subset_derivational_variety_10* ($t(14.55) = -3.85$, $p = .002$, $d = 1.84$) remained significant after correction. However, *derivational_TTR_10* ($p = .023$) did not survive correction.

Discussion

The findings indicate differences between high school English textbooks and English university entrance exams in Turkey across multiple morphological indices. Many effects were large to very large, although only a subset remained significant after correction. It is also important to acknowledge that high school textbooks and university entrance exams are designed with different aims. For instance, textbooks aim to provide general English competence and communicative skills for a wide student population, while university entrance exams are selective tools designed to identify top-performing candidates through linguistically demanding items. Therefore, differences in complexity are not inherently problematic. However, the scale and systematic patterns, particularly concentrated in derivational and prefix-related measures, suggest that the current curriculum does not adequately scaffold students toward the advanced language needed for these high-stakes exams.

Several limitations should be noted. The present study is limited by the relatively small number of texts included in each corpus (eight textbooks and ten exams), which constrains the generalizability of the findings to a specific time frame in a specific country. The analyses were conducted at the text level, and therefore, the statistical comparisons reflect variation across a limited sample of materials rather than across a broader population of educational texts. As a result, the findings should be interpreted as evidence of systematic differences within the sampled corpora, rather than definitive claims about all high school textbooks or all university entrance exams in Turkey (especially since after 2019). At the same time, it is important to note that the observed effects were consistent across multiple indices and often associated with large effect sizes, particularly for derivational and prefix-related measures. This convergence suggests that the differences identified are unlikely to be driven solely by idiosyncratic properties of individual texts. Nevertheless, future research should aim to replicate these findings using larger and more diverse corpora, potentially including additional textbook series, exam years, and proficiency levels, in order to establish the robustness and generalizability of the patterns reported here.

Some of the statistically significant differences can be interpreted in practical terms. For instance, the Derivational_Tokens index indicates that, for every 100 words, the exam corpus includes 27 words with derivational affixes, such as “un-” or “-ness,” whereas the high school corpus contains only 21 words. This highlights the greater lexical richness and complexity of the exam corpus. Moreover, higher prefix-related measures (e.g., number_prefixes_per_cw) and some affix-based indices in the exam corpus indicate more frequent use of affixes, which play a significant role in building word complexity. The Tokens_w_Prefixes index further shows that exam texts use more words with prefixes (approximately 8 per 100 words) than high school texts (around 4 per 100 words), emphasizing the greater morphological diversity. These findings primarily indicate that exam texts are more morphologically dense in derivational and prefix-related structure, reflecting higher linguistic demands on learners, and are consistent with the disconnect between high school and exam corpora reported by Gedik and Kolsal (2022).

These quantitative differences can be seen in the actual word forms present in the two corpora. The exam corpus frequently included morphologically dense items such as *unpredictability*, *reinforcement*, and *globalization*, each combining multiple derivational morphemes. In contrast, the textbook corpus contained simpler derivations such as *happiness* or *careful*. Similarly, exam texts made use of layered inflectional forms such as *had been developed* or *were studying*, while textbooks relied on simpler inflections like *walked* or *is playing*. Prefixal diversity was also greater in the exams, with forms like *reconsider*, *inaccurate*, and *misunderstanding*, compared to more transparent forms like *unhappy* in textbooks. A similar contrast was found in suffix usage: exam texts contained complex items like *environmental* and *nationality*, whereas textbooks more often featured everyday suffixes such as *joyful* or *teacher*. Finally, in terms of morpheme density, exam texts included items such as *unbelievably* (four morphemes) while textbooks typically used less complex forms like *running* (two morphemes).

The statistically significant differences in derivational morphology between the two corpora are particularly striking and hold significant implications for SLA. The exam corpus's higher scores for derivational complexity measures (particularly MCI and derivational variety) suggest that these texts require learners to decode and comprehend morphologically complex words at a far more advanced level than what is emphasized in high school textbooks. This gap is problematic for second language learners, as studies have shown that deriving meaning from morphologically complex words is a cognitively demanding task that relies on both explicit knowledge of derivational rules and implicit sensitivity to morphological patterns (e.g., Schmitt & Zimmerman, 2002). Without sufficient exposure to derivational forms during earlier stages of language learning, students may struggle to meet the lexical demands of high-stakes assessments.

The findings on selected morpheme family size and frequency, particularly prefix-related measures, further highlight the potential challenges that L2 learners may face, as has been documented in previous studies (e.g., Gedik & Kolsal, 2022; Underwood, 2010, although this effect was not significant for suffixes). Higher values in several family size and log frequency measures in the exam corpus indicate that these texts rely on morphologically productive and high-frequency forms, which native speakers may process more efficiently but which may pose challenges for L2 learners. Morphological family size has been shown to facilitate word recognition and lexical retrieval in native speakers (Sánchez-Gutiérrez et al., 2017), but second language learners often rely on less efficient processing strategies due to incomplete morphological knowledge and slower lexical access (Jiang, 2000). The limited exposure to such morphologically productive forms in high school materials may hinder learners' ability to generalize and apply morphological rules, resulting in weaker vocabulary acquisition and slower reading comprehension. However, note again that this only applied to prefixes in the current data set and statistical comparisons (see corrections for multiple comparisons in the “Results” section).

The observed differences between the two corpora may also have critical implications for SLA. The higher number of morphemes per lemma (and to a lesser extent per content word) in the exam corpus reflects a reliance on linguistically dense constructions that pack more information into fewer words. While such constructions are characteristic of academic texts, they are particularly challenging for second language learners, who may lack the proficiency needed to unpack these dense structures (Nation, 2001). For learners, decoding morphologically dense language involves understanding individual morphemes and potentially integrating them into larger syntactic and semantic contexts. This arguably requires advanced morphological and syntactic processing skills, which may not be adequately developed through the textbooks available in the Turkish high school context.

Pedagogical Implications

The mismatch between high school textbooks and university entrance exams can be interpreted through the lens of the negative backwash effect (Prodromou, 1995).

Backwash refers to the influence of testing on teaching and learning, and it becomes negative when instructional content does not align with assessment demands. The present findings suggest that this misalignment is not uniform across all aspects of morphology (see statistical differences after correction).

From a pedagogical perspective, this pattern is important. While high school textbooks appear to provide adequate exposure to basic morphological forms, they may not sufficiently prepare students for the derivationally rich and morphologically productive language that characterizes university entrance exams. In particular, the underrepresentation of derivational processes, such as affixation that expands lexical meaning, and the more limited use of prefixes in instructional materials may restrict learners' ability to decompose and interpret complex words under exam conditions.

This gap has direct implications for second language learning. Derivational morphology plays a central role in vocabulary growth, enabling learners to recognize relationships between words (e.g., *happy*, *unhappy*, *happiness*) and to infer meaning in unfamiliar lexical items (e.g., Kieffer & Lesaux, 2012). The stronger and more consistent differences observed in derivational complexity measures suggest that exam texts require learners to engage in more advanced morphological processing than what is systematically practiced in high school materials. Without sufficient exposure to such structures, learners may rely on surface-level strategies, which are less effective when encountering dense academic language.

Importantly, the results also indicate that not all morphological domains are equally implicated. Differences in inflectional morphology and some frequency-based measures did not remain robust after correction. This distinction is crucial: it may imply that improving alignment between textbooks and exams does not require a wholesale increase in linguistic complexity, but rather a targeted enrichment of derivational and prefixal morphology within instructional materials if the aim of these materials is to prepare students for the university entrance exam.

The educational consequences of this mismatch may be particularly pronounced for students with limited access to supplementary learning resources. Learners who rely primarily on standard textbooks may not encounter sufficient examples of morphologically complex words, placing them at a disadvantage when facing exam texts that assume this knowledge. In contrast, students with access to additional materials or private instruction may develop greater familiarity with such structures, thereby widening existing educational inequalities.

To address this issue, high school English curricula should incorporate more systematic exposure to derivational morphology and prefix-based word formation, both in reading materials and in explicit instruction. This can include integrating morphologically rich texts, designing activities that encourage analysis of word structure, and promoting awareness of how affixes modify meaning. Research in SLA has consistently shown that such targeted interventions can support vocabulary development and reading comprehension (Kieffer & Lesaux, 2012), particularly when learners are guided to actively engage with morphological patterns.

More broadly, improving alignment between instructional materials and assessment demands requires greater coordination between textbook developers, curriculum designers, and exam writers. Rather than increasing overall difficulty, the goal should be to ensure that learners are gradually introduced to the types of

morphological structures they will encounter in high-stakes contexts. Such alignment would not only improve exam preparedness but also support the development of more flexible and efficient language processing skills, which are essential for academic success.

Turkey's university entrance exam system, with its competition and emphasis on language proficiency, creates significant pressure for students to perform well. However, this study reveals a gap between the language skills fostered by the high school curriculum and the advanced language requirements of the exam. For students from urban, well-resourced schools, the disparity may be somewhat mitigated through access to private tutoring, additional learning materials, or immersion opportunities. However, for students in rural or under-resourced schools, the gap in morphological complexity, in its broad sense, is likely to be far more challenging to overcome (see Şahan and Sahan (2021) for a discussion on under-resourced schools in Turkey).

This disparity reflects broader structural inequalities in the Turkish education system, where socioeconomic status often determines the resources available for language learning. Students in underfunded schools are less likely to have access to extracurricular language programs, advanced English materials, or teachers trained to provide explicit instruction in morphology and academic language. This systemic inequality possibly exacerbates the disadvantages faced by these students, making it harder for them to achieve the high language proficiency required to succeed in university entrance exams.

Finally, incorporating explicit instruction in morphological analysis into the curriculum may also be critical. By teaching students how to recognize and manipulate morphemes, educators can equip learners with strategies for decoding complex words, thereby improving their vocabulary breadth and reading comprehension. This approach has been shown to be particularly effective in SLA (Kieffer & Lesaux, 2012), where explicit focus on morphology helps learners internalize patterns that they may otherwise struggle to infer from context alone.

Broader Implications Beyond Turkey

Although the present study focuses on English language education in Turkey, its implications can be extended to other countries. Several educational systems worldwide operate under similar conditions to that of Turkey (e.g., Nur & Islam, 2018; Tai & Chen, 2015; Underwood, 2010): centrally designed curricula, state-approved textbooks, and high-stakes university entrance examinations that exert strong backwash effects on classroom instruction. In such contexts, misalignment between instructional materials and assessment demands is not unique to Turkey but represents a structural risk inherent in exam-driven education systems.

Previous research on textbook–exam correspondence has primarily emphasized vocabulary size, grammatical coverage, or reading difficulty. By contrast, the present study demonstrates that differences in inflectional and, in particular, derivational morphology can substantially increase the linguistic demands placed on learners, even when texts appear superficially similar in topic or length. This suggests that

learners in other EFL contexts may face comparable challenges if textbooks do not sufficiently expose them to morphologically rich and productive forms.

Beyond its empirical findings, the study offers a transferable methodological framework for evaluating linguistic alignment between instructional materials and assessments. The use of automated morphological indices using TAMMI (or other corpus tools) allows researchers, curriculum designers, and policymakers to systematically diagnose mismatches in linguistic complexity across corpora. This corpus approach can be readily applied to other educational settings where morphology plays a central role in academic language proficiency.

Limitations

Several limitations should be acknowledged. First, although multiple comparisons were statistically controlled, the number of indices examined increases the risk of type I and type II error, and results should therefore be interpreted with appropriate caution. Second, the analyses were conducted on a relatively small number of texts (eight textbooks and ten exams), with texts as the unit of analysis. Consequently, the findings reflect systematic differences within the sampled corpora rather than definitive claims about all high school textbooks or all university entrance examinations in Turkey. Future research using larger and more diverse corpora would help establish the robustness and generalizability of the present findings.

Conclusion

This study highlights important differences between high school English textbooks and university entrance exams in Turkey, particularly in derivational and prefix-related morphology. Exam texts consistently exhibited greater derivational complexity and prefix-related morphological richness than the textbook corpus, suggesting a mismatch between instructional materials and the linguistic demands of high-stakes assessments and expanding on previous similar findings of Gedik and Kolsal (2022). These findings emphasize the need for more systematic exposure to morphologically rich language and explicit instruction in word formation processes in order to better prepare students for tertiary-level assessment and academic language use.

Appendix

Inflected_Tokens refers to the average number of inflected tokens per content word in the corpus. Inflectional morphology involves the modification of words to indicate grammatical features such as tense, number, and case. For example, in the word “walked,” the suffix “-ed” is an inflection indicating past tense. A higher number of inflected tokens suggests greater morphological complexity in the text.

Derivational_Tokens calculates the average number of tokens with derivational affixes per content word. Derivational morphology involves the creation of new words by adding prefixes or suffixes to a base word. For example, the word “happiness” is derived from the base word “happy” by adding the suffix “-ness.” A higher count of derivational tokens indicates richer lexical diversity, as texts with more derivational affixes tend to have a greater variety of vocabulary.

Tokens_w_Prefixes measures the average number of tokens with prefixes per content word. Prefixes are morphemes added to the beginning of a root word to alter its meaning. For example, in the word “unhappy,” the prefix “un-” modifies the base word “happy” to convey the opposite meaning. A higher value for this index indicates that the text relies on prefixation to generate new words.

Tokens_w_Affixes calculates the average number of tokens with affixes (which include both prefixes and suffixes) per content word. Affixes are morphemes that are attached to a root word to change its meaning or grammatical function. For instance, “careful” is formed by adding the suffix “-ful” to the base word “care.” A higher count of tokens with affixes suggests a more morphologically complex text.

Compounds refers to the average number of compound words per content word. Compound words are formed by combining two or more base words, such as “bookcase” or “toothbrush.” A greater number of compound words in a text suggests a more advanced lexical structure and reflects the text’s ability to express complex concepts concisely.

number_prefixes_per_cw calculates the average number of prefixes per content word. This index gives insight into the frequency with which prefixes are used in the text. Prefixes often modify the meaning of the base word and can indicate various semantic changes. For example, in the word “replay,” the prefix “re-” suggests the action is done again.

number_roots_per_cw measures the average number of roots per content word. A root is the base form of a word to which affixes can be added. This index helps assess the lexical richness of a text, as texts with many roots typically have a diverse vocabulary. For example, the root of “unhappiness” is “happy,” and the addition of the prefix “un-” and the suffix “-ness” results in a derived form.

number_suffixes_per_cw measures the average number of suffixes per content word. Suffixes are morphemes that are added to the end of a base word. For example, “joyful” is derived by adding the suffix “-ful” to the root word “joy.” A higher number of suffixes per content word indicates a greater use of derivational processes and morphological complexity.

number_affixes_per_cw calculates the average number of affixes (both prefixes and suffixes) per content word. This index reflects the overall morphological complexity of the text, as it accounts for both the prefixes and suffixes added to base words. More affixes typically correlate with higher levels of morphological richness and lexical diversity in the text.

num_morphemes_per_lemma refers to the average number of morphemes per word lemma. A lemma is the base form of a word, used to represent all its inflected forms. The number of morphemes per lemma indicates the degree of morphological complexity in the text. For example, the lemma “run” has one morpheme, while “running” has two morphemes: “run” (the root) and “-ing” (the inflectional suffix).

num_morphemes_per_cw calculates the average number of morphemes per content word. A higher number of morphemes per content word suggests greater linguistic complexity, as it indicates a greater use of both derivational and inflectional processes to build words. For example, in the word “unbelievably,” there are four morphemes: “un-,” “believe,” “-able”, and “-ly.”

affix_freq_per_cw refers to the average frequency of affixes per content word. This measure reflects how often affixes are used in the text, highlighting the degree to which the text relies on affixation to create new words or modify existing ones.

affix_log_freq_per_cw calculates the average log frequency of affixes per content word. This index gives a logarithmic scale for the frequency of affixes in the text, allowing for a more nuanced understanding of how frequently affixed words appear. A higher log frequency suggests that affixed words are more common in the text.

prefix_freq_per_cw refers to the average frequency of prefixes per content word. This measure reflects how often prefixes are used in the text, highlighting the degree to which the text relies on prefixation to create new words or modify existing ones.

prefix_log_freq_per_cw measures the average log frequency of prefixes per content word. Similar to the previous measure, this index evaluates the frequency of words with prefixes, with a logarithmic adjustment to account for distribution. A higher value here suggests a greater reliance on prefixation in the text.

suffix_freq_per_cw refers to the average frequency of suffixes per content word. This measure reflects how often affixes are used in the text, highlighting the degree to which the text relies on suffixation to create new words or modify existing ones.

suffix_log_freq_per_cw calculates the average log frequency of suffixes per content word. This measure, like the prefix log frequency, helps quantify the prominence of suffixes in the text. A higher log frequency of suffixes indicates a text with more frequent use of suffixation.

affix_family_size_per_cw refers to the average family size for affixes per content word. Family size measures the number of distinct words that can be formed from a particular morpheme. A larger family size indicates that a morpheme is more productive and can generate a wide range of derived words.

prefix_family_size_per_cw calculates the average family size for prefixes per content word. This index helps assess how productive prefixes are in the text, with higher family sizes indicating that the prefixes used in the text are able to form many derived words.

suffix_family_size_per_cw measures the average family size for suffixes per content word. Like prefix family size, this index indicates the productivity of suffixes in the text. A higher family size for suffixes suggests that suffixes are frequently used to generate a variety of word forms.

mean_subset_inflectional_variety_10 refers to the average inflectional variety in the text. This index provides an overview of the diversity of inflectional forms used in the text, with a higher value indicating a greater variety of inflectional morphemes and a more complex grammatical structure.

inflectional_TTR_10 measures the Type-Token Ratio for inflectional forms. This index helps assess the range of inflectional morphemes used in the text, with higher values indicating more varied use of inflections.

inflectional_MCI_10 calculates the Inflectional Mean Complexity Index. This measure provides an overall assessment of the complexity of the inflectional system in the text. A higher MCI value suggests more complex and varied inflectional morphology.

mean_subset_derivational_variety_10 measures the average derivational variety in the text. This index indicates how varied the derivational processes are within the corpus, with higher values reflecting greater morphological diversity and lexical creativity.

derivational_TTR_10 assesses the Type-Token Ratio for derivational forms. This measure evaluates the diversity of derivational morphemes used in the text, with a higher TTR indicating greater lexical richness.

derivational_MCI_10 calculates the Derivational Mean Complexity Index. Similar to the inflectional MCI, this index provides a measure of the overall complexity of derivational morphology in the text. Higher values suggest more intricate and varied derivational processes.

Author Contribution Not applicable.

Funding This study was not funded.

Data Availability All data and results can be accessed at this link: https://osf.io/82wzx/overview?view_only=c73c10e5b5a8456ea657d21719d1e6fc.

Declarations

Ethical Approval Not applicable.

Competing Interests The authors declare no competing interests.

References

- Brezina, V., & Pallotti, G. (2019). Morphological complexity in written L2 texts. *Second Language Research*, 35(1), 99–119.
- Bybee, J. (2010). *Language, usage and cognition*. Cambridge University Press.
- Ellis, N. (2006). Language acquisition as rational contingency learning. *Applied Linguistics*, 27(1), 1–24. <https://doi.org/10.1093/applin/ami038>
- Ellis, N. (2013). Construction grammar and second language acquisition. In T. Hoffmann, & G. Trousdale (Eds.), *The Oxford handbook of construction grammar* (pp. 365–378). Oxford University Press. <https://doi.org/10.1093/oxfordhb/9780195396683.013.0020>
- Gedik, T. A. (2021). An analysis of lexicogrammatical development in English textbooks in Turkey: A usage-based construction grammar approach. *Explorations in English Language and Linguistics*, 9(1), 26–55. <https://doi.org/10.2478/exell-2022-0002>
- Gedik, T. A., & Kolsal, Y. S. (2022). A corpus-based approach to English university entrance exams and English high school textbooks in Turkey. *TAPSLA*, 8(1), 157–176. <https://doi.org/10.31261/TAPSLA.9152>
- Gençoğlu, C. (2017). *Republic of Turkey Ministry of National Education* [PowerPoint slides]. COMCEC. <https://www.comcec.org/wp-content/uploads/2021/07/10-POV-PRE-TUR.pdf>
- Hatipoğlu, Ç. (2016). The impact of the university entrance exam on EFL education in Turkey: Pre-service English language teachers' perspective. *Procedia - Social and Behavioral Sciences*, 232, 136–144. <https://doi.org/10.1016/j.sbspro.2016.10.038>

- Jiang, N. (2000). Lexical representation and development in a second language. *Applied Linguistics*, 21(1), 47–77. <https://doi.org/10.1093/applin/21.1.47>
- Kieffer, M. J., & Lesaux, N. K. (2012). Direct and indirect roles of morphological awareness in the English reading comprehension of native English, Spanish, Filipino, and Vietnamese speakers. *Learning and Individual Differences*, 22(6), 768–779. <https://doi.org/10.1111/j.1467-9922.2012.00722.x>
- Lu, X. (2012). The relationship of lexical richness to the quality of ESL learners' oral narratives. *Modern Language Journal*, 96(2), 190–208. <https://doi.org/10.1111/j.1540-4781.2011.01232.1.x>
- Millî Eğitim Bakanlığı. (2018). Ortaöğretim İngilizce Dersi Öğretim Programı. Retrieved April 1, 2026, from <http://mufredat.meb.gov.tr/ProgramDetay.aspx?PID=342>
- Morita, M., Uchida, S., & Takahashi, Y. (2019). The frequency of affixes and affixed words in Japanese senior high school English textbooks: A corpus study. *International Journal of English Linguistics*, 9(1), 81–97. https://doi.org/10.20581/arele.32.0_81
- Nation, I. S. P. (2001). *Learning vocabulary in another language*. Cambridge University Press.
- Nur, S., & Islam, M. (2018). The (dis)connection between secondary English education assessment policy and practice: Insights from Bangladesh. *International Journal of English Language Education*, 6(1), 100–132.
- Prodomou, L. (1995). The backwash effect: From testing to teaching. *ELT Journal*, 49(1), 13–25. <https://doi.org/10.1093/elt/49.1.13>
- Şahan, Ö., & Sahan, K. (2021). What challenges do novice EFL teachers face in under-resourced contexts in Turkey? An exploratory study. In K. M. Bailey, & D. Christian (Eds.), *Research on teaching and learning English in under-resourced contexts* (pp. 87–100). Routledge.
- Sánchez-Gutiérrez, C. H., Mailhot, H., Deacon, S. H., & Wilson, M. A. (2017). MorphoLex: A derivational morphological database for 70,000 English words. *Behavior Research Methods*, 50, 1568–1580. <https://doi.org/10.3758/s13428-017-0981-8>
- Schmitt, N., & Zimmerman, C. B. (2002). Derivative word forms: What do learners know? *TESOL Quarterly*, 36(2), 145–171. <https://doi.org/10.2307/3588328>
- Tai, S., & Chen, H.-J. (2015). Are teachers test-oriented? A comparative corpus-based analysis of the English entrance exam and junior high school English textbooks. In F. Helm, L. Bradley, M. Guarda, & S. Thouéšny (Eds.), *Critical CALL – Proceedings of the 2015 EUROCALL conference, Padova, Italy* (pp. 518–522). Research-publishing.net. <https://doi.org/10.14705/rpnet.2015.000386>
- Tywniwi, R., & Crossley, S. (2020). Morphological complexity of L2 discourse. In E. Friginal, & J. A. Hardy (Eds.), *The Routledge handbook of corpus approaches to discourse analysis* (pp. 269–297). Routledge.
- Underwood, P. (2010). A comparative analysis of MEXT English reading textbooks and Japan's National Center Test. *RELC Journal*, 41(2), 165–182. <https://doi.org/10.1177/0033688210373128>

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Springer Nature or its licensor (e.g. a society or other partner) holds exclusive rights to this article under a publishing agreement with the author(s) or other rightsholder(s); author self-archiving of the accepted manuscript version of this article is solely governed by the terms of such publishing agreement and applicable law.

Authors and Affiliations

Tan Arda Gedik^{1,2} 

✉ Tan Arda Gedik
tan.gedik@fau.de; tan.gedik@bilkent.edu.tr

¹ Language and Cognition Lab, FAU Erlangen-Nürnberg, Bismarckstr. 6, 91054 Erlangen, Germany

² Department of Psychology, Bilkent University, 06800 Ankara, Turkey